

Björn D. Salz

Ringe, topologische Vektorräume, Algebren und die Daugavet-Eigenschaft in JB^* -Tripeln

ABSTRACT. A normed vector space X is said to have the Daugavet property if each bounded, linear map $T: X \rightarrow X$ of rank 1 satisfies $\|\text{Id} + T\| = 1 + \|T\|$. A geometric formulation of this property is well known. BECERRA GUERRERO and MARTÍN (J. Funct. Anal. 224(2) (2005) 316-337) revealed that if X , $X \neq \{0\}$, is a complex JB^* -triple or a real form of a complex JB^* -triple, X has the Daugavet property if and only if X does not possess any division tripotent. Hence a complex JBW^* -triple or a real form of it has the Daugavet property if and only if it does not possess any atom. In particular, therefore, any real or complex von Neumann algebra, if of type II or of type III, has the Daugavet property. In order to be able — even without special previous knowledge — to relate to the details of these statements and of related statements, and to have available a consistent terminology, all of the required mathematical structures will be explained systematically from scratch — such as nonassociative rings and algebras as well as triple systems. Structures over a scalar field will be treated both for the field of real and of complex numbers.

Date and Place: June 30, 2012, Berlin, Germany.

Electronic Edition. Version from July 1, 2025.

Keywords. *Primary:* JB^* -triple, triple system, minimal tripotent, division tripotent, minimal idempotent, skew minimal idempotent, nonassociative ring, C^* -algebra, von Neumann algebra, Daugavet property, rough norm, Fréchet differentiability; *Secondary:* Galois connection, annihilator system, dual system, unit ball annihilator, facear operation, radical, ring with operator domain, nonassociative algebra, Jordan algebra, topological vector space, predual, Radon-Nikodým property, semi-continuity, atom, realification, complexification, real form, involution, Banach manifold, bounded symmetric domain.

AMS Mathematics Subject Classification (2010). *Primary:* 46-02. *Secondary:* 46B20, 46Lxx, 06F25, 16W10, 17A01, 17C27, 17C65; 58Bxx, 06A15, 46B04, 46B22, 47A12.

ZUSAMMENFASSUNG. Ein normierter Vektorraum X heißt die Daugavet-Eigenschaft aufweisend, wenn jede beschränkte, lineare Rang-Eins-Abbildung $T: X \rightarrow X$ die Gleichung $\|\text{Id} + T\| = 1 + \|T\|$ erfüllt. Eine geometrische Charakterisierung dieser Eigenschaft ist wohlbekannt. BECERRA GUERRERO und MARTÍN (J. Funct. Anal. 224(2) (2005) 316-337) zeigten, dass, wenn X , $X \neq \{0\}$, ein komplexes JB^* -Tripel oder eine reelle Form eines komplexen JB^* -Tripels ist, X genau dann die Daugavet-Eigenschaft aufweist, wenn X kein divisionsminimal tripotentes Element aufweist. Somit weist ein komplexes JBW^* -Tripel oder eine reelle Form desselben genau dann die Daugavet-Eigenschaft auf, wenn es kein Atom aufweist — insbesondere weisen somit die reellen und komplexen von Neumann-Algebren, die vom Typ II oder vom Typ III sind, die Daugavet-Eigenschaft auf. Um diese und die damit im Zusammenhang stehenden Aussagen auch ohne spezielle Vorkenntnisse im Detail nachvollziehen zu können, und dabei gleichzeitig eine einheitliche Terminologie zur Verfügung stehend zu haben, werden alle dazu benötigten mathematischen Strukturen systematisch von Grund auf eingeführt und behandelt — so etwa insbesondere nichtassoziative Ringe und Algebren sowie Tripelsysteme. Strukturen über einen Skalarenkörper werden sowohl über den Körper der reellen als auch der komplexen Zahlen betrachtet.

Stichworte. *Primär:* JB^* -Tripel, Tripelsystem, minimal tripotent, divisionsminimal tripotent, schief-minimal idempotent, nichtassoziativer Ring, C^* -Algebra, von Neumann-Algebra, Daugavet-Eigenschaft, raue Norm, Fréchet-Differenzierbarkeit; *Sekundär:* Galois-Verbindung, Annihilatorsystem, Dualsystem, Kugelannihilator, Radikal, Ring mit Operatorenbereich, nichtassoziative Algebra, Jordan-Algebra, topologischer Vektorraum, Prädual, Radon-Nikodým-Eigenschaft, halbstetig, Atom, reelle Strukturierung, Komplexifizierung, reelle Form, Involution, Banach-Mannigfaltigkeit, beschränktes symmetrisches Gebiet.

Vorwort

Konstituierend für den vorliegenden Text war die von mir an der Freien Universität Berlin angefertigte Diplomarbeit *Die Daugavet-Eigenschaft in JB^* -Tripeln*, die hier hauptsächlich durch die Kapitel 4 und 5 dargestellt wird.

Auf einige Konventionen und Bezeichnungen sei besonders hingewiesen. Zwecks einer rationelleren Darstellung sind einige Aussagen, die ein identisches Textmuster aufweisen, durch entsprechende Klammersetzung zusammengefasst. Dafür werden zwei verschiedene Syntaxen verwendet:

(a) Begriffe von mehreren parallelen Fällen werden ab dem zweiten Fall das Kürzel *bzw.* für das Wort *beziehungsweise* vorangestellt, dann per Kommas getrennt in Klammern nach dem Begriff des ersten Falles gesetzt. So ist zum Beispiel die Definition *A (bzw. B, bzw. C) ist definiert als D (bzw. E, bzw. F)* zu lesen als die drei Definitionen *A ist definiert als D. B ist definiert als E. C ist definiert als F.*

(b) Hat man zwei parallele Fälle, wobei der Begriff des zweiten Falles durch Hinzunahme einiger Wörter (meist Adjektive) zu dem Begriff des ersten Falles erhalten werden kann, so wird die Aussage des zweiten Falles formuliert und die Worte werden dann so eingeklammert, dass, wenn man alle eingeklammerten Worte weglässt, man die Aussage des ersten Falles bekommt. So ist zum Beispiel die Definition *(A) B ist definiert als (C) D* zu lesen als die beiden Definitionen *A B ist definiert als C D. B ist definiert als D.*

Besonders wenn keine Missverständnisse zu befürchten sind, wird auch nur einer der Fälle explizit ausformuliert und danach nur eine Bemerkung der Form *Die anderen Fälle entsprechend* oder auch eine Formulierung mit der dafür üblichen lateinischen Floskel *mutatis mutandis*, abgekürzt *m.m.*, zu Deutsch *mit den nötigen Abänderungen*, angefügt.

Isomorphismen sind grundsätzlich alle surjektiv. Ist dementsprechend zum Beispiel X ein Vektorraum, U ein Unterraum von X und T die identische Abbildung von U in sich, so wird hier also nicht, wie bei manchen Autoren üblich, gesagt, dass T ein Isomorphismus von U nach X sei.

Ist X eine Menge, so wird die Fixpunktmenge einer Funktion f auf X mit $X_{(f)}$ bezeichnet. Diese Bezeichnungsweise wird hauptsächlich für die Involutionen auf X verwendet und auch beibehalten, wenn X eine C^* -Algebra ist. Da die Involution einer C^* -Algebra meist mit $*$ bezeichnet wird, schreibt sich somit der selbstadjungierte Teil einer C^* -Algebra X als $X_{(*)}$, wofür andere Autoren Bezeichnungen wie zum Beispiel X_{sa} , X_{ad} , X_{h} , $\text{Re}(X)$ oder $\text{Sym}(X)$ verwenden. Dies sei besonders im Hinblick darauf erwähnt, dass in der Banachraumtheorie und auch im vorliegenden Text ein Prädual eines Banachraumes X mit X_* bezeichnet wird. In der Theorie der JB^* -Tripel findet man in der Literatur oftmals eine mit τ bezeichnete Involution und die entsprechende Fixpunktmenge von X mit $X^{(\tau)}$ symbolisiert.

Topologische Vektorräume und insbesondere auch die lokal konvexen Räume werden nicht als Hausdorff'sch vorausgesetzt. Ebenso wird auch von vornherein nicht verlangt, dass Dualsysteme trennend sind.

Sowohl das neutrale Element von Gruppen als auch die Eins in Ringen und Algebren wird im Allgemeinen durchweg mit e bezeichnet; das Symbol e meint also im Allgemeinen *nicht* die Eulersche Zahl. Die Exponentialabbildung wird durchgehend mit $\exp$ bezeichnet.

Bei Strukturen über einem Skalarenkörper, wie zum Beispiel Vektorräume, Algebren, Mannigfaltigkeiten oder Tripel, wird in der Regel der Skalarenkörper spezifiziert. Dies folgt allgemein dem Muster ... *der Vektorraum über dem Körper $\mathbb{K}$* Ist dabei $\mathbb{K}$ der Körper der reellen Zahlen, so kann man dazu vollkommen synonym auch ... *der reelle Vektorraum* ... sagen, und das entsprechende gilt, falls $\mathbb{K}$ der Körper der komplexen Zahlen ist. Dabei ist zu beachten, dass C^* -Algebren, von Neumann-Algebren und gewisse Tripel ursprünglich über dem Skalarenkörper der komplexen Zahlen definiert worden sind und es verschiedene Möglichkeiten gibt, reelle Analoga zu definieren. So ist es zum Beispiel von vornherein nicht klar, was eine reelle von Neumann-Algebra oder ein reelles JB^* -Tripel sein soll. In der Literatur ist es verbreitet, dass bei diesen Strukturen immer diejenige über den Körper der komplexen Zahlen gemeint ist, wenn dieser nicht näher spezifiziert worden ist. Die je nach Autor verschieden definierten reellen Analoga dieser Strukturen werden dann von den Autoren mit dem Adjektiv *reell* versehen. Bei Jordan-Algebren ist es umgekehrt: Hier war zuerst die reelle Struktur definiert worden und dementsprechend ist es in der Literatur auch üblich, diese zu meinen, wenn der Skalarenkörper nicht mit angegeben ist.

In jeder Definition eines neuen Begriffes wird derselbe kursiv gesetzt. Insbesondere sind daran auch Definitionen als solche zu erkennen, auch wenn nicht explizit das Wort Definition verwendet wird. Dass die kursiv-Setzung auch anderweitig verwendet wird, dürfte dabei nicht zu Missverständnissen führen.

Die Abkürzung *ebd.* steht für *ebenda*, lateinisch *ibidem*.

Literaturstellen werden nach dem Muster NAME(N) [Index](Jahr) zitiert. Während dabei NAME(N) und (Jahr) in erster Linie einer unmittelbaren Orientierung dienen, welche Literaturstelle konkret gemeint ist, wird die eindeutige Zuordnung zu dem Eintrag in das Literaturverzeichnis vom [Index] geleistet. Grundsätzlich sind alle Aussagen, die nicht von mir stammen und noch nicht verbreiteten Eingang in die Buch-Literatur gefunden haben, von mir mit der dazugehörigen Zitierung der Literaturstelle versehen. Es sei aber ausdrücklich darauf hingewiesen, dass im Allgemeinen die von mir genannten Literaturstellen nicht den Anspruch erheben, originär zu sein. Vielmehr sind die Literaturstellen meist zu so verstehen, dass meine dazugehörigen Ausführungen größtenteils oder ganz auf den jeweils angegebenen Literaturstellen aufbauen und man weitere Details — wie zum Beispiel Beweise oder originäre und weiterführende Literaturstellen — dort finden kann. Insbesondere bei den etablierten Standardaussagen und Definitionen spiegeln die von mir zitierten Literaturstellen oftmals nur einige von mir ganz persönlich präferierte Literaturstellen zu der jeweiligen Thematik wieder.

Bezüglich der Sortierung im Stichwortverzeichnis sei auf folgendes hingewiesen: Umlaute werden wie die jeweils zugrunde liegenden nicht umgelauteten Selbstlaute behandelt, also $\ddot{a}$, $\ddot{o}$, $\ddot{u}$ und $\ddot{au}$ als a , o , u und au . Griechische Buchstaben und das Zeichen $*$ werden lautsprachlich behandelt. So ist zum Beispiel σ -kompakt als *sigma-kompakt* und C^* -Algebra als *CStern-Algebra* eingeordnet.

Mein besonderer Dank gebührt Herrn Professor Dr. Dirk Werner, der mir besonders in Fragen zu der Banachraumtheorie stets wertvolle Hinweise und Ratschläge geben konnte.

Berlin, den 30. Juni 2012

Björn D. Salz

E-Mail: Bjoern.Salz@FU-Berlin.de

URL: <https://www.facebook.com/Bjoern.Salz>

Inhaltsverzeichnis

1	Ordnungen und Topologie	1
1.1	Relationen und Ordnungen	1
1.2	Topologie	17
2	Algebren und topologische Vektorräume	28
2.1	Ringe	28
2.2	Vektorräume	57
2.3	Algebren	70
2.4	Topologische Vektorräume	79
2.5	Der Bipolarensatz	96
2.6	Beschränkte Mengen	104
2.7	Topologische Algebren	109
2.8	Darstellungstheorie	110
2.9	Skalarbereichsänderung von Moduln und Algebren	112
2.10	Präduale	121
2.11	Stützfunktional, Kugelsystem, numerischer Bereich und Hermitezität	124
2.12	Spektrum	133
3	Algebren mit Involutionen	138
3.1	Involutionen auf Vektorräumen	138
3.2	Involutionen auf Ringen und Algebren	142
3.3	Cayley-Dickson-Algebren	149
3.4	Hilberträume	152
3.5	C*-Algebren	160
3.6	Operator-Topologien	173
3.7	W*-Algebren	179
4	Holomorphie und Tripelsysteme	197
4.1	Analytische und differenzierbare Abbildungen	197
4.2	Analytische Banach-Mannigfaltigkeiten	204
4.3	Operation und Darstellung einer Gruppe	209
4.4	Lokal linearisierbare Abbildungen	213
4.5	Beschränkte und symmetrische Gebiete	214
4.6	Tripelsysteme	222
4.7	JB*-Tripel	229
5	Die Daugavet-Eigenschaft	260
5.1	Definition, Beispiele und Eigenschaften von Daugavet-Objekten	260
5.2	Die Daugavet-Eigenschaft in JB*-Tripeln	265
	Literatur	270
	Symbol-, Namen- und Sachverzeichnis	286

Kapitel 1

Ordnungen und Topologie

1.1 Relationen und Ordnungen

1.1.1. Zugrunde gelegt sei die Mengenlehre nach der Zermelo-Fraenkel-Axiomatik mit Fundierungsaxiom und Auswahlaxiom, in Zeichen $\mathbf{ZFC}$, siehe zum Beispiel ENDERTON [90](1977) oder in gestraffter Form KUNEN [191](1980), Kapitel 1. Eine Darstellung der Mengenlehre nach dem unfundierten Zermelo-Fraenkel-Axiomensystem, in Zeichen $\mathbf{ZF}^-$, mit dem Auswahlaxiom stellt HALMOS [124](1969) dar.

Auf $\mathbf{ZF}^-$ lassen sich die beiden Unvollständigkeitssätze von Kurt Gödel aus dem Jahr 1931 anwenden: Nach dem zweiten Gödelschen Unvollständigkeitssatz kann man innerhalb von $\mathbf{ZF}^-$ nicht die Widerspruchsfreiheit von $\mathbf{ZF}^-$ zeigen. Nach dem ersten Gödelschen Unvollständigkeitssatz ist $\mathbf{ZF}^-$, wenn man es als widerspruchsfrei annimmt, nicht vollständig, soll heißen, dass es Aussagen in $\mathbf{ZF}^-$ gibt, die innerhalb von $\mathbf{ZF}^-$ weder beweisbar noch widerlegbar sind — also sogenannte von $\mathbf{ZF}^-$ unabhängige Aussagen sind.

Nach Gödel, 1938, ist mit $\mathbf{ZF}^-$ auch das fundierte $\mathbf{ZF}^-$, in Zeichen $\mathbf{ZF}$, und mit diesem wiederum auch $\mathbf{ZFC}$ mit der verallgemeinerten Kontinuumshypothese widerspruchsfrei. Dabei lautet die *verallgemeinerte Kontinuumshypothese*, dass für alle Ordinalzahlen α die Gleichung $2^{\aleph_\alpha} = \aleph_{\alpha+1}$ gilt, wobei hier die Cantorsche Aleph-Bezeichnung verwendet wird: Für jede Ordinalzahl α bezeichnet $\aleph_\alpha$ die α -te unendliche Kardinalzahl. Als die *Kontinuumshypothese* an sich bezeichnet man die Aussage, dass die Gleichung $2^{\aleph_0} = \aleph_1$ gilt.

Paul Joseph Cohen zeigte 1963 mit seiner Erzwingungsmethode (im Englischen *forcing*), dass bei angenommener Widerspruchsfreiheit von $\mathbf{ZF}^-$ jeweils auch

- (a) $\mathbf{ZF}^-$ mit der Negation des Fundierungsaxioms,
- (b) $\mathbf{ZF}$ mit der Negation des Auswahlaxioms und
- (c) $\mathbf{ZFC}$ mit der Kontinuumshypothese $2^{\aleph_0} = \aleph_2$, mithin also auch mit der Negation der Kontinuumshypothese, und damit auch mit der Negation der verallgemeinerten Kontinuumshypothese

widerspruchsfrei ist, wofür ihm 1966 die Fields-Medaille verliehen wurde. Es folgt, dass das Auswahlaxiom von $\mathbf{ZF}$ unabhängig ist, und, dass die verallgemeinerte Kontinuumshypothese von $\mathbf{ZFC}$ unabhängig ist - jeweils bei angenommener Widerspruchsfreiheit von $\mathbf{ZF}^-$.

Man bemerke, dass es in $\mathbf{ZFC}$ außer Mengen keine Objekte gibt. Insbesondere gibt es keine Urelemente und Klassen. Urelemente, auch Individuen genannt, findet man in der unverzweigten Typentheorie, wo es eine restriktive Einteilung von Mengen in Stufen gibt. (In den Principia Mathematica (P.M.) von Bertrand Russell und Alfred North Whitehead

liegt die ursprüngliche verzweigte Typentheorie vor.) SUPPES [273](1972) stellt **ZFC** mit in dieser Theorie eingefügten Urelementen — die Zermelo 1908 ursprünglich auch selbst verwendete — dar. Bezüglich Klassen sei auf das von Neumann-Bernays-Gödelsche System, in Zeichen **NBG**, hingewiesen; es ist fundiert und enthält im Allgemeinen das Auswahlaxiom. Eine Einführung in die **NBG**-Mengenlehre stellen FRIEDRICHS DORF und PRESTEL [103](1985) dar. In ENDERTON [90](1977) findet man einen kurzen Vergleich von **NBG** mit der gegenüber **NBG** verbesserten Erweiterung von **NBG**, dem auch Klassen behandelnden *Morse-Kelley-System*, welches im Anhang von KELLEY [184] präsentiert wird. Einen ausführlicheren Vergleich findet man in FRAENKEL und BAR-HILLEL [97]. Des Weiteren sei auf OBERSCHELP [226](1994) aufmerksam gemacht, der als Basislogik die Klassenlogik aus GLUBRECHT, OBERSCHELP und TODT [110] (1983) verwendet, welche laut Oberschelp den Vorteil hat, dass ihre Sprache der üblicherweise in der Mathematik verwendeten Sprache näher kommt als eine prädikatenlogische Sprache, und, [226], Seite 23, „den geeigneten Rahmen [darstellt], in dem sich die Zermelo-Fraenkelsche Mengenlehre, aber auch beliebige alternative Ansätze, die mit Klassen umgehen, entfalten können“. Oberschelp [226](1994) konstruiert eine von ihm so benannte Allgemeine Mengenlehre, die, wenn sie per seinem Reinheitsaxiom (Jedes Element einer Menge ist wieder eine Menge, das heißt, es werden Urelemente ausgeschlossen), Fundierungsaxiom und Klassenextensionalitätsschema eingeschränkt wird, genau **ZFC** darstellt, siehe OBERSCHELP [226](1994), §41. Dabei zeigt Oberschelp auch, wie man die klassenlogisch formulierte **ZFC** in die übliche prädikatenlogische Form bringt.

Die bei Oberschelp als reale Klassen bezeichneten Objekte werden hier *Klassen* genannt. Ist eine solche Klasse keine Menge, so wird sie als *echte Klasse* bezeichnet (bei Oberschelp reale eigentliche Klasse).

In der prädikatenlogisch formulierten **ZFC** haben Klassen zwar keine formale Existenz, können dort aber sprachlich emuliert werden, indem jede Klasse für eine Formel steht, siehe so zum Beispiel JECH [155](2003) und KUNEN [191](1980).

In OBERSCHELP [225](1973) dagegen sind Klassenterme als eigenständige Ausdrücke zugelassen und es liegen neben Mengen (u.a.) echte Klassen vor, und Klassen können auch echte Klassen als Elemente besitzen. Die Mengen für sich genommen sind dabei genau im Sinne der **ZFC**-Mengenlehre zu verstehen.

Der allgemein hin akzeptierte Glaube an die Widerspruchsfreiheit von **ZFC** sei im Folgenden stets vorausgesetzt.

In der Kategorientheorie möchte man sogenannte *Konglomerate* von Klassen bilden können, mit denen man dann wieder wie mit Mengen rechnen kann. Dies kann man in **ZFC** beispielsweise durch die axiomatische Forderung der Existenz einer bestimmten Menge — genannt *Universum* — und einer anschließenden Umbenennung von Mengen erreichen. Im Wesentlichen ist dies äquivalent zu der axiomatisch geforderten Existenz einer stark unerreichbaren Kardinalzahl: Bezeichnet V_α für eine Ordinalzahl α die α -te von Neumann-Zermelo-Stufe (JECH [155](2003), Seite 64) und ist κ eine stark unerreichbare Kardinalzahl, so ist V_κ ein Modell für **ZFC** (JECH [155](2003), Seite 167). Dem entsprechend nehme man die folgenden drei Umbenennungen gleichzeitig vor:

- jede Teilmenge von V_κ nenne man eine Klasse,
- jede Menge nenne man ein Konglomerat,
- jedes Element von V_κ nenne man eine Menge.

Man führt also im Rahmen eines Modells der Zermelo-Fraenkel-Mengenlehre eine Konstruktion durch, um ein Modell der Zermelo-Fraenkel-Mengenlehre zu erhalten, welches im erstgenannten Modell eingebettet ist (HERRLICH und STRECKER [132](1973), Seiten 9-12 und 329-331. Siehe auch ADÁMEK, HERRLICH und STRECKER [4](2004), Seiten 13-17).

Auf dem Weg zu einer Interpretation der Begriffe Menge, Klasse und „enthalten in“ ist die Umschreibung des Grundbegriffes der Menge von GEORG CANTOR aus dem Jahr 1898 ein guter Wegweiser; nach ihm ist eine *Menge* (bei ihm auch *konsistente Vielheit*) jede Zusammenfassung von wohlbestimmten und wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens (welche die *Elemente* genannt werden) zu einem Ganzen; dabei setzt er dieser Zusammensetzung zu einem Ganzen als ein Ding voraus, dass alle besagten Elemente ohne Widerspruch als „zusammenseiend“ gedacht werden können. Geht letzteres nicht, so bezeichnet Cantor die Vielheit aller besagten Elemente als *absolut unendlich* oder *inkonsistente Vielheit*, was dem in der heutigen Terminologie üblichen Begriff der echten Klasse nahe kommt (und noch näher dem Begriff der eigentlichen Klassen bei Oberschelp).

1.1.2 (JECH [155]; VON QUERENBURG [244](1979); STORCH und WIEBE [269] (2003)). Die leere Menge wird mit $\emptyset$ bezeichnet. Seien A , B und I Mengen. Ist A eine Teilmenge von B , so wird dies mit $A \subseteq B$ oder mit $B \supseteq A$ symbolisiert. Falls A dabei nicht gleich B ist, heißt A eine *echte* (im Englischen *strict* oder *proper*) Teilmenge von B und dies kann durch $A \subsetneq B$ oder $B \supsetneq A$ symbolisiert werden. Die anderorts auch üblichen Symbole $\subset$ und $\supset$ werden nicht verwendet. Mit $\mathcal{P}(A)$ wird die *Potenzmenge* von A bezeichnet. Sind A und B zwei Teilmengen einer Menge, so wird mit $A \setminus B$ die Menge $\{x \in A : x \notin B\}$ bezeichnet. Ist A eine nicht leere Menge und p ein Element von A , so heißt das Paar (A, p) eine *punktierte Menge mit Basispunkt* p . Wenn keine Konfusion möglich erscheint, wird der Basispunkt nicht weiter erwähnt und es wird einfach nur von der punktierten Menge A gesprochen. Ist (A, p) eine punktierte Menge, so wird für jede Teilmenge M von A mit $M^\times$ die Menge $M \setminus \{p\}$ bezeichnet. Die Menge der natürlichen Zahlen inklusive der Null wird mit dem Doppelstrichbuchstaben $\mathbb{N}$ bezeichnet und wird stets auch als eine punktierte Menge $(\mathbb{N}, 0)$ betrachtet.

Sei A_ν für jedes $\nu \in I$ eine Teilmenge von A . Dann gilt $\bigcup_{\nu \in \emptyset} A_\nu = \emptyset$ und $\bigcap_{\nu \in \emptyset} A_\nu = A$. Eine *Abbildung* f von A nach B (im Englischen *from A to B*; oder auch *von A in B*, im Englischen *of A into B*) ist eine Vorschrift, die jedem Element x aus A genau ein Element von B zuordnet, welches im Allgemeinen mit $f(x)$ bezeichnet wird; suggestiv wird solch eine Abbildung wie üblich kurz als $f: A \rightarrow B$ geschrieben und auch als $x \mapsto f(x)$ angegeben. Die Menge aller Abbildungen f von I nach B mit $f(\nu) \in A_\nu$ für alle $\nu \in I$ heißt das (*mengentheoretische*) *Produkt* der Mengen A_ν , $\nu \in I$, und wird mit $\prod_{\nu \in I} A_\nu$ symbolisiert; im Folgenden wird es stets das *kartesische Produkt* der Mengen A_ν , $\nu \in I$, genannt; dabei ist es üblich, es so einzurichten, dass B gleich der Menge $\bigcup_{\nu \in I} A_\nu$ ist. Im Fall von $A_\nu = A$ für alle $\nu \in I$ schreibt man A^I statt $\prod_{\nu \in I} A_\nu$; ist dabei auch noch $I = \{1, \dots, n\}$ für ein $n \in \mathbb{N}^\times$, so schreibt man statt A^I einfach nur A^n .

Sei f eine Abbildung von A nach B . Sei M eine Teilmenge von A . Die Menge $\{f(x) : x \in M\}$ heißt das *Bild* von M unter f und wird mit $f[M]$ bezeichnet; wenn keine Missverständnisse möglich sind, insbesondere also, wenn M kein Element von A ist oder zumindest nicht im üblichen Sinne als ein solches aufgefasst wird, wird das Bild von M unter f meist auch einfach mit $f(M)$ bezeichnet. Das Bild von A unter f wird nach dem englischen Wort *range* für Wertebereich mit $\text{ran}(f)$ bezeichnet. Falls $\text{ran}(f) = B$ vorliegt, sagt man, die Abbildung f sei *surjektiv* oder eine *Surjektion*; in diesem Fall sagt man dann auch, dass f eine Abbildung von A *auf* (im Englischen *onto*, im Französischen *sur*) B sei. Man sagt, die Abbildung f sei *injektiv* (im Englischen auch 1-1 oder *one-to-one*; im Französischen auch *biunivoque*) oder eine *Injektion*, falls für je zwei verschiedene Elemente x und y aus A stets auch die beiden Elemente $f(x)$ und $f(y)$ aus B voneinander verschieden sind. Genau dann, wenn f injektiv ist, ist durch $f(A) \rightarrow A$, $f(x) \mapsto x$ die mit f^{-1} bezeichnete *Umkehrabbildung* von f erklärt. Unabhängig davon, ob f injektiv ist, ist aber stets die ebenfalls mit f^{-1} bezeichnete Umkehrabbildung $\mathcal{P}(B) \rightarrow \mathcal{P}(A)$, $N \mapsto \{x \in A : f(x) \in N\}$ definiert. Die mit $f|M$

bezeichnete Abbildung $M \rightarrow B$, $x \mapsto f(x)$ heißt die *Restriktion von f auf M* . Ist N eine Teilmenge von B mit $f(A) \subseteq N$, so heißt die mit $N|f$ bezeichnete Abbildung $A \rightarrow N$, $x \mapsto f(x)$ die *Astriktion von f auf N* . Um Missverständnisse zu vermeiden, zum Beispiel eine Verwechslung mit dem Absolutbetrag (3.5.45) oder für den Fall, dass M oder N Mengen von Funktionen sind, wird anstelle von $f|M$ auch genauer $f \upharpoonright M$ und anstelle von $N|f$ auch genauer $N \upharpoonright f$ geschrieben.

Eine Menge M heißt *endlich*, wenn es eine bijektive Abbildung auf ein $n \in \mathbb{N}$, n aufgefasst als eine Ordinalzahl, gibt. M heißt *unendlich*, wenn M nicht endlich ist. Man beachte, dass mit dieser Definition auch die leere Menge eine endliche Menge ist. Ist M eine Menge, so sagt man, dass eine Teilmenge N von M *fast alle* Elemente von M enthalte, wenn $M \setminus N$ endlich ist. Eine bijektive Abbildung von A auf sich selbst heißt eine *Permutation* von A . Ist f eine Abbildung von A nach A — also eine Abbildung der Menge A in sich selbst, eine sogenannte *Selbstabbildung* der Menge A —, so heißt die Menge $\{x \in A : f(x) = x\}$ die *Fixpunktmenge* der Abbildung f ; sie wird mit $A_{(f)}$ bezeichnet. Es sei an die dazu im Vorwort gemachte Bemerkung erinnert. Die identische Selbstabbildung $x \mapsto x$ der Menge A wird mit Id_A oder bei klarem Kontext einfach mit Id bezeichnet.

Wie allgemein hin üblich wird das kartesische Produkt $\prod_{\nu \in I} A_\nu$ im Fall einer endlichen, nicht leeren Indexmenge I , etwa $I = \{1, \dots, n\}$ für ein $n \in \mathbb{N}^\times$, auch aufgefasst als die mit $A_1 \times \dots \times A_n$ symbolisierte Menge aller geordneten n -Tupel $(a_1, \dots, a_n)$ mit $a_i \in A_i$ für $i = 1, \dots, n$; siehe zum Beispiel HUNGERFORD [143](1980), Seite 8.

1.1.3 Definition. Seien X , Y und Z drei Mengen. Ist R eine Teilmenge von $X \times Y$, so heißt R eine *Relation von X nach Y* (oder auch *zwischen X und Y*); bei $X = Y$ sagt man üblicherweise einfach auch, R sei eine *Relation auf X* . Ist R eine Relation von X nach Y , so schreibt man synonym für $(x, y) \in R$ auch xRy ; des Weiteren heißt dann die mit R^{-1} symbolisierte Menge $\{(y, x) \in Y \times X : (x, y) \in R\}$ die *zu R duale* oder *inverse* oder *konverse Relation*; ist überdies S eine Relation von Y nach Z , so ist mit $S \circ R$ die Relation $\{(x, z) \in X \times Z : \exists y \in Y : (x, y) \in R \text{ und } (y, z) \in S\}$ von X nach Z gemeint. (Es sei darauf hingewiesen, dass manche Autoren $R \circ S$ für diese Menge schreiben.)

Sei R eine Relation von X nach Y . Sei $M \subseteq X$ und $N \subseteq Y$. Dann heißt die Menge all der Elemente von Y , die jeweils mit jedem Element von M in Relation R stehen, also die Menge

$$M^\perp := \{y \in Y : (x, y) \in R \text{ für alle } x \in M\},$$

der *rechte Träger* von M (bezüglich der Relation R). Entsprechend heißt die Menge

$$N_\perp := \{x \in X : (x, y) \in R \text{ für alle } y \in N\},$$

der *linke Träger* von N (bezüglich der Relation R). Die Menge M heißt *trägerabgeschlossen* (bezüglich der Relation R), falls $M = M^\perp_\perp$ gilt; $M^\perp_\perp$ heißt die *trägerabgeschlossene Hülle* von M (bezüglich der Relation R); N jeweils entsprechend. Man bemerke, dass mit dieser Symbolik sowohl $y \in \{x\}^\perp$ als auch $x \in \{y\}_\perp$ äquivalent zu xRy , sprich $(x, y) \in R$, ist. Für eine Liste von elementaren Rechenregeln und Eigenschaften der Trägerbildung siehe 1.1.19.

Die *Diagonale* von X ist die Menge $\{(x, x) \in X \times X : x \in X\}$ und wird mit Δ_X bezeichnet oder bei klarem Kontext auch einfach nur mit Δ .

Sei R eine Relation auf X . Die Relation R heißt auf X

- (a) *reflexiv*, wenn $\Delta \subseteq R$.
- (b) *irreflexiv*, wenn $\Delta \cap R = \emptyset$.
- (c) *symmetrisch*, wenn $R \subseteq R^{-1}$, also $R = R^{-1}$.

- (d) *antisymmetrisch*, wenn $R \cap R^{-1} \subseteq \Delta$.
- (e) *asymmetrisch*, wenn R irreflexiv und antisymmetrisch ist, das heißt, wenn $R \cap R^{-1} = \emptyset$.
- (f) *transitiv*, wenn $R \circ R \subseteq R$.
- (g) eine *Präordnung* oder eine *Quasiordnung*, wenn R reflexiv und transitiv ist. Dies ist wegen $R \circ R \subseteq \Delta \cup (R \circ R) \subseteq R = R \circ \Delta \subseteq R \circ R$ genau dann der Fall, wenn $R = \Delta \cup (R \circ R)$ gilt.
- (h) eine *partielle Ordnung* oder eine *Halbordnung*, wenn R eine antisymmetrische Präordnung ist.
- (i) eine *strikte partielle Ordnung* oder eine *strikte Halbordnung*, wenn R irreflexiv und transitiv ist.
- (j) eine *Totalordnung*, wenn $(X \times X) \setminus \Delta \subseteq R \cup R^{-1}$, wobei R eine strikte partielle Ordnung oder eine partielle Ordnung ist. Eine *Kette* ist eine total geordnete Teilmenge einer prägeordneten Menge.
- (k) eine *Richtung*, wenn R eine Präordnung ist und je zwei beliebige Punkte von X eine obere Schranke haben, das heißt, $\forall x, y \in X \exists z \in X : (x, z) \in R$ und $(y, z) \in R$. Eine Präordnung mit dieser Eigenschaft wird auch *aufwärts filtrierend* (im Englischen *upward-filtering*) genannt. Für $X \neq \emptyset$ heißt (X, R) dann eine *gerichtete Menge*. Eine Teilmenge J einer gerichteten Menge $(I, \leq)$ heißt *kofinal* in I , wenn für jedes $i \in I$ ein $j \in J$ mit $i \leq j$ existiert.
- (l) eine *Äquivalenzrelation* oder eine *Faserung* von X , wenn R reflexiv, symmetrisch und transitiv ist.

1.1.4 Bemerkungen. Sei R eine Relation auf einer Menge.

(a) Ist R eine strikte partielle Ordnung, dann ist der sogenannte reflexive Abschluss von R , das heißt, $R \cup \Delta$, eine partielle Ordnung. Ist R eine partielle Ordnung, dann ist die reflexive Reduktion von R , das heißt, $R \setminus \Delta$, eine strikte partielle Ordnung.

(b) Sei E eine der in 1.1.3 definierten elementaren Eigenschaften (a) bis (f), formal also $E \in \{ \text{reflexiv, irreflexiv, symmetrisch, antisymmetrisch, asymmetrisch, transitiv} \}$. Dann gilt: Die Relation R weist genau dann die Eigenschaft E auf, wenn die zu R duale Relation R^{-1} sie aufweist.

1.1.5. Seien A und B zwei Mengen. Jede Abbildung $f: A \rightarrow B$ wird hier in umkehrbar eindeutiger Weise auch aufgefasst als eine Relation R von A nach B mit der Eigenschaft, dass jedes Element von A die erste Komponente von genau einem geordneten Paar in R ist; diese Relation heißt auch der *Graph* der Abbildung f ,

$$\text{graph}(f) := \{(x, y) \in A \times B : y = f(x)\};$$

vergleiche 1.1.6. In diesem Sinne bemerke man in der Situation von 1.1.2, dass $\prod_{\nu \in \emptyset} A_\nu = \{\emptyset\}$ gilt. Allgemeiner ist das Auswahlaxiom der Mengenlehre äquivalent zu der Aussage, dass, wenn alle A_ν , $\nu \in I$, nicht leer sind, deren kartesisches Produkt nicht leer ist.

Eine Relation R von A nach B ist genau dann eine Abbildung, wenn $R \circ R^{-1} \subseteq \Delta_B$ gilt; denn $R \circ R^{-1} = \{(b_1, b_2) \in B \times B : \exists a \in A : (a, b_1) \in R \text{ und } (a, b_2) \in R\}$.

1.1.6 Mengenwertige Abbildungen (AUBIN und FRANKOWSKA [11](1990)). Seien A und B zwei Mengen. Sei $f: A \rightarrow \mathcal{P}(B)$ eine Abbildung. Für die Theorie der mengenwertigen Abbildungen hat es sich als fruchtbar erwiesen, solch ein f als eine Teilmenge von

$A \times B$ aufzufassen, und zwar analog, wie ja auch ursprünglich jede übliche Abbildung $A \rightarrow B$ als $\text{graph}(f)$ aufgefasst worden ist. Die mengenwertige Abbildung f wird (im sogenannten *graphical approach*) charakterisiert durch die folgende Teilmenge von $A \times B$:

$$\text{Graph}(f) := \{(x, y) \in A \times B : y \in f(x)\}.$$

(Durch das großgeschriebene „G“ ist somit auch keine Verwechslung mit dem in 1.1.5 definierten $\text{graph}(f)$ möglich.) Gemäß der Theorie der mengenwertigen Abbildungen schreibt man f auch als $A \rightsquigarrow B$ oder $A \rightrightarrows B$. Im Englischen gängige Bezeichnungen für eine mengenwertige Abbildung sind: *multifunction*, *multivalued map*, *multimap* oder schlicht *multi*.

Die Charakterisierung einer mengenwertigen Abbildung durch $\text{Graph}(f)$ legt die folgende Definition der Inversen von f als einer mengenwertigen Abbildung nahe, die hier mit f^{-1} bezeichnet wird, um eine Verwechslung mit den in 1.1.2 definierten Umkehrabbildungen f^{-1} zu vermeiden:

$$f^{-1}: B \rightsquigarrow A, \quad y \mapsto \{x \in A : y \in f(x)\}.$$

Ist $M \subseteq A$, so setzt man $f\langle M \rangle := \bigcup_{x \in M} f(x)$; insbesondere setzt man $\text{Ran}(f) := f\langle A \rangle$. Ist $N \subseteq B$, so hat man also $f^{-1}\langle N \rangle = \bigcup_{y \in N} \{x \in A : y \in f(x)\}$, also

$$f^{-1}\langle N \rangle = \{x \in A : f(x) \cap N \neq \emptyset\}.$$

1.1.7 Definition (DUNFORD und SCHWARTZ [78](1970), VON QUERENBURG [244](1979)). Sei X eine Menge, R eine transitive Relation auf X und M eine Teilmenge von X .

- (a) $x \in X$ ist ein *kleinstes* (oder *erstes*, im Englischen *least* oder *smallest*) *Element* von X , wenn $\{x\} \times (X \setminus \{x\}) \subseteq R$.
- (b) $x \in X$ ist ein *größtes* (im Englischen *greatest*) *Element* von X , wenn $(X \setminus \{x\}) \times \{x\} \subseteq R$.

Ist die Relation R zusätzlich reflexiv (also eine Präordnung), dann definiert man weiter:

- (c) $x \in X$ ist ein *minimales Element* von X , wenn $\forall y \in X : ((y, x) \in R \Rightarrow (x, y) \in R)$.
- (d) $x \in X$ ist ein *maximales Element* von X , wenn $\forall y \in X : ((x, y) \in R \Rightarrow (y, x) \in R)$.
Ein $x \in X$ ist also genau dann kein maximales Element von X , wenn ein $y \in X$ existiert mit $(x, y) \in R$ und $(y, x) \notin R$.

Man beachte, dass in dieser Form der Definition des minimalen und maximalen Elementes keine Gleichheits- oder Ungleichheitsforderungen verwendet werden; siehe auch 1.1.12. Dadurch sind die in Definition 2.1.47 definierten schief-minimal idempotenten Elemente einer Algebra auch wirklich minimale Elemente bezüglich einer gewissen prägeordneten Menge. Die in (c) und (d) angegebenen Definitionen von einem minimalen und maximalen Element sind auch gültig, wenn R transitiv und irreflexiv ist.

Da für jede nicht leere Teilmenge A von X mit R auch $R \cap (A \times A)$ eine Präordnung ist, sind mit (c) und (d) auch die beiden Begriffe minimales und maximales Element von M definiert, siehe 1.1.11.

- (e) $x \in X$ ist eine *untere Schranke* (oder *Minorante*) von M , wenn $\{x\} \times M \subseteq R$.
- (f) $x \in X$ ist eine *obere Schranke* (oder *Majorante*) von M , wenn $M \times \{x\} \subseteq R$.
- (g) Eine untere Schranke $x \in X$ von M heißt ein *Infimum* von M , wenn x ein größtes Element der Menge der unteren Schranken von M ist. Es soll nur genau dann von einem $x \in X$ als *dem* Infimum von M die Rede sein, in Zeichen $x = \inf M$, wenn es ein Infimum von M ist und kein anderes Element von X ein Infimum von M ist.

- (h) Eine obere Schranke $x \in X$ heißt ein *Supremum* von M , wenn x ein kleinstes Element der Menge der oberen Schranken von M ist. Es soll nur genau dann von einem $x \in X$ als *dem* Supremum von M die Rede sein, in Zeichen $x = \sup M$, wenn es ein Supremum von M ist und kein anderes Element von X ein Supremum von M ist.
- (i) X heißt *Verband* (im Englischen *lattice*), wenn für jede zweielementige Teilmenge $\{x, y\}$ von X das Supremum $x \vee y := \sup \{x, y\}$ und das Infimum $x \wedge y := \inf \{x, y\}$ existiert. Ist X ein Verband, so heißt er *vollständig* (im Englischen *complete*), wenn für jede Teilmenge von X sowohl das Infimum als auch das Supremum in X existieren.
- Enthält X nur die unumgänglichen Suprema, soll heißen, nur jeweils das Suprema von zwei vergleichbaren Elementen, so heißt X ein *Antiverband* (im Englischen *antilattice*). In Zeichen: X Antiverband $:\Leftrightarrow \forall x, y \in X : ((x, y) \in R \text{ oder } (y, x) \in R) \Leftrightarrow \sup\{x, y\} \text{ existiert}$.
- (j) X heißt *induktiv geordnet*, wenn jede Kette in X eine obere Schranke hat. (Man bemerke, dass somit die leere Menge keine induktiv geordnete Menge ist.) Man beachte, dass nicht verlangt wird, dass besagte obere Schranke eine kleinste obere Schranke, also ein Supremum, sein soll.

1.1.8 Bemerkungen. (a) Die Relation eines Verbandes ist antisymmetrisch. (b) Jeder vollständiger Verband X hat per $\inf X$ ein kleinstes und per $\sup X$ ein größtes Element. (c) Der Satz 1.1.21 wird zeigen, dass es in der Definition eines vollständigen Verbandes genügt, für jede Teilmenge von X die Existenz entweder durchgehend jeweils nur des Infimums oder des Supremums zu fordern.

1.1.9 (DUNFORD und SCHWARTZ [78](1958), Seite 6). Das zu dem Auswahlaxiom äquivalente Lemma von Zorn lautet: Jede induktiv geordnete Menge hat ein maximales Element.

1.1.10 Maximal-Prinzip (KELLEY [184](1955), Seite 33). Sei X eine Familie von Mengen und $\leq$ die durch $A \leq B :\Leftrightarrow A \subseteq B$, $A, B \in X$, definierte partielle Ordnung auf X . Dann besagt das zu dem Auswahlaxiom äquivalente Maximal-Prinzip folgendes: Existiert für jede Kette K in $(X, \leq)$ ein $A \in X$ mit $\bigcup K \subseteq A$, so existiert in $(X, \leq)$ ein maximales Element.

1.1.11 Wohlfundiert (KUNEN [191](1980), Seite 98, 101). Sei X eine Menge und R eine Relation auf X . Ein $x \in X$ heißt ein *R -minimales Element* von X , falls es kein $y \in X$ mit $(y, x) \in R$ gibt. Die Relation R heißt *wohlfundiert* (im Englischen *well-founded*), wenn jede nicht leere Teilmenge A von X ein R -minimales Element $x \in A$ hat.

Fügt man das Fundierungsaxiom (auch Regularitätsaxiom genannt, bei Bernays: axiom of restriction) beim Aufbau der Mengenlehre nach Zermelo-Fraenkel als letztes Axiom hinzu, so kann man es dabei wie folgt formulieren (Zermelo, 1930): Jede nicht leere Menge hat ein $\in$ -minimales Element; mit anderen Worten ist $\in$ auf jeder Menge wohlfundiert.

Sei nun R eine strikte partielle Ordnung auf X . Dann ist ein Element von X genau dann ein minimales Element von X , wenn es ein R -minimales Element von X ist. R -maximal alles entsprechend.

Beweis. Sei x ein minimales Element von X . Dies ist äquivalent zu: $\forall y \in X : (\neg(y, x) \in R) \vee ((x, y) \in R)$. („ $\neg$ “ bezeichnet im vorliegenden Text stets die logische Verneinung.) Sei $y \in X$. Wäre $(y, x) \in R$ mit $(x, y) \in R$, dann wäre wegen der Transitivität $(x, x) \in R$,

im *Widerspruch* zur Irreflexivität. Also ist die erste Klammer im Ausdruck $(\neg(y, x) \in R) \vee ((x, y) \in R)$ für alle $y \in X$ wahr, die zweite Klammer kann weggelassen werden und es liegt R -Minimalität vor. $\square$

1.1.12. Sei X eine Menge, R eine partielle Ordnung auf X und $x \in X$. Dann gilt: **(a)** x ist ein R -minimales Element von $X \Rightarrow x$ ist ein minimales Element von $X \Leftrightarrow \forall y \in X : ((y, x) \in R \Rightarrow x = y)$; **(b)** x ist ein R -maximales Element von $X \Rightarrow x$ ist ein maximales Element von $X \Leftrightarrow \forall y \in X : ((x, y) \in R \Rightarrow x = y)$.

Beweis. Für die erste Implikation beachte nur die bereits im Beweis von 1.1.11 erwähnte formale Formulierung der Minimalität; Maximalität entsprechend.

Sowohl in (a) als auch in (b) folgt in der Äquivalenz die Implikation „ $\Rightarrow$ “ aus der Antisymmetrie von R und „ $\Leftarrow$ “ aus der Reflexivität von R . $\square$

Des Weiteren ist jedes kleinste Element von X ein minimales Element von X , und jedes größte Element von X ein maximales Element X ; wegen der Antisymmetrie kann X höchstens ein kleinstes Element und ebenso höchstens ein größtes Element besitzen.

1.1.13 Kettenbedingungen (ROWEN [250](1988); KUROSCHE [193](dt.)(1964)). Sei X eine Menge und R eine partielle Ordnung auf X . Man sagt, (X, R) erfülle die *Minimalbedingung* (im Englischen *minimal condition*), wenn jede nicht leere Teilmenge $M \subseteq X$ ein minimales Element $m \in M$ hat. Und analog sagt man: (X, R) erfülle die *Maximalbedingung* (im Englischen *maximal condition*), wenn jede nicht leere Teilmenge $M \subseteq X$ ein maximales Element $m \in M$ hat.

Im Folgenden sei die Relation R als $\leq$ notiert. Wie üblich, bezeichne dann „ $x < y$ “ die Situation „ $x \leq y$ und $x \neq y$ “. Man sagt, $(X, \leq)$ erfülle die *fallende Ketten-Bedingung* oder die *Bedingung des Abbrechens absteigender Ketten* (im Englischen *descending chain condition*, abgekürzt *DCC*), wenn es keine unendliche echt fallende Kette $x_1 > x_2 > x_3 > \dots$ in X gibt, mit anderen Worten, wenn in X jede fallende Kette $x_1 \geq x_2 \geq x_3 \geq \dots$ endlich ist, also ein $n \in \mathbb{N}^\times$ mit $x_n = x_{n+1} = x_{n+2} = \dots$ existiert. Man sagt dann auch, X sei eine *fundierte Menge*. Analog sagt man, $(X, \leq)$ erfülle die *aufsteigende Ketten-Bedingung* (im Englischen *ascending chain condition*, abgekürzt *ACC*), wenn es in X keine unendliche echt aufsteigende Kette $x_1 < x_2 < x_3 < \dots$ gibt.

Es gilt (BIRKHOFF [25](1973), Seite 180, und ROWEN [250](1988)):

(a) (X, R) genügt genau dann der Maximalbedingung, wenn jede nicht leere Kette $K \subseteq X$ ein maximales Element besitzt, insbesondere also genau dann, wenn (X, R) der ACC genügt.

(b) (X, R) genügt genau dann der Minimalbedingung, wenn jede nicht leere Kette $K \subseteq X$ ein minimales Element besitzt, insbesondere also genau dann, wenn (X, R) der DCC genügt.

Beweis. Da für $X = \emptyset$ die genannten Äquivalenzen trivialerweise erfüllt sind, sei $X \neq \emptyset$. Außerdem genügt es (a) zu zeigen, da die in (b) über (X, R) gemachte Aussage durch Anwenden von (a) auf (X, R^{-1}) folgt.

Genüge (X, R) nicht der ACC. Es gibt also eine unendliche aufsteigende Kette K . Somit gibt es eine nicht leere Teilkette $L \subseteq K$, so dass L kein größtes Element hat, also, da L eine Kette ist, kein maximales Element hat und die Maximalbedingung ist nicht erfüllt.

Genüge (X, R) nun der ACC. Sei R wieder als $\leq$ notiert. Sei $\emptyset \neq M \subseteq X$. Sei $K \neq \emptyset$ eine Kette in M . K hat ein maximales Element, also eine obere Schranke, denn: Sei $k_1 \in K$. Ist k_1 ein maximales Element von K , so ist man fertig. Ansonsten existiert ein $k_2 \in K \setminus \{k_1\}$ mit $k_1 \leq k_2$, also $k_1 < k_2$. Ist k_2 ein maximales Element, so ist man fertig, ansonsten verfähre man entsprechend weiter. Man erhält so eine echt aufsteigende Kette

$k_1 < k_2 < k_3 < \dots$, die nach Voraussetzung endlich sein muss, das heißt, es gibt für ein $n \in \mathbb{N}^\times$ ein $k_n \in K$ mit $k_1 < k_2 < \dots < k_n$, so dass kein $k \in K$ existiert mit $k_n < k$; k_n ist also ein maximales Element von K .

Somit ist M induktiv geordnet und nach dem Lemma von Zorn hat M ein maximales Element. $\square$

Genügt $(X, \leq)$ der ACC, so wird $(X, \leq)$ oder wie üblich, falls keine Missverständnisse möglich sind, einfach auch nur X , *noethersch* (im Englischen *Noetherian*) genannt. Und genügt $(X, \leq)$ der DCC, so wird X *artinsch* (im Englischen *Artinian*) genannt. Namenspatrone dieser Bezeichnungen waren Emil Artin und Emmy Noether.

Es sei noch auf die folgende Äquivalenz aufmerksam gemacht: $(X, \leq)$ ist genau dann artinsch, wenn $(X, \leq)$ die folgende sogenannte *Induktionsbedingung* erfüllt: Sei $\mathcal{S}$ eine Eigenschaft, die ein Element der mit $\leq$ partiell geordneten Menge haben kann; haben dann alle (insofern überhaupt vorhandenen) minimalen Elemente die Eigenschaft $\mathcal{S}$ und kann man aus dem Vorhandensein der Eigenschaft $\mathcal{S}$ bei allen Elementen $x \in X$ mit $x < a$ für ein gewisses $a \in X$ auf das Bestehen dieser Eigenschaft beim Element a selbst geschlossen werden, so besitzen alle Elemente von X die Eigenschaft $\mathcal{S}$.

Beweis. Offensichtlich schreibt sich die Induktionsbedingung etwas formaler als:

$$\forall S \subseteq X \text{ mit } \{x \in X : x \text{ minimal in } (X, \leq)\} \subseteq S \text{ und } \left(\forall a \in X : (\{x \in X : x < a\} \subseteq S \Rightarrow a \in S) \right) : S = X.$$

Erfülle $(X, \leq)$ die Minimalbedingung und sei $\mathcal{S}$ eine Eigenschaft, die die Voraussetzungen der Induktionsbedingung erfüllt. Bezeichne S die Menge aller $x \in X$, die die Eigenschaft $\mathcal{S}$ besitzen. Angenommen, es sei $S \neq X$. Dann ist die Menge $N := X \setminus S$ nicht leer und besitzt nach Voraussetzung ein minimales Element, etwa a . Nach der ersten Voraussetzung der Induktionsbedingung kann a kein minimales Element für ganz X sein, die Menge $V := \{x \in X : x < a\}$ ist also nicht leer. Da a in $N = X \setminus S$ minimal ist, gilt $V \subseteq S$. Nach der zweiten Voraussetzung der Induktionsbedingung ist $a \in S$, *Widerspruch*. Also ist in Wirklichkeit $S = X$ und alle Elemente von X besitzen die Eigenschaft $\mathcal{S}$.

Erfülle nun $(X, \leq)$ die Induktionsbedingung. Sei konkret wie folgt eine Eigenschaft $\mathcal{S}$ definiert: Ein $x \in X$ hat die Eigenschaft $\mathcal{S}$, wenn jede bei x beginnende, echt fallende Kette $x > x_1 > x_2 > \dots$ in X nach endlich vielen Gliedern abbricht. Bezeichne S die Menge aller $x \in X$, die diese Eigenschaft $\mathcal{S}$ erfüllen. Offensichtlich sind vorhandene minimale Elemente von X in S enthalten. Sei nun $a \in X$ ein Element von X mit $\{x \in X : x < a\} \subseteq S$. Sei dann $a > x_1 > x_2 > \dots$ eine beliebige bei a beginnende, echt fallende Kette in X . Dann ist $x_1 \in S$, x_1 hat also die Eigenschaft $\mathcal{S}$ und somit hat auch a die Eigenschaft $\mathcal{S}$. Aus der Induktionsbedingung folgt, dass alle $x \in X$ die Eigenschaft $\mathcal{S}$ haben, das heißt, $(X, \leq)$ erfüllt die DCC. $\square$

Man bemerke, dass die Induktionsbedingung es ermöglicht, induktive Beweise und induktive, rekursive Konstruktionen durchzuführen. Ist zum Beispiel die Menge X eine Ordinalzahl, so liegt das Beweisprinzip der *transfiniten Induktion* vor. Bezüglich der *transfiniten Rekursion* siehe KÜROSCH, ebd.

Eine artinsche Totalordnung bezeichnet man als *Wohlordnung*.

1.1.14 Netz. Sei X eine Menge. Ein *Netz* in X ist eine Abbildung x von einer gerichteten Menge I nach X ; Schreibweise: $(x_\nu)_{\nu \in I}$ oder einfach (x_ν) , wobei jeweils $x_\nu := x(\nu)$ für alle $\nu \in I$ gesetzt wird; um dabei einfach deutlich zu machen, dass das Netz in X liegt, wird hier gelegentlich auch die Schreibweise $(x_\nu)_{\nu \in I} \subseteq X$ beziehungsweise einfach nur $(x_\nu) \subseteq X$ verwendet — sie dürfte zu keinen Missverständnissen führen. I heißt der

Index des Netzes. Andere übliche Bezeichnungen für ein Netz sind: *Moore-Smith-Folge*, *verallgemeinerte Folge* (im Englischen *generalized sequence*) oder *gerichtete Familie*.

Ein Netz $y: J \rightarrow X$ heißt ein *Teilnetz* eines Netzes $x: I \rightarrow X$ oder *feiner* als ein Netz $x: I \rightarrow X$, wenn eine Abbildung $\varphi: J \rightarrow I$ mit $y = x \circ \varphi$ existiert (äquivalent: $y_j = x_{\varphi(j)}$ für alle $j \in J$), so dass für jedes $i \in I$ ein $j_0 \in J$ mit $i \leq \varphi(j)$ für alle $j \geq j_0$ existiert.

Ist $x: I \rightarrow X$ ein Netz, J eine gerichtete Menge, $\varphi: J \rightarrow I$ eine *monotone Abbildung* — soll heißen, aus $j_1 \leq j_2$ in J folgt $\varphi(j_1) \leq \varphi(j_2)$ in I — und $\varphi(J)$ in I *kofinal*, so ist $y := x \circ \varphi$ ein Teilnetz von x .

Eine Folge ist ein Netz mit $\mathbb{N}$ als Indexmenge. Eine Teilfolge einer Folge $x: I \rightarrow X$ ist ein Teilnetz $y: J \rightarrow X$ des Netzes x , wobei die Abbildung $\varphi: J \rightarrow I$ mit $y = x \circ \varphi$ *monoton* ist.

1.1.15 Monotones Netz. Sei X eine Menge mit einer Präordnung $\leq$. Ein Netz $(x_\nu)_{\nu \in I}$ heißt *monoton steigend*, wenn für alle $\mu, \nu \in I$ aus $\mu \leq \nu$ die Ungleichung $x_\mu \leq x_\nu$ folgt. Es heißt *streng monoton steigend* (im Englischen *strong monotone increasing* oder auch *strictly increasing*), wenn für alle $\mu, \nu \in I$ aus $\mu < \nu$ die Ungleichung $x_\mu < x_\nu$ folgt. (*Streng*) *monoton fallend* (im Englischen (*strong*) *monotone decreasing* oder auch *strictly decreasing*) wird entsprechend definiert.

Ein Netz $(x_\nu)_{\nu \in I}$ heißt *nach oben ordnungsbeschränkt* (oder einfach nur *nach oben beschränkt*, wenn kein Missverständnis möglich ist), wenn die Menge $\{x_\nu : \nu \in I\}$ in X eine obere Schranke hat. *Nach unten (ordnungs)beschränkt* wird entsprechend definiert.

Sei M eine Teilmenge von X . M heißt (*streng*) *monoton steigend vollständig* (im Englischen (*strong*) *monotone increasing complete*), wenn jedes nach oben ordnungsbeschränkte, (*streng*) monoton steigende Netz $(x_\nu)_{\nu \in I}$, $x_\nu \in M$ für alle $\nu \in I$, ein Supremum $x := \sup\{x_\nu : \nu \in I\}$ mit $x \in M$ hat. Entsprechend ist (*streng*) *monoton fallend vollständig* (im Englischen (*strong*) *monotone decreasing complete*) definiert. Setze

$$M^m := \{x \in X : \text{Es gibt ein nach oben ordnungsbeschränktes,} \\ \text{monoton steigendes Netz } (x_\nu)_{\nu \in I} \text{ mit } x_\nu \in M \\ \text{für alle } \nu \in I \text{ und } x \in X \text{ als Supremum} \}$$

und

$$M_m := \{x \in X : \text{Es gibt ein nach unten ordnungsbeschränktes,} \\ \text{monoton fallendes Netz } (x_\nu)_{\nu \in I} \text{ mit } x_\nu \in M \\ \text{für alle } \nu \in I \text{ und } x \in X \text{ als Infimum} \}.$$

Dann gilt $M \subseteq M^m \cap M_m$ und es gelten die beiden Implikationen

$$M \text{ monoton steigend vollständig} \Rightarrow M = M^m$$

und

$$M \text{ monoton fallend vollständig} \Rightarrow M = M_m.$$

Dabei gilt im Allgemeinen jeweils nicht die umgekehrte Implikation.

Beweis. Es wird nur die erste Implikation gezeigt, da die zweite analog bewiesen werden kann. Sei M monoton steigend vollständig, dann ist $M^m \subseteq M$ zu zeigen. Sei $x \in M^m$. Es existiert also ein nach oben ordnungsbeschränktes, monoton steigendes Netz $(x_\nu)_{\nu \in I}$ mit $x_\nu \in M$ für alle $\nu \in I$ und $x := \sup\{x_\nu : \nu \in I\}$ als Supremum. Nach Voraussetzung ist dieses Supremum in M , also $x \in M$. Dass die umgekehrte Implikation nicht immer gelten muss, sieht man, wenn zum Beispiel $X = M$ und M die Menge der rationalen Zahlen ist, versehen mit der üblichen Ordnungsstruktur. $\square$

Weiteres zu monotonen Netzen findet man ab 2.2.15.

1.1.16 Filter (BOURBAKI [32](1966); PREUSS [243](1975), Seite 70; VON QUERENBURG [244](1979); KELLEY [184](1967), Seite 83).

Sei X eine Menge. Ein *Filter* (auf X) ist eine Menge $\mathcal{F}$ von Teilmengen von X , die die folgenden drei Eigenschaften besitzt:

(a) $\mathcal{F}$ ist abgeschlossen bezüglich der Bildung von Obermengen, das heißt, mit $A \in \mathcal{F}$, $A \subseteq B \subseteq X$ ist auch $B \in \mathcal{F}$.

(b) $\mathcal{F}$ ist abgeschlossen bezüglich der Bildung von endlichen Durchschnitten, das heißt, mit $A, B \in \mathcal{F}$ ist auch $A \cap B \in \mathcal{F}$.

(c) $\mathcal{F}$ enthält nicht die leere Menge $\emptyset$ als ein Element.

Man beachte, dass nach der Definition des Begriffes *endlich* in 1.1.2 aus (b) folgt, dass $X \in \mathcal{F}$ gilt, jeder Filter also insbesondere eine nicht leere Menge ist.

Die beiden Eigenschaften (a) und (b) lassen sich zu einer einzigen Eigenschaft zusammenfassen:

(a+b) Für alle $A, B \subseteq X$ gilt die Äquivalenz $A, B \in \mathcal{F} \Leftrightarrow A \cap B \in \mathcal{F}$.

Sind $\mathcal{F}_1$ und $\mathcal{F}_2$ zwei Filter auf X , so heißt $\mathcal{F}_1$ *feiner* als $\mathcal{F}_2$, falls $\mathcal{F}_1 \supseteq \mathcal{F}_2$ gilt. Ein Filter $\mathcal{F}$ auf X heißt ein *Ultrafilter*, falls es keinen von $\mathcal{F}$ verschiedenen Filter auf X gibt, der feiner als $\mathcal{F}$ ist. Mit anderen Worten ist ein Ultrafilter ein maximales Element der durch Inklusion geordneten Menge aller Filter auf X . Nach dem Lemma von Zorn (1.1.9) ist jeder Filter in einem Ultrafilter enthalten. Ein Ultrafilter $\mathcal{F}$ heißt *frei*, wenn der Durchschnitt aller Mengen $F \in \mathcal{F}$ leer ist, er heißt *fixiert*, wenn dieser Durchschnitt nicht leer ist. Ist zum Beispiel X eine nicht leere Menge und x ein Element von X , so ist $\mathcal{F} = \{\{x\}\}$ ein fixierter Ultrafilter auf X .

Ist $(x_\nu)_{\nu \in I}$ ein Netz in X , so ist $\{M \in \mathcal{P}(X) : \exists \gamma \in I \forall \nu \in I, \nu \geq \gamma : x_\nu \in M\}$ ein Filter auf X ; er heißt der zu dem Netz $(x_\nu)_{\nu \in I}$ *gehörige Filter* oder genauer auch *Abschnittsfilter* (im Englischen *section filter*); ist $I = \mathbb{N}$, so heißt er der zu dem Netz $(x_\nu)_{\nu \in I}$ *gehörige Elementarfilter*.

Sei $\mathcal{F}$ ein Filter auf X . Setze $I := \{(x, F) \in X \times \mathcal{F} : x \in F, F \in \mathcal{F}\}$. Durch $(x, F) \leq (y, G) :\Leftrightarrow G \subseteq F$ wird I zu einer gerichteten Menge $(I, \leq)$ und $\mathcal{F}$ ist der zu dem Netz $I \rightarrow X, (x, F) \mapsto x$ gehörige Abschnittsfilter.

Bezüglich Filter auf topologischen Räumen siehe 1.2.24.

1.1.17 Ordnungserhaltend. Seien $(X, \leq)$ und $(Y, \leq)$ zwei partiell geordnete Mengen. Eine Abbildung $T: X \rightarrow Y$ heißt *ordnungserhaltend* (im Englischen *order-preserving*), wenn aus $a \leq b$ in X stets $T(a) \leq T(b)$ in Y folgt, *isoton* (im Englischen *isotonic*) oder *streng ordnungserhaltend*, wenn aus $a < b$ in X stets $T(a) < T(b)$ in Y folgt, *antiisoton*, wenn aus $a < b$ in X stets $T(b) < T(a)$ in Y folgt, und *ordnungsumkehrend* (im Englischen *order-reversing*) oder *antiton*, wenn aus $a \leq b$ in X stets $T(b) \leq T(a)$ in Y folgt. Die Mengen X und Y heißen *ordnungsisomorph*, falls es eine Bijektion $T: X \rightarrow Y$ derart gibt, dass sowohl T als auch T^{-1} isoton sind; in diesem Fall nennt man T einen Ordnungsisomorphismus. *Antiordnungsisomorph* entsprechend.

1.1.18 Galois-Verbindungen (SMITH [264](2010)). Seien X und Y zwei Mengen. Sei $\leq$ eine partielle Ordnung auf X . Seien $\preceq$ und $\sqsubseteq$ zwei partielle Ordnungen auf Y . Seien $f^*: X \rightarrow Y$ und $f_*: Y \rightarrow X$ zwei Abbildungen.

(a) Man sagt, das Paar (f^*, f_*) bilde eine *Galois-Verbindung* (im Amerikanischen *Galois connection*, im Britischen *Galois connexion*) zwischen $(X, \leq)$ und $(Y, \preceq)$, falls die Äquivalenz

$$f^*(x) \preceq y \quad \Leftrightarrow \quad x \leq f_*(y) \quad (1.1)$$

für alle $x \in X, y \in Y$ gilt; in diesem Fall heißt dann f^* die *Linksadjungierte* von f_* , und f_* die *Rechtsadjungierte* von f^* . Ist zum Beispiel eine Abbildung $f: X \rightarrow Y$ ein

Ordnungsisomorphismus, so bildet das Paar (f, f^{-1}) eine Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$.

(b) Das Paar (f^*, f_*) bildet genau dann eine Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$, wenn die beiden folgenden Bedingungen erfüllt sind:

(i) f^* und f_* sind beide ordnungserhaltend.

(ii) Für alle $x \in X, y \in Y$ gilt sowohl

$$x \leq f_*(f^*(x)) \quad \text{als auch} \quad f^*(f_*(y)) \preceq y. \quad (1.2)$$

Beweis. Die beiden Bedingungen sind notwendig: Bilde dazu das Paar (f^*, f_*) eine Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$. Nach (1.1) gilt für alle $x \in X$ die Äquivalenz $f^*(x) \preceq f^*(x) \Leftrightarrow x \leq f_*(f^*(x))$. Da $\preceq$ reflexiv ist, ist die Aussage auf der linken Seite dieser Äquivalenz wahr. Für alle $x \in X$ gilt also $x \leq f_*(f^*(x))$. Die andere Ungleichung in (1.2) wird analog gezeigt. Seien nun $x_1, x_2 \in X$ mit $x_1 \leq x_2$. Wie eben gesehen, gilt $x_2 \leq f_*(f^*(x_2))$, also $x_1 \leq f_*(f^*(x_2))$. Nach (1.1) gilt $f^*(x_1) \preceq f^*(x_2)$ genau dann, wenn $x_1 \leq f_*(f^*(x_2))$ gilt. Also gilt $f^*(x_1) \preceq f^*(x_2)$, und f^* ist ordnungserhaltend. f_* entsprechend.

Die beiden Bedingungen sind hinreichend: Seien die beiden Aussagen (i) und (ii) erfüllt. Sei $x \in X$ und $y \in Y$ mit $f^*(x) \preceq y$. Da f_* nach (i) ordnungserhaltend ist, gilt $f_*(f^*(x)) \leq f_*(y)$. Da nach (1.2) die Ungleichung $x \leq f_*(f^*(x))$ gilt, hat man $x \leq f_*(y)$. Die andere Implikation von (1.1) geht analog. $\square$

Wegen Eigenschaft (i) nennt man die in (a) definierten Galois-Verbindungen auch *monotone Galois-Verbindungen*.

(c) Nach Bemerkung 1.1.4(b) ist auch $\preceq^{-1}$ eine partielle Ordnung auf Y . Mit ihr lässt sich die Äquivalenz (1.1) auch wie folgt formulieren:

$$y \preceq^{-1} f^*(x) \quad \Leftrightarrow \quad x \leq f_*(y).$$

Da nun die beiden Abbildungen f^* und f_* genau dann beide ordnungserhaltend bezüglich $\leq$ und $\preceq$ sind, wenn sie ordnungsumkehrend bezüglich $\leq$ und $\preceq^{-1}$ sind, definiert man: Man sagt, das Paar (f^*, f_*) bilde eine *antitone Galois-Verbindung* zwischen $(X, \leq)$ und $(Y, \sqsubseteq)$, wenn für alle $x \in X, y \in Y$ die Äquivalenz

$$y \sqsubseteq f^*(x) \quad \Leftrightarrow \quad x \leq f_*(y) \quad (1.3)$$

erfüllt ist, und in dem Sinn dieser Variante hatte Évariste Galois seine Untersuchungen vorgenommen; siehe 4.3.4. Das Paar (f^*, f_*) bildet also genau dann eine monotone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$, wenn es eine antitone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq^{-1})$ bildet.

Bisweilen ist es praktisch neben einer Aussage über monotone Galois-Verbindungen auch die — insofern zutreffende — entsprechende Aussage über antitone Galois-Verbindungen parat zu haben; sie folgt meist unmittelbar und wird im Folgenden jeweils auch angegeben. Des Weiteren wird zur Vermeidung von Verwechslungen das Adjektiv *monoton* bei den monotonen Galois-Verbindungen stets beibehalten.

(d) Die antitone Variante der in (b) aufgeführten Charakterisierung einer monotonen Galois-Verbindung lautet offensichtlich wie folgt. Das Paar (f^*, f_*) bildet genau dann eine antitone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \sqsubseteq)$, wenn die folgenden beiden Bedingungen erfüllt sind:

(i) f^* und f_* sind beide ordnungsumkehrend.

(ii) Für alle $x \in X$, $y \in Y$ gilt sowohl

$$x \leq f_*(f^*(x)) \quad \text{als auch} \quad y \sqsubseteq f^*(f_*(y)). \quad (1.4)$$

(e) Ist das Paar (f^*, f_*) eine monotone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$, so gilt:

- (i) Für jedes $x \in X$ ist $f^*(x)$ das kleinste Element der Menge $\{y \in Y : x \leq f_*(y)\}$.
- (ii) Für jedes $y \in Y$ ist $f_*(y)$ das größte Element der Menge $\{x \in X : f^*(x) \preceq y\}$.

Beweis. Es wird hier nur (i) bewiesen; der Beweis von (ii) geht analog. Bilde das Paar (f^*, f_*) eine monotone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$. Sei $x \in X$. Nach (1.2) gilt $x \leq f_*(f^*(x))$. Es ist also $f^*(x) \in \{y \in Y : x \leq f_*(y)\}$. Sei nun $v \in Y$ mit $x \leq f_*(v)$. Da f^* ordnungserhaltend ist, gilt $f^*(x) \preceq f^*(f_*(v))$. Nach (1.2) ist $f^*(f_*(v)) \preceq v$, also gilt $f^*(x) \preceq v$, und $f^*(x)$ ist das kleinste Element von $\{y \in Y : x \leq f_*(y)\}$. $\square$

Monotone Galois-Verbindungen kann man statt wie in (a) auch äquivalent über die beiden in (b) aufgeführten Bedingungen (i) und (ii) definieren. Eine sich aus diesem Blickwinkel aufbauende Theorie über monotone Galois-Verbindungen ist die im Englischen heißende *residuation theory*, siehe zum Beispiel BLYTH und JANOWITZ [27](1972). Bildet das Paar (f^*, f_*) eine monotone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$, so heißt in dieser Residuation Theory die Abbildung f^* im Englischen *residuated*, und die Abbildung f_* *the residual of f^** . Ein Grund, den monotonen Galois-Verbindungen solch eine Aufmerksamkeit zukommen zu lassen, ist die Eigenschaft, dass die Komposition zweier residuated Abbildungen stets wieder residuated ist, was dagegen für den antitonen Fall im Allgemeinen nicht zutrifft (BLYTH und JANOWITZ [27](1972), Seite 19, Übung 2.8, Bemerkung).

(f) Die entsprechende antitone Variante von (e) lautet: Ist das Paar (f^*, f_*) eine antitone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \sqsubseteq)$, so gilt:

- (i) Für jedes $x \in X$ ist $f^*(x)$ das größte Element der Menge $\{y \in Y : x \leq f_*(y)\}$.
- (ii) Für jedes $y \in Y$ ist $f_*(y)$ das größte Element der Menge $\{x \in X : f^*(x) \sqsubseteq y\}$.

(g) Bildet ein Paar (f, g) eine monotone oder antitone Galois-Verbindung zwischen zwei partiell geordneten Mengen, so gilt

$$(i) \quad f \circ g \circ f = f \quad \text{und} \quad g \circ f \circ g = g. \quad (1.5)$$

- (ii) Es ist genau dann $y \in \text{ran}(f)$, wenn y ein Fixpunkt der Abbildung $f \circ g$ ist. Es ist genau dann $x \in \text{ran}(g)$, wenn x ein Fixpunkt der Abbildung $g \circ f$ ist.

$$(iii) \quad \text{ran}(f) = f(\text{ran}(g)) \quad \text{und} \quad \text{ran}(g) = g(\text{ran}(f)).$$

Beweis. Sei das Paar (f^*, f_*) eine monotone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$.

(i) Nach (1.2) gilt für alle $x \in X$ die Aussage $x \leq f_*(f^*(x))$, und, da f^* ordnungserhaltend ist, somit auch die Aussage $f^*(x) \preceq f^*(f_*(f^*(x)))$. Nach (1.1) ist die Aussage $f^*(f_*(f^*(x))) \preceq f^*(x)$ äquivalent zu der Aussage $f_*(f^*(x)) \leq f_*(f^*(x))$, die wegen der Reflexivität von $\leq$ wahr ist. Mit der Antisymmetrie von $\leq$ folgt $f^*(x) = f^*(f_*(f^*(x)))$. Die andere Gleichung entsprechend.

(ii) Sei $y \in f^*(X)$. Es gibt also ein $x \in X$ mit $y = f^*(x)$. Somit $f^*(f_*(y)) = f^*(f_*(f^*(x))) \stackrel{(i)}{=} f^*(x) = y$ und y ist ein Fixpunkt von $f^* \circ f_*$. Die andere Aussage geht analog. Die umgekehrte Richtung ist jeweils klar.

(iii) Ist $y \in f^*(X)$, so gilt nach (ii) die Gleichung $y = f^*(x)$ für $x := f_*(y)$, also $f^*(X) \subseteq f^*(f_*(Y))$. Da $f_*(Y) \subseteq X$ gilt, ist die umgekehrte Inklusion klar. Die andere Gleichung analog.

Damit sind die Aussagen (i), (ii) und (iii) für den Fall einer monotonen Galois-Verbindung gezeigt. Bildet nun das Paar (f^*, f_*) eine antitone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \sqsubseteq)$, so bildet es — wie in (c) festgestellt — gleichzeitig auch eine monotone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \sqsubseteq^{-1})$. Da also die Abbildungen f^* und f_* zwischen den beiden Trägermengen X und Y bei beiden Galois-Verbindungs-Varianten jeweils dieselben sind, und in den Aussagen (i), (ii) und (iii) nur über diese Abbildungen etwas ausgesagt wird, ist auch der antitone Teil gezeigt. $\square$

(h) Bildet ein Paar (f, g) eine monotone oder antitone Galois-Verbindung zwischen $(X, \leq)$ und einer partiell geordneten Menge, so gelten für die Selbstabbildung $\varphi := g \circ f$ der Trägermenge X die folgenden drei Aussagen:

- (i) Für alle $x \in X$ gilt $x \leq \varphi(x)$.
- (ii) φ ist ordnungserhaltend.
- (iii) Für alle $x \in X$ gilt $\varphi(\varphi(x)) = \varphi(x)$.

Beweis. Zuerst sei der monotone Fall betrachtet. Sei das Paar (f^*, f_*) eine monotone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \preceq)$ bildend. Setze $\varphi := f_* \circ f^*$. Wegen (1.2) ist (i) erfüllt. Sei $a \leq b$ in X . Da f^* ordnungserhaltend ist, gilt $f^*(a) \preceq f^*(b)$. Da auch f_* ordnungserhaltend ist, gilt $f_*(f^*(a)) \leq f_*(f^*(b))$, und es gilt (ii). Nach (1.5) gilt $f^* \circ f_* \circ f^* = f^*$, also auch $(f^* \circ f_*) \circ (f^* \circ f_*) = (f^* \circ f_*)$, und somit (iii).

Bilde nun das Paar (f^*, f_*) eine antitone Galois-Verbindung zwischen $(X, \leq)$ und $(Y, \sqsubseteq)$. Setze wieder $\varphi := f_* \circ f^*$. Wegen (1.4) ist (i) erfüllt. Sei $a \leq b$ in X . Da f^* ordnungsumkehrend ist, gilt $f^*(b) \sqsubseteq f^*(a)$. Da auch f_* ordnungsumkehrend ist, gilt $f_*(f^*(a)) \leq f_*(f^*(b))$, und somit (ii). Die Aussage (iii) gilt genau aus dem gleichen Grund wie im monotonen Fall. $\square$

Eine Abbildung φ von der partiell geordneten Menge $(X, \leq)$ nach einer partiell geordneten Menge, die die eben aufgeführten drei Eigenschaften (i), (ii) und (iii) besitzt, nennt man auch einen *Hüllenoperator* oder einen *Abschlussoperator* (im Englischen *closure operator*) (und ist ein Spezialfall eines allgemeineren Konzeptes der *Hüllenbildung*, worauf hier nicht näher eingegangen wird).

(i) Seien U und V zwei Mengen. Sei R eine Relation von U nach V . Betrachte die in 1.1.3 erklärten rechten und linken Träger bezüglich der Relation R . Sei f^* die Abbildung $\mathcal{P}(U) \rightarrow \mathcal{P}(V)$, $M \mapsto M^\perp$, und f_* die Abbildung $\mathcal{P}(V) \rightarrow \mathcal{P}(U)$, $N \mapsto N_\perp$. Dann bildet das Paar (f^*, f_*) eine antitone Galois-Verbindung zwischen $(\mathcal{P}(U), \subseteq)$ und $(\mathcal{P}(V), \subseteq)$.

Beweis. Liege $X = \mathcal{P}(U)$ und $Y = \mathcal{P}(V)$ vor und außerdem, dass $\leq$ (bzw. $\sqsubseteq$) die auf X (bzw. Y) durch Mengeninklusion $\subseteq$ induzierte kanonische partielle Ordnung sei — also $A \leq$ (bzw. $\sqsubseteq$) $B \Leftrightarrow A \subseteq B$ für alle Elemente A, B aus X (bzw. Y) gelte. Für alle $M \subseteq U$ und $N \subseteq V$ gilt die folgende Äquivalenzkette: $(N \subseteq f^*(M)) \Leftrightarrow (N \subseteq M^\perp) \Leftrightarrow (\forall y \in N \forall x \in M : (x, y) \in R) \Leftrightarrow (\forall x \in M \forall y \in N : (x, y) \in R) \Leftrightarrow (M \subseteq N_\perp) \Leftrightarrow (M \subseteq f_*(N))$. Das heißt, es gilt die Äquivalenz (1.3), und das Paar (f^*, f_*) bildet eine antitone Galois-Verbindung zwischen $(\mathcal{P}(U), \leq)$ und $(\mathcal{P}(V), \sqsubseteq)$. $\square$

Jede Relation zwischen zwei Mengen induziert also in natürlicher Weise eine antitone Galois-Verbindung zwischen deren jeweils durch die mengentheoretische Inklusion kanonisch partiell geordneten Potenzmengen.

(j) (ERNÉ, KOSŁOWSKI, MELTON und STRECKER [91](1993), Satz 7(2)). Im folgenden Sinn gilt auch die Umkehrung von (i). Seien U und V zwei Mengen. Sei das Paar (f^*, f_*) eine antitone Galois-Verbindung zwischen $(\mathcal{P}(U), \subseteq)$ und $(\mathcal{P}(V), \subseteq)$. Dann wird diese antitone Galois-Verbindung durch die Relation

$$R := \{(u, v) \in U \times V : v \in f^*(u)\}$$

(was wegen (1.3) gleich $\{(u, v) \in U \times V : u \in f_*(v)\}$ ist) wie bei (i) beschrieben induziert.

1.1.19. Seien X und Y zwei Mengen. Sei R eine Relation von X nach Y . Es sei an die in 1.1.3 erklärten rechten und linken Träger bezüglich der Relation R erinnert. Wie eben in 1.1.18(i) gezeigt, bildet das aus der Bildung der rechten und der linken Träger bestehende Abbildungspaar zwischen den kanonisch durch Mengeninklusion partiell geordneten Potenzmengen von X und Y eine antitone Galois-Verbindung. Seien G und M Teilmengen von X . Für jedes $\nu \in I$, I eine beliebige Indexmenge, sei M_ν eine Teilmenge von X . Sei N eine Teilmenge von Y . Dann gilt:

- (a) $M^\perp = \bigcap_{m \in M} \{m\}^\perp$ und $N_\perp = \bigcap_{n \in N} \{n\}_\perp$.
- (b) $\left(\bigcup_{\nu \in I} M_\nu\right)^\perp = \bigcap_{\nu \in I} (M_\nu^\perp)$, speziell also $(G \cup M)^\perp = G^\perp \cap M^\perp$.
- (c) Aus $G \subseteq M$ folgt $M^\perp \subseteq G^\perp$.
- (d) $(M \times N \subseteq R) \Leftrightarrow (M \subseteq N_\perp) \Leftrightarrow (N \subseteq M^\perp)$. Da für $M^\perp$ nach Definition die Inklusion $M \times M^\perp \subseteq R$ gilt, gilt insbesondere $M \subseteq M^{\perp\perp}$. Und wegen $N_\perp \times N \subseteq R$ gilt analog $N \subseteq N_\perp^\perp$.
- (e) $M^\perp = M^{\perp\perp\perp}$ und $N_\perp = N_\perp^{\perp\perp}$.
- (f) $M^\perp \setminus (X^\perp) \subseteq Y \setminus ((X \setminus M)^\perp)$.
- (g) Für $\emptyset \subseteq X$ gilt $\emptyset^\perp = Y$.
- (h) $G^\perp \cup M^\perp \subseteq (G \cap M)^\perp$.
- (i) Falls G , M und $G^\perp \cup M^\perp$ trägerabgeschlossen sind, gilt $G^\perp \cup M^\perp = (G \cap M)^\perp$.
- (j) $M^{\perp\perp} = \{x \in X : (M \cup \{x\})^\perp_\perp = M^{\perp\perp}\}$.
- (k) $M^{\perp\perp} = \{x \in X : (M \cup \{x\})^\perp = M^\perp\}$.
- (l) M ist genau dann trägerabgeschlossen, wenn für alle $x \in X \setminus M$ ein $y \in M^\perp$ mit $(x, y) \notin R$ existiert.
- (m) Aus $M \subseteq G \subseteq M^{\perp\perp}$ folgt $M^\perp = G^\perp$.

Beweis. (a) ist klar. (d) ist auch klar; es sei an die Äquivalenzkette im Beweis von 1.1.18(i) erinnert; $M \subseteq M^{\perp\perp}$ ist auch schon per 1.1.18(h), dortige Unteraussage (i), bewiesen.

$$(b) \left(\bigcup_{\nu \in I} M_\nu\right)^\perp \stackrel{(a)}{=} \bigcap_{m \in \bigcup_{\nu \in I} M_\nu} \{m\}^\perp = \bigcap_{\nu \in I} \bigcap_{m \in M_\nu} \{m\}^\perp.$$

(c) Klar nach 1.1.18(d), die dortige Bedingung (i). Die Aussage kann aber auch hier mit (a) gezeigt werden: Sei $G \subseteq M$, dann gilt nach (a): $M^\perp = \left(\bigcap_{g \in G} \{g\}^\perp \right) \cap \left(\bigcap_{m \in M} \{m\}^\perp \right) = G^\perp \cap M^\perp \Leftrightarrow M^\perp \subseteq G^\perp$.

(e) Das ist (1.5) in 1.1.18(g)

(f) $y \in M^\perp \setminus (X^\perp) \Leftrightarrow (y \in M^\perp \text{ und } y \notin X^\perp) \Leftrightarrow (\forall x \in M : (x, y) \in R \text{ und } \exists x \in X : (x, y) \notin R) \Rightarrow (\exists x \in X \setminus M : (x, y) \notin R) \Leftrightarrow y \notin (X \setminus M)^\perp \Leftrightarrow y \in Y \setminus ((X \setminus M)^\perp)$.

(g) Nach (a) gilt $\emptyset^\perp = \bigcap_{x \in \emptyset} \{x\}^\perp$, nach 1.1.2 also $\emptyset^\perp = Y$.

(h) Man bemerke nur, dass $G \cap M$ sowohl eine Teilmenge von G als auch von M ist, und wende dann (c) an. Dass $G^\perp \cup M^\perp$ eine echte Teilmenge von $(G \cap M)^\perp$ sein kann, sieht man sofort mit (g) für den Fall $G \cap M = \emptyset$.

(i) $G^\perp \cup M^\perp = (G^\perp \cup M^\perp)^{\perp\perp} \stackrel{(b)}{=} (G^{\perp\perp} \cap M^{\perp\perp})^\perp = (G \cap M)^\perp$.

(j) Sei $x \in \{\xi \in X : (M \cup \{\xi\})^\perp_\perp = M^\perp_\perp\}$. Wegen (b) ist dies äquivalent zu $(M^\perp \cap \{x\}^\perp)_\perp = M^\perp_\perp$. Mit (h) folgt $M^\perp_\perp \cup \{x\}^\perp_\perp \subseteq M^\perp_\perp$, also $x \in M^\perp_\perp$. Sei nun umgekehrt $x \in M^\perp_\perp$. Dann ist die Inklusion $M^\perp \subseteq (M \cup \{x\})^\perp$ zu zeigen. Sei dazu $n \in M^\perp$ und $\xi \in M \cup \{x\}$. Für die dann beiden möglichen Fälle $\xi \in M$ oder $\xi = x$ sieht man sofort, dass jeweils $(\xi, n) \in R$ vorliegt, das heißt, $n \in (M \cup \{x\})^\perp$.

(k) Folgt per (e) sofort aus (j).

(l) Die angegebene Bedingung ist für die Trägerabgeschlossenheit von M notwendig. Sei dazu M trägerabgeschlossen, also $M = M^\perp_\perp$, und sei $x \in X \setminus M$. Also $x \notin M^\perp_\perp$. Das heißt, es gibt ein $y \in M^\perp$ mit $(x, y) \notin R$.

Umgekehrt ist für die Trägerabgeschlossenheit von M die angegebene Bedingung auch hinreichend. Gelte dazu die besagte Bedingung. Ist dann $x \in X \setminus M$, so gibt es also ein $y \in M^\perp$ mit $(x, y) \notin R$, das heißt, es ist $x \notin M^\perp_\perp$. Also hat man $M^\perp_\perp \subseteq M$, und mit (d) folgt die Trägerabgeschlossenheit von M .

(m) Gelte $M \subseteq G \subseteq M^\perp_\perp$. Mit (c) und (e) also $M^\perp = M^{\perp\perp\perp} \subseteq G^\perp \subseteq M^\perp$. $\square$

Gemäß (k) ist also eine Teilmenge M von X genau dann trägerabgeschlossen, wenn es kein von M verschiedenes Element $x \in X$ (mehr) gibt, so dass $M \cup \{x\}$ den gleichen rechten Träger wie M liefern würde. Oder mit anderen Worten: $M \subseteq X$ ist genau dann nicht trägerabgeschlossen, wenn es (noch) ein von M verschiedenes Element $x \in X$ gibt, das man zu der Menge M hinzufügen darf, ohne deren rechten Träger zu verändern.

1.1.20 Satz (THRON [276](1966), Seite 7, Satz 2.2). *Sei $(X, \leq)$ eine prägeordnete Menge. Dann sind die folgenden beiden Aussagen äquivalent:*

(a) *Jede nicht leere Teilmenge von X , welche eine untere Schranke aufweist, besitzt ein Infimum in X .*

(b) *Jede nicht leere Teilmenge von X , welche eine obere Schranke aufweist, besitzt ein Supremum in X .*

Beweis. (a) $\Rightarrow$ (b): Gelte (a). Sei $M \subseteq X$ nicht leer mit einer oberen Schranke. Sei O die Menge aller oberen Schranken von M . (Also $O = M^\perp$.) Die Menge O ist nicht leer und hat (per jedem $m \in M$) eine untere Schranke. (Letzteres ist wegen 1.1.19(d) auch formal klar: $M \subseteq M^\perp_\perp = O^\perp$.) Nach Voraussetzung hat O ein Infimum $o^* = \inf O$. Angenommen, o^* wäre keine obere Schranke von M . Dann würde ein von o^* verschiedenes $m \in M$ existieren mit $o^* \leq m$, also eine untere Schranke von O , die echt größer als o^* wäre; o^* ist aber nach Definition die größte aller unteren Schranken von O , *Widerspruch*. Also ist $o^* \in O$. Da $o^* = \inf O$, ist o^* die kleinste obere Schranke von M , das heißt, $o^* = \sup M$ und es gilt (b).

Vollkommen analog zeigt man (b) $\Rightarrow$ (a): Gelte (b). Sei $M \subseteq X$ nicht leer mit einer unteren Schranke. Sei U die Menge aller unteren Schranken von M . (Also $U = M_\perp$.)

Die Menge U ist nicht leer und hat (per jedem $m \in M$) eine obere Schranke. (Wieder ist letzteres wegen 1.1.19(d) auch formal klar: $M \subseteq M_{\perp}^{\perp} = U^{\perp}$.) Nach Voraussetzung hat U ein Supremum $u^* = \sup U$. Angenommen, u^* wäre keine untere Schranke von M . Dann würde ein von u^* verschiedenes $m \in M$ existieren mit $m \leq u^*$, also eine obere Schranke von U , die echt kleiner als u^* wäre; u^* ist aber nach Definition die kleinste aller oberen Schranken von U , Widerspruch. Also ist $u^* \in U$. Da $u^* = \sup U$, ist u^* die größte untere Schranke von M , das heißt, $u^* = \inf M$ und es gilt (a). $\square$

1.1.21 Satz. Sei $(X, \leq)$ eine prägeordnete Menge. Dann sind die folgenden beiden Aussagen äquivalent:

- (a) Für jede Teilmenge von X existiert das Infimum in X .
- (b) Für jede Teilmenge von X existiert das Supremum in X .

Beweis. (a) $\Rightarrow$ (b): Gelte (a). Insbesondere existiert somit das Infimum der leeren Menge, das heißt, X besitzt ein größtes Element. Damit weist jede nicht leere Teilmenge von X eine obere Schranke auf. Da nach Voraussetzung die Aussage (a) von Satz 1.1.20 gilt, folgt mit eben diesem Satz, dass jede nicht leere Teilmenge von X ein Supremum besitzt. Da des Weiteren nach Voraussetzung das Infimum von X existiert, besitzt X ein kleinstes Element; mithin existiert per diesem das Supremum der leeren Menge, und es gilt (b). Ganz entsprechend zeigt man die Implikation (b) $\Rightarrow$ (a).

Alternativ kann man den Beweis auch direkt, fast wortwörtlich wie im Beweis von Satz 1.1.20, führen: Gelte (a). Sei $M \subseteq X$. Sei O die Menge aller oberen Schranken von M . (Also $O = M^{\perp}$.) Nach Voraussetzung hat O ein Infimum $o^* = \inf O$. Da jedes Element von M eine untere Schranke von O ist, folgt $o^* \in O$ (sonst Widerspruch). Da $o^* = \inf O$, ist o^* die kleinste obere Schranke von M , das heißt, $o^* = \sup M$ und es gilt (b). Die andere Implikation entsprechend. $\square$

1.2 Topologie

1.2.1 Definition (BOURBAKI [32](1966); VON QUERENBURG [244](1979)). Sei X eine Menge und τ eine Menge von Teilmengen von X . Dann heißt τ eine *Topologie auf X* , wenn die folgenden beiden Bedingungen erfüllt sind:

- (a) Jede Vereinigung von Mengen aus τ ist in τ enthalten.
- (b) Jeder Durchschnitt endlich vieler Mengen aus τ ist in τ enthalten.

Ist τ eine Topologie auf X , so heißen die Mengen aus τ *offene Mengen*. Ein *topologischer Raum* ist eine Menge X , versehen mit einer Topologie τ auf X , in Zeichen (X, τ) .

$\tau = \{\emptyset, X\}$ heißt die *indiskrete* oder auch *leere Topologie*. $\tau = \mathcal{P}(X)$ heißt die *diskrete Topologie* auf X und (X, τ) heißt dann ein *diskreter topologischer Raum*.

1.2.2 Topologien und Ordnung. Sei X eine Menge. Sei $\leq$ die auf $\mathcal{P}(X)$ durch Mengeneinklusion $\subseteq$ induzierte kanonische partielle Ordnung — also $A \leq B \Leftrightarrow A \subseteq B$ für alle Teilmengen A, B von X . Offensichtlich ist dann $(\mathcal{P}(X), \leq)$ ein vollständiger Verband.

Sei τ eine Topologie auf X . Mit Satz 1.1.21 in Verbindung mit Definition 1.2.1 (a) sieht man, dass $(\tau, \leq)$ ebenfalls ein vollständiger Verband ist.

Sind $\tau_i, i \in I$, Topologien auf X , dann ist offensichtlich auch $\bigcap \{\tau_i : i \in I\}$ eine Topologie auf X . Betrachtet man also die Menge $\mathcal{T}$ aller Topologien auf X , versehen mit der durch Mengeneinklusion $\subseteq$ induzierten partiellen Ordnung — also $\tau_2 \leq \tau_1 \Leftrightarrow \tau_2 \subseteq \tau_1$ für alle $\tau_1, \tau_2 \in \mathcal{T}$ —, so ist $(\mathcal{T}, \leq)$ nach Satz 1.1.21 ein vollständiger Verband.

1.2.3 G_δ - und F_σ -Mengen (VON QUERENBURG [244](1979)). Sei X ein topologischer Raum. Eine Teilmenge von X heißt eine G_δ -Menge, wenn sie als Durchschnitt abzählbar vieler offener Mengen dargestellt werden kann. Eine Teilmenge von X heißt eine F_σ -Menge, wenn sie als eine Vereinigung von abzählbar vielen abgeschlossenen Mengen dargestellt werden kann.

1.2.4 Grenzpunkte und isolierte Punkte. Sei (X, τ) ein topologischer Raum, $A \subseteq X$ und $x \in X$. Eine *Umgebung* des Punktes x ist eine Teilmenge U von X , die eine offene Menge O umfasst, die x enthält, in Zeichen also $x \in O \subseteq U$. Eine *Umgebung* von A ist eine Teilmenge U von X , die eine offene Menge O umfasst, welche A umfasst, in Zeichen also $A \subseteq O \subseteq U$. Das Wort *Umgebung* wird mit *Umg* abgekürzt. Demgemäß wird *Umgebung von x* als $Umg(x)$ und *Umgebung von A* als $Umg(A)$ geschrieben.

x heißt *innerer Punkt* von A , wenn A eine Umgebung von x ist. Die Menge der inneren Punkte von A heißt das *Innere* (oder auch der *Kern*; im Englischen *interior*) von A und wird mit $\text{int}(A)$ bezeichnet; die dafür sonst auch übliche Bezeichnung A° wird hier nicht verwendet, um Missverständnisse mit der weiter unten definierten Polaren einer Menge zu vermeiden.

Der Punkt x heißt *Berührungspunkt* von A (im Englischen *adherence point of A* oder auch *adherent to A*), wenn jede Umgebung von x mit A einen nicht leeren Durchschnitt hat. Die Menge aller Berührungspunkte von A heißt die *abgeschlossene Hülle* von A oder auch der *Abschluss* von A in X und wird mit $\text{cl}(\tau; A)$ oder mit $\text{cl}(\text{Schlüssel}; A)$ bezeichnet, wobei *Schlüssel* für die Topologie selbst oder für eine abkürzende Bezeichnung der verwendeten Topologie steht. Ist dabei klar, welche Topologie gemeint ist, wird auch einfach nur $\text{cl}(A)$ geschrieben. Besonders in Formeln mit den Allquantor- und Existenz-Operatoren $\forall$ und $\exists$ wird für das Wort *abgeschlossen* die Abkürzung *abg* verwendet. A heißt *dicht bezüglich einer Menge $B \subseteq X$* , wenn $B \subseteq \text{cl}(A)$ gilt, mit anderen Worten, wenn also für jedes $x \in B$ jede Umgebung von x einen nicht leeren Durchschnitt mit A hat. Man sagt dann auch einfach nur, A sei *dicht in B* , dies aber immer im Fall von $A \subseteq B$. X heißt *separabel*, wenn X eine abzählbare, dichte Teilmenge enthält.

x heißt *Randpunkt* von A , wenn x sowohl von A als auch von $X \setminus A$ ein Berührungspunkt ist. Die Menge aller Randpunkte von A heißt der *Rand* (im Englischen *boundary* oder *frontier*) von A und wird mit ∂A bezeichnet. Andere übliche Bezeichnungen für den Rand von A sind $\text{b}(A)$, $\text{Bd}(A)$, $\text{Fr}(A)$, ρA und $\text{Rd}(A)$.

Der Punkt x heißt *Häufungspunkt* (im Englischen *accumulation point*) von A , wenn x ein Berührungspunkt von $A \setminus \{x\}$ ist, in Zeichen $x \in \text{cl}(A \setminus \{x\})$, oder mit anderen Worten, wenn in jeder Umgebung von x mindestens ein von x verschiedener Punkt aus A liegt. Ist X ein T_1 -Raum (siehe 1.2.25), so ist x genau dann ein Häufungspunkt von A , wenn in jeder Umgebung von x unendlich viele Punkte aus A liegen, und in diesem Fall (und bei manchen Autoren auch sonst) heißt x auch ein *Kondensationspunkt* (im Englischen *cluster point*) von A . Die Menge aller Häufungspunkte von A heißt die *erste Ableitung* oder auch die *Derivierte* von A (im Englischen *the derived set of A*) und wird mit A^d bezeichnet. Andere verbreitete, hier aber nicht verwendete, Bezeichnungen für A^d sind A' und $A^{(1)}$.

Sowohl die Berührungspunkte als auch die Häufungspunkte heißen im Englischen manchmal *limit points*.

Der Punkt x heißt *isolierter Punkt* von A , wenn $x \in A \setminus A^d$ gilt, oder äquivalent dazu, wenn es eine Umgebung von x gibt, in der x der einzige Punkt von A ist. Der Punkt x ist genau dann ein isolierter Punkt von X , wenn die Menge $\{x\}$ in X offen ist.

Es gilt $\text{cl}(A) = A \cup A^d$, genauer: $\text{cl}(A) = (A \setminus A^d) \dot{\cup} A^d$. A ist genau dann abgeschlossen, wenn $A^d \subseteq A$ gilt.

Als direkte Konsequenz der Definition der abgeschlossenen Hülle einer Teilmenge

eines topologischen Raumes hat man:

1.2.5 Bemerkung. Sei X mit den Topologien τ_1 und τ_2 jeweils ein topologischer Raum. Sei A eine Teilmenge von X . Dann sind äquivalent:

- (a) $\text{cl}(\tau_1; A) = \text{cl}(\tau_2; A)$.
- (b) A ist τ_1 -abgeschlossen $\Leftrightarrow A$ ist τ_2 -abgeschlossen.
- (c) Für jede Teilmenge S von X gilt: A ist τ_1 -dicht in $S \Leftrightarrow A$ ist τ_2 -dicht in S .

1.2.6 Lokal dicht. Sei X ein topologischer Raum, $A \subseteq X$ und $x_0 \in X$. A heißt *bei x_0 lokal dicht*, wenn eine Umgebung U von x_0 existiert, so dass A dicht in U ist. Mit dieser Terminologie ist somit A genau dann dicht (in X), wenn A bei jedem $x \in X$ lokal dicht ist.

Beweis. Sei A dicht in X . Dann ist offensichtlich $U := X$ für jedes $x \in X$ eine Umgebung von x , so dass A dicht in U ist.

Sei nun umgekehrt A bei jedem $x \in X$ lokal dicht. Für jedes $x \in X$ existiert also eine Umgebung $U(x)$ mit $U(x) \subseteq \text{cl}(A)$. Insbesondere ist also jedes $x \in X$ ein Berührungspunkt von A , das heißt, A ist in X dicht. $\square$

Folglich ist A genau dann nicht dicht (in X), wenn A nicht bei allen $x \in X$ lokal dicht ist. Existiert gar kein Punkt aus X , bei dem A lokal dicht ist, so heißt A *nirgends dicht* (in X). Man überlegt sich leicht, dass A genau dann nirgends dicht (in X) ist, wenn $\text{cl}(A)$ keinen inneren Punkt enthält:

Beweis. Sei A nirgends dicht. *Angenommen*, $\text{cl}(A)$ hätte einen inneren Punkt, etwa x_0 . Dann gäbe es also eine Umgebung $U(x_0)$ von x_0 mit $U(x_0) \subseteq \text{cl}(A)$, das heißt, A wäre in der Umgebung $U(x_0)$ eines Punktes von X , nämlich ja von x_0 , dicht, *Widerspruch*.

Enthalte nun umgekehrt $\text{cl}(A)$ keinen inneren Punkt. *Angenommen*, A wäre nicht nirgends dicht. Dann existierte ein $x_0 \in X$ bei dem A lokal dicht wäre. Es gäbe also eine Umgebung $U(x_0)$ von x_0 , in der A dicht wäre, das heißt, $U(x_0) \subseteq \text{cl}(A)$, *Widerspruch*. $\square$

1.2.7 Definition. Seien X und Y topologische Räume. Sei $f: X \rightarrow Y$ eine Abbildung und $x \in X$. Die Abbildung f heißt *stetig im Punkt x* , wenn zu jeder Umgebung V von $f(x)$ eine Umgebung U von x existiert, so dass gilt: $x \in U$ impliziert $f(x) \in V$. Ist f in jedem Punkt $x \in X$ stetig, so heißt f *stetig*. Die Abbildung f heißt *offen*, wenn das Bild unter f von jeder offenen Menge in X offen in Y ist.

Eine bijektive Abbildung $f: X \rightarrow Y$ zwischen topologischen Räumen heißt ein *Homöomorphismus*, wenn f und f^{-1} stetig sind; in diesem Fall heißen X und Y *homöomorph*.

1.2.8 Halbstetige Abbildungen (KURATOWSKI [192](1966), Seite 173). Seien X und Y zwei topologische Räume. Sei $f: X \rightarrow \mathcal{P}(Y)$ eine Abbildung und $x_0 \in X$.

Die Abbildung f heißt *oberhalbstetig* (auch *nach oben halbstetig*, *von oben halbstetig* oder *halbstetig von oben*; im Englischen *upper semi-continuous*) *im Punkt x_0* , wenn für jede offene Teilmenge G von Y mit $f(x_0) \subseteq G$ der Punkt x_0 ein innerer Punkt von $f^{-1}(\mathcal{P}(G)) = \{x \in X : f(x) \subseteq G\}$ ist. Die Abbildung f ist also genau dann oberhalbstetig bei x_0 , wenn für jede offene Teilmenge G von Y mit $f(x_0) \subseteq G$ eine Umgebung U von x_0 existiert, so dass für alle $x \in U$ die Inklusion $f(x) \subseteq G$ erfüllt ist. Ist f in jedem Punkt von X oberhalbstetig, so heißt f *oberhalbstetig*.

Die Abbildung f heißt *unterhalbstetig* (auch *nach unten halbstetig*; im Englischen *lower semi-continuous*) *im Punkt x_0* , wenn für jede abgeschlossene Teilmenge K von Y , für die x_0 ein Berührungspunkt der Menge $f^{-1}(\mathcal{P}(K)) = \{x \in X : f(x) \subseteq K\}$ ist, die

Inklusion $f(x_0) \subseteq K$ gilt; mit anderen Worten, wenn für jede abgeschlossene Teilmenge K von Y der Punkt x_0 genau dann in der Menge $f^{-1}(\mathcal{P}(K))$ enthalten ist, wenn er ein Berührungspunkt dieser Menge ist. Die Abbildung f ist genau dann unterhalbstetig bei x_0 , wenn für jede offene Teilmenge G von Y mit $f(x_0) \cap G \neq \emptyset$ eine Umgebung U von x_0 existiert, so dass für alle $x \in U$ der Durchschnitt $f(x) \cap G$ nicht leer ist.

Beweis. Sei f unterhalbstetig bei x_0 . Sei G eine offene Teilmenge von Y mit $f(x_0) \cap G \neq \emptyset$. Das heißt, $f(x_0)$ ist keine Teilmenge von $Y \setminus G$ und dies wiederum ist äquivalent zu $x_0 \notin f^{-1}(\mathcal{P}(Y \setminus G))$. Nach Voraussetzung gilt also auch $x_0 \notin \text{cl } f^{-1}(\mathcal{P}(Y \setminus G))$; das heißt, $U := X \setminus \text{cl } f^{-1}(\mathcal{P}(Y \setminus G))$ ist eine Umgebung von x_0 . Nun sieht man leicht, dass für jedes $x \in U$ auch jeweils die umgekehrten Implikationen gelten, insbesondere also $f(x) \cap G \neq \emptyset$.

Existiere nun umgekehrt für jede offene Teilmenge G von Y mit $f(x_0) \cap G \neq \emptyset$ eine Umgebung U von x_0 , so dass für alle $x \in U$ die Ungleichung $f(x) \cap G \neq \emptyset$ gelte. Sei K eine abgeschlossene Teilmenge von Y mit $x_0 \notin f^{-1}(\mathcal{P}(K))$. Betrachte die offene Menge $G := X \setminus K$. Dann gilt also $f(x_0) \not\subseteq K$, also $f(x_0) \cap G \neq \emptyset$. Nach Voraussetzung gibt es eine Umgebung U von x_0 , so dass für alle $x \in U$ die Ungleichung $f(x) \cap G \neq \emptyset$ erfüllt ist. Das heißt, für alle $x \in U$ gilt $f(x) \not\subseteq K$, sprich $x \notin f^{-1}(\mathcal{P}(K))$. Somit ist U eine Umgebung von x_0 , die einen leeren Durchschnitt mit der Menge $f^{-1}(\mathcal{P}(K))$ hat, das heißt, x_0 ist kein Berührungspunkt dieser Menge, in Zeichen $x_0 \notin \text{cl } f^{-1}(\mathcal{P}(K))$. $\square$

Ist f in jedem Punkt von X unterhalbstetig, so heißt f *unterhalbstetig*.

Die Abbildung f ist genau dann unterhalbstetig bei x_0 , wenn für jedes $y \in f(x_0)$ gilt: Für jede Umgebung $V(y)$ von y existiert eine Umgebung $U(x_0)$ von x_0 , so dass für alle $x \in U(x_0)$ der Durchschnitt $f(x) \cap V(y)$ nicht leer ist.

Beweis. Sei f unterhalbstetig bei x_0 . Sei $y \in f(x_0)$ und dann $V(y)$ eine Umgebung von y . Dann sei G eine offene Menge von Y mit $y \in G \subseteq V(y)$. Wegen $y \in f(x_0) \cap G$ also $f(x_0) \cap G \neq \emptyset$. Nach Voraussetzung gibt es dann eine Umgebung $U(x_0)$ von x_0 , so dass für alle $x \in U(x_0)$ die Ungleichung $f(x) \cap G \neq \emptyset$ und somit auch $f(x) \cap V(y) \neq \emptyset$ gilt. Dies zeigt, dass die Unterhalbstetigkeit von f bei x_0 hinreichend ist für die behauptete Eigenschaft.

Sie ist auch notwendig; dazu gelte nun die besagte Eigenschaft. Sei $G \subseteq Y$ offen mit $f(x_0) \cap G \neq \emptyset$. Sei $y \in f(x_0) \cap G$. Dann ist $V(y) := G$ eine Umgebung von y und somit gibt es nach Voraussetzung eine Umgebung $U(x_0)$ von x_0 , so dass für alle $x \in U(x_0)$ die Ungleichung $f(x) \cap V(y) \neq \emptyset$, sprich $f(x) \cap G \neq \emptyset$ gilt. $\square$

Die Abbildung f ist genau dann oberhalbstetig, wenn für jede offene Teilmenge A von Y die Menge $f^{-1}(\mathcal{P}(A))$, also die Menge $\{x \in X : f(x) \subseteq A\}$, offen ist; dies ist genau dann der Fall, wenn für jede abgeschlossene Teilmenge A von Y die Menge $X \setminus f^{-1}(\mathcal{P}(X \setminus A))$, also die Menge $\{x \in X : f(x) \cap A \neq \emptyset\}$, sprich $f^{-1}(\mathcal{P}(A))$, abgeschlossen ist; unterhalbstetig entsprechend mit „offen“ jeweils vertauscht mit „abgeschlossen“.

Siehe auch 1.2.9 und 2.4.40.

1.2.9 Halb stetige Abbildungen nach geordneten Räumen (CHOQUET [45](1969), Seite 29). Sei X ein topologischer Raum. Sei Y eine Menge, versehen mit einer Totalordnung. Sei $x_0 \in X$. Eine Abbildung $f: X \rightarrow Y$ heißt *unterhalbstetig* (auch *nach unten halb stetig*; im Englischen *lower semi-continuous*) *im Punkt* x_0 , wenn für jedes $y \in Y$ mit $f(x_0) > y$ eine Umgebung U von x_0 existiert, so dass für alle $x \in U$ die Ungleichung $f(x) > y$ gilt. *Oberhalbstetig* (oder auch synonym *nach oben halb stetig*) *im Punkt* x_0 wird entsprechend mit jeweils $<$ anstatt $>$ definiert. Falls eine Abbildung $f: X \rightarrow Y$ in jedem Punkt von X unterhalbstetig ist, heißt sie *unterhalbstetig*. *Oberhalbstetig* entsprechend.

Sei $f: X \rightarrow Y$ eine Abbildung. Dann ist f genau dann unterhalbstetig, wenn für jedes $y \in Y$ die Menge $\{x \in X : f(x) > y\}$ offen ist.

Beweis. Sei f unterhalbstetig. Sei $y \in Y$. Setze $A := \{x \in X : f(x) > y\}$. Liege der Fall vor, dass A nicht leer ist. Sei $x_0 \in A$ beliebig, aber fest. Nach Voraussetzung existiert eine Umgebung U von x_0 mit $U \subseteq A$, das heißt, x_0 ist ein innerer Punkt von A .

Sei nun umgekehrt für jedes $y \in Y$ die Menge $\{x \in X : f(x) > y\}$ offen. Sei $x_0 \in X$. Sei $y \in Y$ mit $f(x_0) > y$. Die Menge $A := \{x \in X : f(x) > y\}$ enthält also x_0 und ist nach Voraussetzung offen. Es existiert also eine offene Menge $U \subseteq A$ mit $x_0 \in U$. $\square$

Ganz analog gilt: f ist genau dann oberhalbstetig, wenn für jedes $y \in Y$ die Menge $\{x \in X : f(x) < y\}$ offen ist.

Siehe auch 2.9.32.

1.2.10 Grenzwert einer Funktion. Seien X und Y topologische Räume, sei $A \subseteq X$, $f: A \rightarrow Y$ eine Abbildung, $a \in X$ ein Berührungspunkt von A und $y \in Y$. Dann heißt das y ein *Grenzwert* von f bezüglich A im Punkt $a \in \text{cl}(A)$, in Zeichen $\lim_{x \rightarrow a, x \in A} f(x) = y$, wenn zu jeder Umgebung V von y in Y derart eine Umgebung U von a in X existiert, dass $f(U \cap A) \subseteq V$ erfüllt ist. Ist g die durch $g|_A := f$, $g(a) := y$ definierte Abbildung $A \cup \{a\} \rightarrow Y$, so gilt genau dann $g(a) = \lim_{x \rightarrow a, x \in A} g(x)$ (was nach Definition von g äquivalent zu $y = \lim_{x \rightarrow a, x \in A} f(x)$ ist), wenn g in a stetig ist.

Somit gilt: Seien X und Y topologische Räume, sei $g: X \rightarrow Y$ eine Abbildung und $a \in X$ ein Berührungspunkt von $X \setminus \{a\}$. Dann ist die Abbildung g genau dann in a stetig, wenn $g(a) = \lim_{x \rightarrow a, x \in X \setminus \{a\}} g(x)$ gilt.

Die hier gegebene Definition des Grenzwertes einer Funktion findet sich zum Beispiel bei den Autoren BOURBAKI, DIEUDONNÉ, SERGE LANG und AMANN und ESCHER. Bei manchen Autoren wie zum Beispiel RUDIN, HEUSER und KÖNIGSBERGER findet man eine andere Definition des Grenzwertes einer Funktion $f: A \rightarrow Y$: Es wird $a \in X$ als Häufungspunkt von A vorausgesetzt und der Grenzwert definiert als $\lim_{x \rightarrow a} f(x) := \lim_{x \rightarrow a, x \in A \setminus \{a\}} f(x)$; diese Variante ist die klassische Definition, die auf Weierstrass zurückgeht.

1.2.11 Definition. Sei (X, τ) ein topologischer Raum. Wenn B eine Familie von Umgebungen eines $x \in X$ (bzw. einer Teilmenge A von X) ist und für jede Umgebung V von x (bzw. von A) ein $U \in B$ mit $U \subseteq V$ existiert, dann heißt B eine *Umgebungsbasis* von x (bzw. von A) (im Englischen auch *fundamental system of neighborhoods* of x (bzw. of A)) und wird mit $\mathfrak{B}(x)$ (bzw. $\mathfrak{B}(A)$) bezeichnet; genauer schreibt man auch $\mathfrak{B}(\tau; x)$ (bzw. $\mathfrak{B}(\tau; A)$). Falls für jedes $x \in X$ eine abzählbare Umgebungsbasis von x existiert, sagt man, X erfülle das *erste Abzählbarkeitsaxiom*.

Ein $\beta \subseteq \mathcal{P}(X)$ heißt eine *Basis der Topologie* τ , falls $\beta \subseteq \tau$ und jedes Element von τ eine Vereinigung von Elementen von β ist: $\forall O \in \tau \quad \exists \gamma \subseteq \beta : O = \bigcup \gamma$. Falls X eine abzählbare Basis besitzt, sagt man, X erfülle das *zweite Abzählbarkeitsaxiom*. Ein $\sigma \subseteq \mathcal{P}(X)$ heißt eine *Subbasis der Topologie* τ , wenn das aus allen endlichen Durchschnitten von Elementen aus σ gebildete Mengensystem eine Basis von τ ist.

1.2.12 Definition. Sei (X, τ) ein topologischer Raum und $x \in X$. Ein $S \subseteq \mathcal{P}(X)$ heißt eine *Umgebungssubbasis* von x , falls $\{\bigcap F : F \subseteq S \text{ endlich}\}$ eine Umgebungsbasis von x ist.

1.2.13. Sei (X, τ) ein topologischer Raum. Für $\beta \subseteq \mathcal{P}(X)$ sind äquivalent:

- (a) β ist eine Basis von τ .
- (b) $\beta \subseteq \tau$ und $\forall x \in X \quad \forall V \text{ Umg}(x) \quad \exists U \in \beta : x \in U \subseteq V$.
- (c) $\beta \subseteq \tau$ und $\forall x \in X : \{O \in \beta : x \in O\}$ ist eine Umgebungsbasis von x .

Beweis. (a) $\Rightarrow$ (b): Gelte (a). Sei $x \in X$ und sei V eine Umgebung von x . Es gibt also ein $W \in \tau$ mit $x \in W \subseteq V$. Nach Voraussetzung existiert ein $\gamma \subseteq \beta$ mit $W = \bigcup \gamma$, insbesondere also ein $U \in \gamma$ mit $x \in U$, also $x \in U \subseteq V$.

(b) $\Rightarrow$ (a): Gelte (b). Sei $V \in \tau$. Setze $\gamma := \{O \in \beta : O \subseteq V\}$. Sei $x \in V$. Nach Voraussetzung existiert ein $U \in \beta$ mit $x \in U \subseteq V$. Somit ist $U \in \gamma$ und $x \in \bigcup \gamma$, also $V \subseteq \bigcup \gamma$. Da offensichtlich $\bigcup \gamma \subseteq V$, also $V = \bigcup \gamma$, das heißt, es gilt (a).

(b) $\Rightarrow$ (c): Gelte (b). Sei $x \in X$ und V eine Umgebung von x . Dann existiert nach Voraussetzung ein $U \in \beta$ mit $x \in U \subseteq V$, das heißt, ein Element U aus der Menge $\{O \in \beta : x \in O\}$ mit $U \subseteq V$ und es gilt (c).

(c) $\Rightarrow$ (b): Gelte (c). Sei $x \in X$ und V eine Umgebung von x . Nach Voraussetzung existiert ein Element U aus der Menge $\{O \in \beta : x \in O\}$ mit $U \subseteq V$, das heißt, $x \in U \subseteq V$ und es gilt (b).

Obwohl damit die Äquivalenz von (a), (b) und (c) gezeigt ist, sei hier noch der direkte Beweis der Äquivalenz von (a) und (c) angefügt:

(a) $\Rightarrow$ (c): Gelte (a). Wörtlich wie bei (a) $\Rightarrow$ (b) und dann beachte man noch, dass U ein Element von $\{O \in \beta : x \in O\}$ ist.

(c) $\Rightarrow$ (a): Gelte (c). Sei $O \in \tau$. Nach Voraussetzung existiert für jedes $x \in O$ ein $U_x \in \beta$ mit $x \in U_x \subseteq O$. Somit ist $O = \bigcup \{U_x \in \beta : x \in O\}$, das heißt, es gilt (a). $\square$

1.2.14. Sei X eine Menge. $\beta \subseteq \mathcal{P}(X)$ ist genau dann eine Basis einer Topologie auf X , wenn $X = \bigcup \beta$ und für alle endlichen $F \subseteq \beta$ der Durchschnitt $\bigcap F$ Vereinigung von Mengen aus β oder leer ist.

1.2.15 Satz (KELLEY [184](1955), Seite 47). *Eine Familie β von Mengen ist genau dann eine Basis für eine Topologie auf $\bigcup \beta$, wenn $\forall U, V \in \beta \quad \forall x \in U \cap V \quad \exists W \in \beta : x \in W \subseteq U \cap V$.*

1.2.16 Satz und Definition (PREUSS [243](1975), Seite 41, Lemma 1.3.19). Sei X eine Menge und σ eine nicht leere Teilmenge von $\mathcal{P}(X)$. Dann ist σ eine Subbasis einer eindeutig bestimmten Topologie auf X .

Beweis. Setze $\beta = \{ \bigcap F : F \subseteq \sigma \text{ endlich} \}$. Nach 1.1.2 ist $X \in \beta$, also $\bigcup \beta = X$. Der Durchschnitt von je zwei Elementen von β ist wieder ein Element von β . Nach Satz 1.2.15 ist β eine Basis einer Topologie τ auf X , mithin σ eine Subbasis von τ . Da eine Topologie von X mit σ als Subbasis stets β als eine Basis besitzt, ist die Topologie τ eindeutig bestimmt. $\square$

Die demonstrierte Topologie τ wird *die von σ erzeugte Topologie* genannt, in Zeichen $\tau = (\sigma)$.

1.2.17. Sei X eine Menge und seien τ_1 und τ_2 zwei Topologien auf X . τ_1 heißt *feiner* als τ_2 , und τ_2 *gröber* als τ_1 , wenn $\tau_1 \supseteq \tau_2$ gilt. Sei für $i = 1, 2$ β_i eine Basis für τ_i und σ_i eine Subbasis für τ_i . Sei für $x \in X$ und $i = 1, 2$ $\mathfrak{B}(\tau_i; x)$ eine Umgebungsbasis von x in (X, τ_i) . Folgende Aussagen sind äquivalent:

- (a) τ_1 ist feiner als τ_2 .
- (b) Die Abbildung $(X, \tau_1) \rightarrow (X, \tau_2), x \mapsto x$ ist stetig.
- (c) Es gibt eine Subbasis von τ_2 , die Teilmenge von τ_1 ist.
- (d) $\forall x \in X \quad \forall U \in \mathfrak{B}(\tau_2; x) : U \text{ ist Umg}(\tau_1; x)$.
- (e) Jede τ_2 -Umgebung ist eine τ_1 -Umgebung.
- (f) $\forall x \in X \quad \forall V \in \mathfrak{B}(\tau_2; x) \quad \exists U \in \mathfrak{B}(\tau_1; x) : U \subseteq V$.

- (g) $\forall A \subseteq X : \text{cl}(\tau_1; A) \subseteq \text{cl}(\tau_2; A)$.
- (h) $\forall A \subseteq X \tau_2\text{-abg} : A \text{ ist } \tau_1\text{-abg}$.
- (i) $\forall O_2 \in \beta_2 \quad \forall x \in O_2 \quad \exists O_1 \in \beta_1 : x \in O_1 \subseteq O_2$.
- (j) $\forall O_2 \in \sigma_2 \quad \forall x \in O_2 \quad \exists O_1 \in \beta_1 : x \in O_1 \subseteq O_2$.
- (k) $\forall A \subseteq X : \text{int}(\tau_1; A) \supseteq \text{int}(\tau_2; A)$.
- (l) $\forall A \subseteq X \quad \forall B \subseteq X : \left(A \text{ } \tau_1\text{-dicht bzgl. } B \Rightarrow A \text{ } \tau_2\text{-dicht bzgl. } B \right)$.

Beweis. (a) $\Rightarrow$ (i): Gelte (a). Sei $O_2 \in \beta_2$. Sei $x \in O_2$. Nach Voraussetzung ist $O_2 \in \tau_1$ und daher existiert ein $\gamma_1 \subseteq \beta_1$ mit $O_2 = \bigcup \gamma_1$, insbesondere ein $O_1 \in \gamma_1$ mit $x \in O_1$, also $x \in O_1 \subseteq O_2$ und $O_1 \in \beta_1$, das heißt, es gilt (i).

(i) $\Rightarrow$ (a): Gelte (i). Sei $U_2 \in \tau_2$. Dann existiert ein $\gamma_2 \subseteq \beta_2$ mit $U_2 = \bigcup \gamma_2$. Sei $O_2 \in \beta_2$. Nach Voraussetzung existiert für jedes $x \in O_2$ ein $O_x \in \beta_1$ mit $x \in O_x \subseteq O_2$. Also ist $O_2 = \bigcup \{O_x \in \beta_1 : x \in O_x\}$. Folglich lässt sich jedes $O_2 \in \beta_2$, insbesondere jedes $O_2 \in \gamma_2$ und damit auch U_2 , als Vereinigung von Elementen aus β_1 darstellen, das heißt, $U_2 \in \tau_1$ und es gilt (a). $\square$

1.2.18 Subbasis. Sei (X, τ) ein topologischer Raum. Für $\sigma \subseteq \mathcal{P}(X)$ sind äquivalent:

- (a) σ ist eine Subbasis von τ . (Definition 1.2.11)
- (b) $\{ \bigcap F : F \subseteq \sigma \text{ endlich} \}$ ist eine Basis von τ .
- (c) $\sigma \subseteq \tau$ und jedes Element von τ ist die Vereinigung von endlichen Schnitten von Elementen aus σ .
- (d) τ ist die von σ erzeugte Topologie. (Satz und Definition 1.2.16)
- (e) τ ist die größte Topologie auf X mit $\sigma \subseteq \tau$.

1.2.19. Seien (X_1, τ_1) und (X_2, τ_2) zwei topologische Räume. Ist $f: X_1 \rightarrow X_2$ eine Abbildung und $x \in X_1$. Sei für $i = 1, 2$ β_i eine Basis für τ_i und σ_i eine Subbasis für τ_i . Für $i = 1, 2$ und $x_i \in X_i$ bezeichne $\mathfrak{B}(\tau_i; x_i)$ eine Umgebungsbasis und $\mathfrak{S}(\tau_i, x_i) \subseteq \mathfrak{B}(\tau_i; x_i)$ eine Umgebungssubbasis von x_i in X_i . Dann gilt:

(a) Die Abbildung f ist genau dann stetig, wenn für alle $O_2 \in \sigma_2$ ein $O_1 \in \tau_1$ mit $f(O_1) \subseteq O_2$ existiert.

(b) Die Abbildung f ist genau dann in x stetig, wenn für alle $V \in \mathfrak{S}(\tau_2; f(x))$ ein $U \in \mathfrak{B}(\tau_1; x)$ mit $f(U) \subseteq V$ existiert.

Beweis. Dass für ein in x stetiges f die angegebene Bedingung folgt, ist klar. Sei nun umgekehrt die angegebene Bedingung erfüllt. Sei V eine Umgebung von $f(x)$. Dann existieren für ein $n \in \mathbb{N}^\times$ $V_1, \dots, V_n \in \mathfrak{S}(\tau_2; f(x))$ mit $V_1 \cap \dots \cap V_n \subseteq V$. Nach Voraussetzung existiert zu jedem $V_i, i = 1, \dots, n$, ein $U_i \in \mathfrak{B}(\tau_1; x)$ mit $f(U_i) \subseteq V_i$. Damit gilt: $f(U_1 \cap \dots \cap U_n) \subseteq f(U_1) \cap \dots \cap f(U_n) \subseteq V_1 \cap \dots \cap V_n \subseteq V$. $\square$

1.2.20 Definition. Sei X ein topologischer Raum. X ist *kompakt*, wenn jede offene Überdeckung von X eine endliche Teilüberdeckung enthält. $A \subseteq X$ ist *kompakt*, wenn der Unterraum A kompakt ist. $A \subseteq X$ ist *relativ kompakt* (bei DUNFORD und SCHWARTZ [78](1958) im Englischen *conditionally compact*), wenn $\text{cl}(A)$ kompakt ist. X ist *abzählbar kompakt*, wenn jede abzählbare offene Überdeckung von X eine endliche Teilüberdeckung enthält. X ist ein *Lindelöf-Raum*, wenn jede offene Überdeckung von X eine abzählbare Teilüberdeckung enthält. X ist *lokal kompakt*, wenn jeder Punkt von X mindestens eine

kompakte Umgebung hat. X ist σ -kompakt, wenn X lokal kompakt ist und als eine Vereinigung von höchstens abzählbar vielen kompakten Räumen betrachtet werden kann. Siehe auch die Bemerkungen 1.2.28 weiter unten.

1.2.21 Definition. Seien X_ν , $\nu \in I$, I eine Indexmenge, topologische Räume. Sei M eine Menge. Es sei an Bemerkung 1.2.2 erinnert.

(a) Für jedes $\nu \in I$ sei $f_\nu: M \rightarrow X_\nu$ eine Abbildung. Die grösste Topologie auf M bezüglich der alle f_ν , $\nu \in I$, stetig sind, heißt die *Initialtopologie* (von M) bezüglich (der topologischen Räume X_ν , $\nu \in I$, und) der Familie f_ν , $\nu \in I$.

(b) Für jedes $\nu \in I$ sei $g_\nu: X_\nu \rightarrow M$ eine Abbildung. Die feinste Topologie auf M bezüglich der alle g_ν , $\nu \in I$, stetig sind, heißt die *Finaltopologie* (von M) bezüglich (der topologischen Räume X_ν , $\nu \in I$, und) der Familie g_ν , $\nu \in I$.

1.2.22 Produkttopologie (CAROTHERS [43](2005), Seite 170). Seien X_ν , $\nu \in I$, I eine Indexmenge, topologische Räume. Die *Tychonoff-Topologie* oder *Produkttopologie* auf dem kartesischen Produkt $\prod_{\nu \in I} X_\nu$, aufgefasst als die Menge aller mit I indizierten Familien $x = (x_\nu)_{\nu \in I}$ mit $x_\nu \in X_\nu$ für alle $\nu \in I$, ist die Initialtopologie der Funktionenfamilie aller kanonischen Abbildungen $\pi_\mu: \prod_{\nu \in I} X_\nu \rightarrow X_\mu$, $(x_\nu)_{\nu \in I} \mapsto x_\mu$, $\mu \in I$. Gemäß 1.2.18 bilden die $\pi_\nu^{-1}(U_\nu)$, $\nu \in I$, U_ν offen in X_ν , somit eine Subbasis im Produkt. Ist von einem kartesischen Produkt $\prod_{\nu \in I} X_\nu$ als topologischem Raum die Rede, so ist, falls keine andere Topologie als die Produkttopologie näher spezifiziert ist, der Raum als mit der Produkttopologie ausgestattet zu verstehen.

1.2.23 Würfel. Bezeichne I das abgeschlossene Einheitsintervall $[0, 1]$ der reellen Zahlengerade und sei J eine beliebige Indexmenge. Dann heißt $I^J := \prod_{\nu \in J} I_\nu$ mit $I_\nu = I$ für alle $\nu \in J$, versehen mit der Produkttopologie, ein *Würfel* (im Englischen *cube*) oder ein *Parallelotop*. Es ist verbreitet, für J eine Kardinalzahl $\mathfrak{m}$ anzunehmen. Speziell nennt man einen Würfel $I^{\mathfrak{m}}$ einen *Tychonoff-Würfel*, wenn $\mathfrak{m}$ eine unendliche Kardinalzahl ist, und man nennt ihn einen *euklidischen Würfel*, wenn $\mathfrak{m}$ eine endliche Kardinalzahl ist.

1.2.24 Filter. In 1.1.16 wurden Filter auf Mengen X definiert. Sei nun X ein topologischer Raum. Die Menge aller Umgebungen einer Menge $A \subseteq X$ ist ein Filter; er wird der *Umgebungsfilter von A* genannt; bei einelementigem A , etwa $A = \{x\}$, sagt man einfach *Umgebungsfilter von x* . Für jedes $x \in X$ ist der Umgebungsfilter von x ein fixierter Filter. Man sagt, ein Filter $\mathcal{F}$ auf X *konvergiere gegen $x \in X$* , in Zeichen $\mathcal{F} \rightarrow x$, wenn $\mathcal{F}$ feiner als der Umgebungsfilter von x ist; x heißt dann ein *Limespunkt* von $\mathcal{F}$. Ist dabei x der einzige Limespunkt von $\mathcal{F}$, so schreibt man auch $\lim \mathcal{F} = x$. Diese Konvergenzbegriffe und -symbole übertragen sich auf Netze, indem man deren Abschnittsfilter betrachtet.

Man sagt, dass eine indizierte Familie $x_\nu \in X$, $\nu \in I$, *nach x bezüglich eines Filters $\mathcal{F}$ konvergiere*, in Zeichen $x = \lim_{\mathcal{F}} x_\nu$, falls für jede offene Menge $G \subseteq X$, die x enthält, die Menge $\{\nu \in I : x_\nu \in G\}$ zu $\mathcal{F}$ gehört (JOHNSON und LINDENSTRAUSS [158](2001), Seite 55).

Die folgenden Trennungsaxiome und Äquivalenzen entstammen den Topologie-Büchern von BOTO VON QUERENBURG, KELLEY, KURATOWSKI und DUGUNDJI.

1.2.25 Trennungsaxiome. Sei $(X; \tau)$ ein topologischer Raum. X heißt ein

(a) T_0 -Raum

$:\Leftrightarrow$ Für je zwei verschiedene Punkte von X , hat mindestens einer der beiden Punkte eine Umgebung, die nicht den anderen Punkt enthält.

$\Leftrightarrow \forall x, y \in X, x \neq y : x \notin \text{cl}(\{y\}) \text{ oder } y \notin \text{cl}(\{x\}).$

$\Leftrightarrow \forall x, y \in X, x \neq y : \text{cl}(\{x\}) \neq \text{cl}(\{y\}).$

(b) T_1 -Raum

- $:\Leftrightarrow$ Für je zwei verschiedene Punkte von X , hat jeder der beiden Punkte eine Umgebung, die nicht die Umgebung des anderen Punktes enthält.
- $\Leftrightarrow$ Jede einpunktige Menge ist abgeschlossen.
- $\Leftrightarrow$ Jede Teilmenge $A \subseteq X$ ist der Durchschnitt aller ihrer Umgebungen.

(c) T_2 -Raum oder Hausdorffraum

- $:\Leftrightarrow \forall p, q \in X, p \neq q \quad \exists G, H \subseteq X$ offen, disjunkt: $p \in G, q \in H$.
- $\Leftrightarrow$ Für jeden Punkt $x \in X$ ist der Durchschnitt all seiner abgeschlossenen Umgebungen gleich der Menge $\{x\}$.
- $\Leftrightarrow$ Die Diagonale von X ist in $X \times X$ abgeschlossen.
- $\Leftrightarrow$ Jeder konvergente Filter auf X besitzt genau einen Limespunkt.
- $\Leftrightarrow$ Jedes Netz konvergiert zu höchstens einen Punkt.

(d) T_3 -Raum oder regulärer Raum

- $:\Leftrightarrow \forall p \in X \quad \forall F \subseteq X$ abg, $p \notin F \quad \exists G_1, G_2 \subseteq X$ offen, disjunkt:
 $p \in G_1, F \subseteq G_2$.
- $\Leftrightarrow \forall p \in X \quad \forall F \subseteq X$ abg, $p \notin F \quad \exists G \subseteq X$ offen:
 $p \in G, cl(G) \cap F = \emptyset$.
- $\Leftrightarrow \forall p \in X \quad \forall G$ Umg(p) $\exists F$ Umg(p), abg: $F \subseteq G$.
- $\Leftrightarrow \forall p \in X$: Die Menge aller abgeschlossenen Umgebungen von p ist eine Umgebungsbasis von p .

(e) T_{3a} -Raum oder vollständig regulärer Raum

- $:\Leftrightarrow \forall p \in X \quad \forall F \subseteq X$ abg, $p \notin F \quad \exists f: X \rightarrow [0, 1]$ stetig, $f(p) = 0 \quad \forall x \in F: f(x) = 1$.
- $\Leftrightarrow$ Die Topologie von X besitzt als Basis das Mengensystem $\{f^{-1}(U) : U \subseteq \mathbb{R}, f: X \rightarrow \mathbb{R} \text{ stetig}\}$.
- $\Leftrightarrow$ Die Nullstellenmengen der stetigen reellwertigen Funktionen auf X bilden eine Basis für die abgeschlossenen Mengen von X .
- $\Leftrightarrow \forall p \in X \quad \forall U$ Umg(p) $\exists f: X \rightarrow [0, 1]$ stetig, $f(p) = 0$
 $\forall x \in X \setminus U: f(x) = 1$.

(f) Tychonoff-Raum

- $:\Leftrightarrow X$ ist T_1 und T_{3a} .
- $\Leftrightarrow X$ ist homöomorph zu einem Unterraum eines Würfels.
- $\Leftrightarrow X$ ist homöomorph zu einem Unterraum eines kompakten Hausdorff-raumes.

(g) T_4 -Raum oder normaler Raum

$$\begin{aligned}
 & :\Leftrightarrow \forall F_0, F_1 \subseteq X \text{ abg, disjunkt} \quad \exists G_0, G_1 \subseteq X \text{ offen, disjunkt:} \\
 & \quad F_0 \subseteq G_0, F_1 \subseteq G_1. \\
 & \Leftrightarrow \forall G_0, G_1 \subseteq X \text{ offen, } X = G_0 \cup G_1 \quad \exists F_0, F_1 \subseteq X \text{ abg:} \\
 & \quad F_0 \subseteq G_0, F_1 \subseteq G_1, F_0 \cup F_1 = X. \\
 & \Leftrightarrow \forall F_0, F_1 \subseteq X \text{ abg, disjunkt, } \neq \emptyset \quad \exists f: X \rightarrow [0, 1] \text{ stetig:} \\
 & \quad f(F_0) = \{0\}, f(F_1) = \{1\}. \\
 & \Leftrightarrow \forall F \subseteq X \text{ abg, } \forall G \subseteq X \text{ offen, } F \subseteq G \quad \exists A \subseteq X \text{ offen:} \\
 & \quad F \subseteq A \subseteq \text{cl}(A) \subseteq G.
 \end{aligned}$$

1.2.26 Bemerkung. Manche Autoren verlangen bei der Definition von den T_3 -, T_{3a} - und T_4 -Räumen zusätzlich, dass sie T_1 - oder T_2 -Räume sind.

1.2.27 Satz. Sei X ein topologischer Raum. Dann gelten die folgenden Implikationen: $T_1 + T_4 \Rightarrow T_1 + T_{3a}$, $T_{3a} \Rightarrow T_3$, $T_1 + T_3 \Rightarrow T_2 \Rightarrow T_1 \Rightarrow T_0$.

1.2.28 Bemerkungen. (a) Manche Autoren verlangen bei der Definition von *kompakt*, dass X ein Hausdorffraum ist. Üblicherweise wird dann das hier definierte *kompakt* als *präkompakt*, *bikompakt* oder *quasikompakt* bezeichnet. Bei diesen Autoren wird dann meist auch für die Definitionen von *abzählbar kompakt* und *lokal kompakt* X als ein Hausdorffraum gefordert.

(b) Man beachte, dass eine kompakte Teilmenge eines topologischen Raumes X nur dann mit Sicherheit abgeschlossen ist, wenn X Hausdorff'sch ist.

(c) Es gibt Autoren, die den Begriff *relativ kompakt* etwas schwächer definieren, in dem sie eine Teilmenge A eines topologischen Raumes X *relativ kompakt* nennen, wenn jede offene Überdeckung von X eine endliche Teilüberdeckung von A enthält; mit dieser Definition lässt sich dann zum Beispiel ein Hausdorffraum X angeben, der eine in diesem Sinne relativ kompakte Teilmenge enthält, die keine kompakte Obermenge in X hat. Siehe zum Beispiel BARTSCH [14](2005).

1.2.29 Satz. Sei X ein topologischer Raum. Dann gelten die folgenden Implikationen: $T_2 + \text{lokal kompakt} \Rightarrow T_1 + T_{3a}$, $T_2 + \text{kompakt} \Rightarrow T_1 + T_4$, $T_2 + \text{lokal kompakt} + \sigma\text{-kompakt} \Rightarrow T_1 + T_4$, $T_3 + \text{Lindelöf} \Rightarrow T_4$.

1.2.30 Zusammenhang (VON QUERENBURG [244](1979); JÄNICH [154](2004)). Sei X ein topologischer Raum. X heißt *zusammenhängend*, wenn X nicht in zwei disjunkte, nicht leere, offene Mengen zerlegt werden kann. Eine stetige Abbildung $[0, 1] \rightarrow X$ heißt ein *Weg* von X . Einen Weg f mit $f(0) = f(1)$ nennt man *geschlossen*. X heißt *wegzusammenhängend*, wenn es zu je zwei Punkten $x, y \in X$ einen Weg f mit $f(0) = x$ und $f(1) = y$ gibt. X heißt *lokal (weg)zusammenhängend*, wenn es zu jedem Punkt $x \in X$ und zu jeder Umgebung U von x eine (weg)zusammenhängende Umgebung V von x mit $V \subseteq U$ gibt.

Zwei Wege f und g von X mit gemeinsamen Anfangspunkt x und gemeinsamen Endpunkt y in X heißen *homotop*, in Zeichen $f \sim g$, wenn es zwischen ihnen eine *Homotopie* gibt, das heißt, eine stetige Abbildung $h: [0, 1] \times [0, 1] \rightarrow X$ mit der Eigenschaft, dass für die einzelnen Wege $h_\lambda: [0, 1] \rightarrow X$, $t \mapsto h(t, \lambda)$, $\lambda \in [0, 1]$, gilt: $h_0 = f$ und $h_1 = g$, und alle h_λ haben den Anfangspunkt x und den Endpunkt y .

Ein geschlossener Weg f heißt *nullhomotop*, wenn $f \sim ([0, 1] \rightarrow X, t \mapsto f(0))$ gilt. X heißt *einfach zusammenhängend*, wenn X wegzusammenhängend ist und alle geschlossenen Wege nullhomotop sind.

1.2.31. Jeder wegzusammenhängende Raum ist zusammenhängend.

Ein topologischer Raum X ist genau dann zusammenhängend, wenn jede stetige Abbildung von X in einen diskreten topologischen Raum, der mindestens zwei Punkte enthält, konstant ist (VON QUERENBURG [244](1979), Seite 50, A4.3).

Die leere Menge und jede einelementige Menge eines topologischen Raumes ist zusammenhängend (BOURBAKI [32](1966), I.11.1).

Ist X ein Hausdorffraum, so ist jede endliche Teilmenge von X mit mehr als einem Element nicht zusammenhängend; allgemeiner ist jede Teilmenge von X mit mehr als einem Element und mindestens einem isolierten Punkt nicht zusammenhängend (BOURBAKI [32](1966), I.11.1).

Kapitel 2

Algebren und topologische Vektorräume

2.1 Ringe

2.1.1 Definition (SCHULZE [258](2006); GOLDHABER und EHRLICH [115] (1970)). Seien X und Y zwei Mengen. Eine Abbildung $f: Y \times X \rightarrow X$ heißt eine *Verknüpfung auf X* (genauer eine *zweistellige Verknüpfung auf X*). Im Fall von $Y \neq X$ heißt f eine *äußere Verknüpfung auf X mit Operatorenbereich Y* . Im Fall von $Y = X$ heißt f eine *innere Verknüpfung auf X* . Wie üblich wird für $x \in X$, $y \in Y$ das Verknüpfungsergebnis $f(y, x)$ meist als yfx geschrieben, was wiederum bei Verwendung einer multiplikativen Schreibweise auch einfach nur als yx geschrieben wird.

Sei $(X, \circ)$ eine Menge X mit einer zweistelligen Verknüpfung $\circ$. Seien A und B zwei Teilmengen von X . Dann wird mit $A \circ B$ die Menge $\{a \circ b \in X : a \in A, b \in B\}$ bezeichnet; sie heißt das *Komplexprodukt* von A , B . Ist eine der Mengen A oder B einelementig, etwa $A = \{a\}$, dann schreibt man $a \circ B$ anstelle von $\{a\} \circ B$. Man beachte 2.1.13.

Sei X eine nicht leere Menge, I eine nicht leere Indexmenge und $V := \{f_\nu : \nu \in I\}$ eine nicht leere Menge von nicht notwendig zweistelligen Verknüpfungen auf X . Dann heißt (X, V) eine *algebraische Struktur*. Für endliches V , etwa $V = \{f_1, \dots, f_n\}$, schreibt man für (X, V) auch $(X, f_1, \dots, f_n)$. Ist $(A, \circ)$ eine algebraische Struktur, deren Verknüpfung $\circ$ eine innere zweistellige Verknüpfung ist, so heißt $(A, \circ)$ ein *Magma* (im Englischen *magma*; bei KUROSCHE [193](1964) ein *Gruppoid*).

Sei $(M, \circ)$ ein Magma. Die Verknüpfung $\circ$ heißt *assoziativ*, wenn für alle $x, y, z \in M$ die Gleichung $(x \circ y) \circ z = x \circ (y \circ z)$ gilt — die beiden Ausdrücke $(x \circ y) \circ z$ und $x \circ (y \circ z)$ also auch klammerfrei als $x \circ y \circ z$ geschrieben werden dürfen; in diesem Fall heißt $(M, \circ)$ eine *Halbgruppe* (im Englischen *semigroup*). Die Verknüpfung $\circ$ heißt *kommutativ*, wenn für alle $x, y \in M$ die Gleichung $x \circ y = y \circ x$ gilt; in diesem Fall heißt $(M, \circ)$ *kommutativ* oder auch *abelsch*. Ein $\ell \in M$ heißt ein *linksneutrales Element* oder eine *Linkseins* von $(M, \circ)$, wenn $\ell \circ a = a$ für alle $a \in M$ gilt. Ein $r \in M$ heißt ein *rechtsneutrales Element* oder eine *Rechtseins* von $(M, \circ)$, wenn $a \circ r = a$ für alle $a \in M$ gilt. Ein Element $e \in M$ heißt ein *neutrales Element* oder eine (*beidseitige*) *Eins* von $(M, \circ)$, falls es sowohl ein links- als auch ein rechtsneutrales Element von $(M, \circ)$ ist. Sei nun $e \in M$ eine Links-, Rechts- oder beidseitige Eins von $(M, \circ)$; für die dann vorliegende Konstellation schreibt man dann auch genauer $(M, \circ, e)$. Gilt für zwei Elemente a und b aus M die Gleichung $a \circ b = e$, so heißt a eine *Linksinverse* von b und b eine *Rechtsinverse* von a . In diesem Fall sagt man auch, dass a *rechts-invertierbar* ist, und dass b *links-invertierbar* ist. Hat ein $a \in M$ sowohl eine Links- als auch eine Rechtsinverse, so heißt a *invertierbar*.

Sei $(M, \circ)$ ein Magma mit einem neutralen Element e . Das neutrale Element heißt auch die *Verknüpfung einer leeren Familie von Elementen aus M* . Man bemerke, dass

damit automatisch Summen und Produkte über eine leere Indexmenge erklärt sind, siehe BOURBAKI [34](1974), I.2.1, Seite 13.

Eine Halbgruppe $(H, \circ)$ mit einem neutralen Element e heißt ein *Monoid* und wird mit $(H, \circ, e)$ bezeichnet.

Sei $(H, \circ, e)$ ein Monoid mit einem invertierbaren $h \in H$. Sei $\ell \in H$ eine Linksinverse und $r \in H$ eine Rechtsinverse von h . Da dann die Gleichungskette $\ell = \ell \circ e = \ell \circ (h \circ r) = (\ell \circ h) \circ r = e \circ r = r$ gilt, definiert man: Sei $(M, \circ, e)$ ein Magma mit neutralem Element e . Stimmt für ein invertierbares $m \in M$ die Linksinverse von m mit der Rechtsinversen von m überein — was, wie eben gesehen, der Fall ist, wenn $(M, \circ, e)$ ein Monoid ist —, so heißt die (somit eindeutig bestimmte) Linksinverse von m die *Inverse* von m ; sie wird meist mit m^{-1} bezeichnet und speziell, wenn die Verknüpfung $\circ$ additiv, also zum Beispiel als $+$, geschrieben wird (was oftmals (nur) der Fall ist, wenn die Verknüpfung kommutativ ist), mit $-m$ bezeichnet.

Eine Halbgruppe $(H, \circ)$ heißt eine *Gruppe*, wenn einerseits eine Linkseins e von H existiert und andererseits zu jedem $h \in H$ eine Linksinverse von h bezüglich e existiert, soll heißen, ein $g \in H$ mit $g \circ h = e$ existiert.

Sei $(G, \circ)$ eine Gruppe und $U \subseteq G$. $(U, U \upharpoonright \circ)$ heißt eine *Untergruppe* von $(G, \circ)$, falls $(U, U \upharpoonright \circ)$ eine Gruppe ist.

Eine Untergruppe U einer Gruppe $(G, \circ)$ heißt *Normalteiler* von G , wenn für alle $g \in G$ die Gleichung $g \circ U = U \circ g$ gilt.

2.1.2 Bemerkung. Sei $(M, \circ)$ ein Magma. Enthält M sowohl eine Linkseins ℓ als auch eine Rechtseins r , so stimmen ℓ und r wegen $\ell = \ell \circ r = r$ überein. M hat also dann eine Eins e mit $e = \ell = r$. Insbesondere stimmen daher zwei Einsen immer überein, so dass M nicht mehr als eine Eins haben kann. Somit sei die Eins eines Magmas M mit Eins immer einheitlich mit e bezeichnet, falls nicht etwas anderes ausdrücklich vereinbart ist.

2.1.3 Definition. Für $n \in \mathbb{N}^\times$ wird die Gruppe aller Permutationen von $\{1, \dots, n\}$ mit $\mathbf{S}_n$ bezeichnet; sie heißt die *symmetrische Gruppe* vom Grad n oder auch vom Index n . $\mathbf{S}_\infty := \bigcup_{n=1}^\infty \mathbf{S}_n$ bezeichnet die Gruppe aller Permutationen von $\mathbb{N}^\times$, die jeweils nur endlich viele Zahlen permutieren.

2.1.4 Annihilatorsysteme. Seien X und Y zwei Mengen und sei (Z, p) eine punktierte Menge. Sei B eine Abbildung $X \times Y \rightarrow Z$. Die so vorliegende Konstellation von Mengen, punktierter Menge und Abbildung heißt ein *Annihilatorsystem*, in Zeichen $(X, Y; B; (Z, p))$ oder kurz $(X, Y; B; Z)$, wenn klar ist, was der Basispunkt ist, und wird bei klarem Kontext auch einfach nur *System* genannt. In diesem Zusammenhang heißt B die *zu dem System $(X, Y; B; Z)$ gehörende Abbildung* (wenn man möchte, auch *Systemabbildung* oder genauer *Annihilatorabbildung*).

Wenn klar ist, was Z ist, wird anstelle von $(X, Y; B; Z)$ einfach nur $(X, Y; B)$ geschrieben. Dies tritt vor allem in der Dualitätstheorie per $Z = \mathbb{K}$, $p = 0$ (siehe in Abschnitt 2.4) und in der Darstellungstheorie per $Z = Y$ (siehe Abschnitt 2.8) auf. Wenn klar ist, welche Abbildung B gemeint ist, wird das Symbol für diese Abbildung weggelassen, also nur $(X, Y; Z)$, oder einfach nur (X, Y) geschrieben. Es ist manchmal praktisch, das verwendete System dem Abbildungssymbol als Index anzufügen, wobei man im Index dann natürlich auf das Abbildungssymbol selbst verzichten kann. Ist $x \in X$ und $y \in Y$, so verwendet man dann also anstelle von $B(x, y)$ die Schreibweise $B_{(X, Y; Z)}(x, y)$ bzw. $B_{(X, Y)}(x, y)$ oder, falls $X = Y$, $B_{(X; Z)}(x, y)$ bzw. $B_X(x, y)$. Oftmals wird anstelle von B das Abbildungssymbol $\langle \cdot, \cdot \rangle$ verwendet; anstelle von zum Beispiel $B_{(X, Y; Z)}(x, y)$ wird dann also $\langle x, y \rangle_{(X, Y; Z)}$ geschrieben.

Die Abbildung B heißt *rechts ausgeartet*, wenn Y eine punktierte Menge ist und ein $y \in Y^\times$ existiert, so dass für alle $x \in X$ die Gleichung $B(x, y) = p$ gilt. Analog heißt die Abbildung B *links ausgeartet*, wenn X eine punktierte Menge ist und ein $x \in X^\times$

existiert, so dass für alle $y \in Y$ die Gleichung $B(x, y) = p$ gilt. Die Abbildung B heißt *ausgeartet*, wenn sie links ausgeartet oder rechts ausgeartet ist.

Sprachlich synonym wird für den Begriff *ausgeartet* der Ausdruck *nicht trennend* verwendet. Dementsprechend ist der Ausdruck *nicht ausgeartet* synonym zu dem Begriff *trennend*. Durch Verwenden der beiden Begriffe *ausgeartet* und *trennend* ist es somit möglich, das Wort *nicht* gezielt vermeiden oder verwenden zu können, wodurch manches klarer ausgedrückt werden kann.

Ebenfalls sprachlich synonym werden die Begriffe *links* und *in X* verwendet. Analog sind die Begriffe *rechts* und *in Y* synonym. Man kann also zum Beispiel anstelle von *rechts nicht ausgeartet* auch *in Y trennend* sagen.

Die oben eingeführte Terminologie der Entartung von Abbildungen wird auf das entsprechende System übertragen. Es ist also zum Beispiel äquivalent zu sagen, *die Abbildung $B: X \times Y \rightarrow \mathbb{K}$ sei trennend*, oder zu sagen, *das System $(X, Y; B)$ sei trennend*.

Sei $(X, Y; B; (Z, p))$ ein Annihilatorsystem. Das System $(X, Y; B; Z)$ heißt *orthosymmetrisch*, falls für alle $x \in X \cap Y$ und alle $y \in X \cap Y$ aus $B(x, y) = p$ stets $B(y, x) = p$ folgt.

Betrachte die Relation

$$R := \{(x, y) \in X \times Y : B(x, y) = p\}$$

von X nach Y . Das heißt, für alle $x \in X, y \in Y$ ist genau dann $(x, y) \in R$, wenn $B(x, y) = p$ ist. Es sei an die in 1.1.3 definierte Bildung der rechten und linken Träger bezüglich einer Relation erinnert. Sei $M \subseteq X$ und $N \subseteq Y$. Dann heißt der rechte Träger von M bezüglich der Relation R der *Rechtsannihilator* von M oder im orthosymmetrischen Fall einfach auch nur der *Annihilator* von M ; es ist also

$$M^\perp = \{y \in Y : B(m, y) = p \quad \text{für alle } m \in M\}.$$

Entsprechend heißt der linke Träger von N bezüglich der Relation R der *Linksannihilator* von N ; es ist also

$$N_\perp = \{x \in X : B(x, n) = p \quad \text{für alle } n \in N\}.$$

Gemäß 1.1.18(i) bildet das Paar $(M \mapsto M^\perp, N \mapsto N_\perp)$ eine antitone Galois-Verbindung zwischen $(\mathcal{P}(X), \subseteq)$ und $(\mathcal{P}(Y), \subseteq)$. Wenn nicht etwas anderes vereinbart ist, ist dieses aus der Rechts- und Linksannihilatorbildung bestehende Paar auch stets als eine solche antitone Galois-Verbindung bildend aufzufassen. Ist das System $(X, Y; B; Z)$ orthosymmetrisch, so schreibt man für $N_\perp$ auch $N^\perp$ und nennt diese Menge ebenfalls einfach nur den *Annihilator* von N . Aber auch sonst, wird anstelle von $N_\perp$ einfach $N^\perp$ geschrieben, wenn keine Missverständnisse möglich sind. Der Vorteil der Unterscheidung von der Hoch- und der Tiefstellung des Annihilatorsymbols zeigt sich aber zum Beispiel bei sich überlappenden Annihilatorsystemen: Sind $(X_1, X_2), (X_2, X_3), (X_3, X_4), \dots$ Annihilatorsysteme, so braucht man nicht für jedes dieser Systeme ein eigenes Annihilatorsymbol zu definieren, sondern kann für alle das eine Symbol $\perp$ verwenden, vorausgesetzt, dass man die Hoch- und Tiefstellung des Annihilatorsymbols konsequent verwendet. Falls M einelementig ist, etwa $M = \{m\}$, so wird bei klarem Kontext auch $m^\perp$ anstelle von $M^\perp$ geschrieben; $M_\perp$ entsprechend; liegt des Weiteren eine Situation vor, wo klar ist, dass $m^\perp$ einelementig ist, also etwa $m^\perp = \{n\}$, so wird, falls keine Missverständnisse zu befürchten sind, anstelle des Symbols n der Ausdruck $m^\perp$ verwendet; Linksannihilator entsprechend. Das Wort Annihilator heißt im Englischen durchweg *annihilator*, aber im Deutschen findet sich dafür zeitweilen, besonders in der Theorie der Ringe oder Moduln, auch das Wort *Annulator*.

Ein $x \in X$ und ein $y \in Y$ heißen *orthogonal (zueinander)*, in Zeichen $x \perp y$, falls $B(x, y) = p$ gilt und im Fall von $x, y \in X \cap Y$ zusätzlich auch $B(y, x) = p$ erfüllt ist.

Sind X und Y zwei punktierte Mengen mit demselben Basispunkt q , so heißt eine Menge $M \subseteq X \cap Y$ *orthogonal*, falls $q \notin M$ und $x \perp y$ für jedes $x, y \in M$ mit $x \neq y$.

Im orthosymmetrischen Fall wird der Rechtsannihilator von M auch die *Orthogonalmenge* von M genannt; Linksannihilator von N entsprechend.

M heißt *annihilatorabgeschlossen*, falls M bezüglich der Relation R trägerabgeschlossen ist, sprich, $M = M^\perp_\perp$ gilt. Die bezüglich der Relation R trägerabgeschlossene Hülle von M , also $M^\perp_\perp$, heißt die *annihilatorabgeschlossene Hülle* von M . N entsprechend. Ist das System $(X, Y; B; Z)$ orthosymmetrisch, so verwendet man anstelle von annihilatorabgeschlossen den Ausdruck *orthogonalabgeschlossen* und entsprechend bezeichnet man dann eine annihilatorabgeschlossene Hülle als eine *orthogonalabgeschlossene Hülle*.

Sei $G \subseteq X$ und weiterhin wie gehabt $M \subseteq X$ und $N \subseteq Y$. Für jedes $\nu \in I$, I eine beliebige Indexmenge, sei M_ν eine Teilmenge von X . Dann gelten alle in 1.1.19 gemachten Aussagen (a) bis (m) zeichengetreu auch hier. Speziell kann man die dortige Aussage (d) zum Beispiel auch wie folgt formulieren:

$(M \times N \subseteq R) \Leftrightarrow (B(M, N) = \{p\}) \Leftrightarrow (M \subseteq N^\perp_\perp) \Leftrightarrow (N \subseteq M^\perp_\perp)$. Da für $M^\perp$ nach Definition die Inklusion $M \times M^\perp \subseteq R$, also die Gleichung $B(M, M^\perp) = \{p\}$ gilt, gilt insbesondere $M \subseteq M^\perp_\perp$. Und wegen $B(N^\perp_\perp, N) = \{p\}$ gilt analog $N \subseteq N^\perp_\perp$.

Die Aussage 1.1.19(e) — also $M^\perp = M^\perp_\perp^\perp$ und $N^\perp = N^\perp_\perp^\perp$ — kann hier auch wie folgt bewiesen werden:

Beweis. Der besseren Lesbarkeit halber sei hier nur der Beweis im orthosymmetrischen Fall gezeigt. Der allgemeine Fall geht genauso, nur dass die $\perp$, die im orthosymmetrischen Fall oben geschrieben werden dürfen, in unterer Indexposition zu setzen sind. Nach 1.1.19(d) ist $M \subseteq M^{\perp\perp}$, also gilt nach 1.1.19(c) die Inklusion $M^{\perp\perp\perp} \subseteq M^\perp$. Ebenfalls nach 1.1.19(d) gilt $M^\perp \subseteq (M^\perp)^\perp = M^{\perp\perp\perp}$. Für N entsprechend. $\square$

2.1.5 Homomorphismen. Eine Abbildung zwischen zwei Mengen mit gleichartiger Struktur, die mit diesen Strukturen verträglich ist, heißt *Homomorphismus*. Ein Homomorphismus, der injektiv ist, heißt *Monomorphismus*. Ein Homomorphismus, der surjektiv ist, heißt *Epimorphismus*. Ein Homomorphismus von einer Menge in sich selbst heißt *Endomorphismus*. Eine bijektive Abbildung, die sowohl selbst als auch ihre Umkehrabbildung ein Homomorphismus ist, heißt *Isomorphismus*. Ein Endomorphismus, der bijektiv ist, heißt *Automorphismus*. Die Menge aller Automorphismen einer (strukturierten) Menge wird mit $\text{Aut}(M)$ bezeichnet.

Seien zum Beispiel $(G, \circ)$ und $(H, \odot)$ zwei Halbgruppen oder auch nur zwei Monoide. Dann heißt eine Abbildung $T: G \rightarrow H$ ein Homomorphismus, falls für alle $x, y \in G$ die Gleichung $T(x \circ y) = T(x) \odot T(y)$ gilt. Die expliziten Definitionen der Homomorphismen gängiger Strukturen werden als bekannt vorausgesetzt, wobei hier in der Regel nicht gefordert wird, dass neutrale Elemente aufeinander abgebildet werden.

2.1.6 Kern. Seien G und H zwei Monoide. Bezeichne e_G bzw. e_H das neutrale Element von G bzw. H . Sei $T: G \rightarrow H$ ein Homomorphismus. Dann heißt die Menge $\{g \in G : T(g) = T(e_G)\}$ der *Kern* von T und wird mit $\ker T$ bezeichnet. Sind G und H zwei Gruppen mit e_G und e_H als ihre jeweiligen neutralen Elemente und ist T ein Gruppenhomomorphismus, so ist der Kern von T definiert als der Kern von der als Monoidhomomorphismus aufgefassten Abbildung T . Man beachte, dass hier T als ein Gruppenhomomorphismus stets automatisch das neutrale Element von G auf das neutrale Element von H abbildet. Der Kern eines Gruppenhomomorphismus ist ein Normalteiler.

2.1.7 Inneres direktes Produkt (VAN DER WAERDEN [285](1971), Seite 157). Seien $U_1, \dots, U_n$ Untergruppen einer Gruppe $(G, \circ)$. Bezeichne e das neutrale Element von

G . Dann heißt G ein *inneres direktes Produkt* der Untergruppen $U_1, \dots, U_n$, in Zeichen $G = U_1 \circ \dots \circ U_n$, wenn die folgenden drei Bedingungen erfüllt sind:

- (i) Alle U_i , $i \in \{1, \dots, n\}$, sind Normalteiler von G ;
- (ii) $G = U_1 \circ \dots \circ U_n$;
- (iii) $G = (U_1 \circ \dots \circ U_{k-1}) \cap U_k = \{e\}$ für alle $k \in \{2, 3, \dots, n\}$.

Ist $\circ$ kommutativ, speziell also, wenn $\circ$ additiv etwa als $+$ geschrieben wird, so spricht man entsprechend von einer *inneren direkten Summe*; speziell schreibt man dann auch $\sum_{i \in \{1, \dots, n\}}^{\oplus} U_i$ oder auch $\sum_{i=1}^n \oplus U_i$ für $U_1 \circ \dots \circ U_n$.

Eine Gruppe $(G, \circ)$ ist genau dann ein inneres direktes Produkt von Untergruppen $U_1, \dots, U_n$ von G , wenn einerseits jedes Element $g \in G$ eindeutig als Produkt $g = u_1 \circ \dots \circ u_n$, $u_k \in U_k$, $k = 1, \dots, n$, darstellbar ist und andererseits jedes Element von U_k mit jedem von U_ℓ , $k \neq \ell$, $k, \ell \in \{1, \dots, n\}$, kommutiert.

2.1.8 Definition. Ein *nichtassoziativer Ring* ist eine Menge R mit zwei Verknüpfungen $+$ und $\cdot$, so dass gilt:

- (a) $(R, +)$ ist eine abelsche Gruppe,
- (b) $a(b + c) = ab + ac$ und $(b + c)a = ba + ca$ für alle $a, b, c \in R$;

er wird wieder mit R bezeichnet, oder gemäß 2.1.1 auch genauer mit $(R, +, \cdot)$. Dabei ist wie üblich das Zeichen für die Multiplikation stärker bindend zu verstehen als das für die Addition (wobei hier in (b) gemäß 2.1.1 das Multiplikationszeichen gar nicht mehr erst explizit hingeschrieben wurde). Das neutrale Element von $(R, +)$ wird mit 0 bezeichnet. Soweit nicht etwas anderes vereinbart ist, wird jeder nichtassoziativer Ring $(R, +, \cdot)$ stets auch als die punktierte Menge R mit Basispunkt 0 betrachtet. Ein nichtassoziativer Ring $(R, +, \cdot)$ heißt *kommutativ*, wenn die Verknüpfung $\cdot$ kommutativ ist. Ist $(R, +, \cdot)$ ein nichtassoziativer Ring, so heißt eine Teilmenge U von R ein *Unterring* von R , falls U , versehen mit den beiden auf U eingeschränkten Verknüpfungen von R , also $(U, + \upharpoonright U, \cdot \upharpoonright U)$, ein nichtassoziativer Ring ist.

Ein *Ring* $(R, +, \cdot)$ ist ein nichtassoziativer Ring, so dass $(R, \cdot)$ eine Halbgruppe ist, mit anderen Worten, dass die Verknüpfung $\cdot$ assoziativ ist.

2.1.9 Definition. (a) Sei $(R, +, \cdot)$ ein nichtassoziativer Ring. Ein $\ell \in R$ heißt *Linkseins*, wenn ℓ eine Linkseins des R zugrunde liegenden Magmas $(R, \cdot)$ ist. *Rechtseins* und *Eins* entsprechend. Insoweit in R vorhanden, wird die Eins mit dem Symbol e bezeichnet. Ein $r \in R$ heißt *idempotent* oder auch ein *Idempotent* von R , wenn $r^2 = r$ gilt. Die Menge aller Idempotenten von R wird mit $\text{Idem}(R)$ bezeichnet.

Mit der Ringverknüpfung als Systemabbildung ist $(R, R; \cdot; R)$ ein Annihilatorsystem; es heißt das zu R *kanonisch assoziierte Annihilatorsystem*. Zwei Elemente r und s von R heißen *orthogonal (zueinander)*, wenn sie es bezüglich des zu R kanonisch assoziierten Annihilatorsystems sind. (Besonderes Interesse erfahren meist die idempotenten Elemente von R , die zueinander orthogonal sind.)

(b) Sei R ein nichtassoziativer Ring und S eine Teilmenge von R . Die Menge $\{r \in R : rs = sr \text{ für alle } s \in S\}$ heißt die *Kommutante* von S und wird mit S' bezeichnet. S heißt *kommutativ*, wenn jedes seiner Elemente mit jedem Element von S kommutiert. S heißt *maximal kommutativ*, wenn es in der durch die Mengeninklusion prägeordneten Menge aller kommutativen Teilmengen von R ein maximales Element ist. (Ein kommutatives S ist also genau dann maximal kommutativ, wenn es keine echte Teilmenge einer kommutativen Teilmenge von R ist. Siehe auch Satz 2.3.26.) Die Menge $S \cap S'$ heißt das *Zentrum* (im Englischen *center*) von S . Ein Element $r \in S \cap S'$ heißt *zentral in S* . Ein idempotentes $r \in R'$ heißt *zentral idempotent (in R)*.

(c) Sei $(R, +, \cdot)$ ein nichtassoziativer Ring mit Eins e . Ein idempotentes $r \in R$ heißt *echt* (im Englischen *proper*), wenn $r \neq e$. Ein Element $r \in R$ heißt *links-invertierbar*, wenn r bezüglich des zugrunde liegenden Magmas $(R, \cdot, e)$ links-invertierbar ist. *Linksinverse*, *rechts-invertierbar*, *Rechtsinverse*, *Inverse* und *invertierbar* entsprechend. Die Menge aller invertierbaren Elemente in R wird mit $G(R)$ bezeichnet. Siehe auch 2.1.20. Auch wenn ein nichtassoziativer Ring keine Eins hat, so kann man aber immer die hier als (multiplikative) *mnemotechnische Eins* bezeichnete 1 verwenden; mit ihr schreibt sich zum Beispiel der Ausdruck $y - xy$ als $(1 - x)y$. Insofern ein nichtassoziativer Ring eine Eins aufweist, gilt für die mnemotechnische Eins natürlich stets $1 = e$.

(d) (PALMER [231](1994), Seite 192). Sei $(R, +, \cdot)$ ein nichtassoziativer Ring. Dann ist das sogenannte *Quasi-Produkt* $\circ$ auf R definiert als:

$$r \circ s := r + s - rs \quad \text{für alle } r, s \in R.$$

Dann ist $(R, \circ)$ ein Magma mit 0 als neutralem Element. Ein Element r des nichtassoziativen Ringes $(R, +, \cdot)$ heißt *links-quasi-invertierbar*, wenn r bezüglich des Magmas $(R, \circ, 0)$ links-invertierbar ist. *Links-quasi-Inverse*, *rechts-quasi-invertierbar*, *Rechts-quasi-Inverse*, *Quasi-Inverse* und *quasi-invertierbar* entsprechend, wobei man bemerke, dass $(R, +, \cdot)$ genau dann assoziativ ist, wenn $(R, \circ, 0)$ assoziativ ist. Insofern existent, wird die Quasi-Inverse von r mit r^q bezeichnet. Die Menge aller quasi-invertierbaren Elemente in R wird mit $G^q(R)$ bezeichnet.

In der Literatur ist noch ein anderes Quasi-Produkt verbreitet: $r \bullet s := r + s + rs$ für alle $r, s \in R$; die Unterschiede der beiden Quasi-Produkte $\circ$ und $\bullet$ sind rein technischer Natur, siehe dazu die Bemerkung bei RICKART [246](1960), Seite 19. Wenn nicht ausdrücklich etwas anderes vereinbart ist, ist mit dem Quasi-Produkt hier immer das mit $\circ$ bezeichnete gemeint.

(e) (BOURBAKI [34](1974), I.8.1). Sei R ein nichtassoziativer Ring. Auf R definiere man die folgenden (im Allgemeinen nicht transitiven) Relationen:

$$\begin{aligned} a|_{\ell}b &: \Leftrightarrow \exists c \in R : & ac = b, \\ a|_rb &: \Leftrightarrow \exists c \in R : & ca = b, \\ a|b &: \Leftrightarrow \exists c \in R : & (a|_{\ell}b \text{ oder } a|_rb), \quad a, b \in R. \end{aligned}$$

$a|_{\ell}b$ liest sich als *a ist Linksteiler von b*, $a|_rb$ entsprechend als *a ist Rechtsteiler von b* und $a|b$ als *a ist Teiler von b*. Speziell für $b = 0$ gilt die folgende sprachliche Übereinkunft: Es ist zwar jedes $a \in R$ ein Teiler von 0, aber als *Linksnulleiler* werden nur genau die $a \in R$ bezeichnet, für die ein $c \in R^\times$ mit $ac = 0$ existiert; *Rechtsnulleiler* entsprechend. Und ein $a \in R$ heißt ein *Nulleiler*, wenn er ein Links- oder ein Rechtsnulleiler ist.

Ein nichtassoziativer Ring R , in dem es außer der Null keinen Nulleiler gibt, das heißt, wenn für alle $a, b \in R$ aus $ab = 0$ stets $a = 0$ oder $b = 0$ folgt, heißt *nulleilerfrei* (im Englischen *without non-zero zero-divisors*). Mit anderen Worten ist ein nichtassoziativer Ring genau dann nulleilerfrei, wenn das Produkt zweier Elemente aus $R^\times$ stets ungleich 0 ist. Auch im Hinblick auf die Definition 2.1.22 sei hier erwähnt, dass zum Beispiel ROWEN [250](1988), Seite 2, einen nulleilerfreien Ring $R \neq \{0\}$ im Englischen als *domain* bezeichnet, und dass JACOBSON [152](1974), Seite 87, einen nulleilerfreien Ring $R \neq \{0\}$ mit Eins im Englischen als *domain* oder auch als *integral domain* bezeichnet.

2.1.10 Kronecker-Delta. Sei R ein nichtassoziativer Ring mit Eins. Sei I eine Indexmenge. Das Symbol δ_{ij} , $i, j \in I$, bezeichnet die 0 von R , falls $i \neq j$, und die Eins von R , falls $i = j$.

2.1.11 Bemerkung (GROVE [122](1983), Seite 48). Sei R ein Ring mit Eins. Weist ein Unterring S von R eine Eins auf, so kann diese von der Eins von R verschieden sein. Als

Beispiel betrachte für R den Matrizenring aller 2×2 -Matrizen mit Einträgen aus $\mathbb{K}$ und für S dessen Matrizenunterring $\left\{ \begin{pmatrix} x & 0 \\ 0 & 0 \end{pmatrix} : x \in \mathbb{K} \right\}$.

2.1.12 Definition. Sei $(R, +, \cdot)$ ein nichtassoziativer Ring und $\mathfrak{a}$ eine Untergruppe von $(R, +)$. $\mathfrak{a}$ heißt ein *Linksideal* von R , wenn $R \cdot \mathfrak{a} \subseteq \mathfrak{a}$ gilt. $\mathfrak{a}$ heißt ein *Rechtsideal* von R , wenn $\mathfrak{a} \cdot R \subseteq \mathfrak{a}$ gilt. $\mathfrak{a}$ heißt ein *Ideal*, wenn es sowohl ein Links- als auch ein Rechtsideal ist. Gelegentlich wird der Deutlichkeit halber ein Ideal als ein *beidseitiges Ideal* bezeichnet und dementsprechend die Links- und Rechtsideale allgemein als *einseitige Ideale*.

Ist A eine nicht leere Teilmenge von R , dann bezeichnet $(A)_\ell$ oder auch $(A|$ das kleinste Linksideal in R mit $A \subseteq (A)_\ell$; es heißt das von A *erzeugte Linksideal*. Analog bezeichnet $(A)_r$ oder $|A)$ das kleinste Rechtsideal in R mit $A \subseteq (A)_r$; es heißt das von A *erzeugte Rechtsideal*. Und (A) bezeichnet das kleinste beidseitige Ideal in R mit $A \subseteq (A)$; es heißt das von A *erzeugte Ideal*. Ist A endlich, so heißt das jeweils erzeugte Links-, Rechts- oder beidseitige Ideal *endlich erzeugt*.

Ist A einelementig, etwa $A = \{a\}$, so schreibt man (a) anstelle von $(\{a\})$; es heißt das von a erzeugte *Hauptideal* (im Englischen *principal ideal*).

2.1.13 Produkte von Teilmengen (DALES [56](2000), Seite 29). Sei $(R, +, \cdot)$ ein nichtassoziativer Ring. Seien A und B zwei nicht leere Teilmengen von R . Dann ist mit dem sogenannten mit AB bezeichneten *Produkt* von A und B die Menge aller Elemente der Form $a_1b_1 + \dots + a_nb_n$ mit $a_k \in A$, $b_k \in B$, $k \in \{1, \dots, n\}$, $n \in \mathbb{N}^\times$ gemeint. Mit A^2 ist AA gemeint. Und falls nicht ausdrücklich etwas anderes vereinbart ist, meint A^3 den Ausdruck $A(A^2)$ und allgemeiner für $n \in \mathbb{N}$ mit $n > 2$ meint A^n den Ausdruck $A(A^{n-1})$ — durch den jeweiligen Kontext wird dabei eine Verwechslung mit dem kartesischen Produkt A^n vermieden werden.

Es sei an das in 2.1.1 definierte Komplexprodukt erinnert. Wieder wird für einelementige Mengen nur das Element selber geschrieben. Man beachte, dass für $a \in R$ zwar immer $a \cdot R = aR$ gilt, aber bereits für assoziatives R im Allgemeinen $R \cdot a \cdot R \neq RaR$ gilt, da sich nicht notwendig jedes Element von der Form $r_1as_1 + \dots + r_nas_n$ mit $r_k, s_k \in R$, $k \in \{1, \dots, n\}$, $n \in \mathbb{N}^\times$ als ein ras mit $r, s \in R$ schreiben lassen muss. Ist speziell R ein Ring und sind $\mathfrak{a}$ und $\mathfrak{b}$ zwei Linksideale in R , so gilt $\mathfrak{a}\mathfrak{b} = (\mathfrak{a} \cdot \mathfrak{b})_\ell \subseteq \mathfrak{b}$; siehe MCCOY [209](1964), Seite 31.

2.1.14 Summen von Idealen (MCCOY [209](1964), Seiten 27 und 33). Sei R ein nichtassoziativer Ring. Seien $\mathfrak{a}$ und $\mathfrak{b}$ zwei Linksideale von R . Dann gilt $\mathfrak{a} + \mathfrak{b} = (\mathfrak{a} \cup \mathfrak{b})_\ell$; man schreibt daher anstelle von $\mathfrak{a} + \mathfrak{b}$ auch $(\mathfrak{a}, \mathfrak{b})_\ell$. Ist $\mathfrak{a}_\nu$, $\nu \in I$, eine Familie von Linksidealen von R , so schreibt man für $(\bigcup_{\nu \in I} \mathfrak{a}_\nu)_\ell$ den Ausdruck $\sum_{\nu \in I} \mathfrak{a}_\nu$; er heißt die von den $\mathfrak{a}_\nu$, $\nu \in I$, *erzeugte Summe* oder auch einfach nur *die Summe der* $\mathfrak{a}_\nu$, $\nu \in I$. Für Rechtsideale und Ideale entsprechend. Ist speziell R ein Ring, so gilt: Für $r_1, \dots, r_n \in R$ gilt: $(r_1, \dots, r_n)_\ell = (r_1)_\ell + \dots + (r_n)_\ell$. Für Rechtsideale und Ideale gilt die analoge Gleichung. Sind $\mathfrak{a}$, $\mathfrak{b}$ und $\mathfrak{c}$ Linksideale eines Ringes, so gilt $\mathfrak{a}(\mathfrak{b} + \mathfrak{c}) = \mathfrak{a}\mathfrak{b} + \mathfrak{a}\mathfrak{c}$.

2.1.15 Definition. Sei R ein nichtassoziativer Ring mit Eins e_R und sei S ein nichtassoziativer Ring mit Eins e_S . Eine Abbildung $T: R \rightarrow S$ heißt *unital*, wenn $T(e_R) = e_S$ gilt.

2.1.16 Definition. Seien R und S zwei nichtassoziative Ringe. Ein *Homomorphismus* $T: R \rightarrow S$ ist ein Homomorphismus zwischen den additiven Gruppen, der multiplikativ ist: $T(rs) = T(r)T(s)$ für alle $r, s \in R$.

Ein *Antihomomorphismus* $T: R \rightarrow S$ ist ein Homomorphismus zwischen den additiven Gruppen, der *antimultiplikativ* ist:

$$T(rs) = T(s)T(r) \quad \text{für alle } r, s \in R.$$

Analog wie in 2.1.5 bildet man Termini wie zum Beispiel *Antiautomorphismus*.

Ist $T: R \rightarrow S$ ein Homomorphismus, so heißt das mit $\ker T$ bezeichnete beidseitige Ideal $\{r \in R : T(r) = 0\}$ der *Kern* von T .

Falls R und S jeweils eine Eins besitzen und $T: R \rightarrow S$ ein Homomorphismus oder ein Antihomomorphismus ist, so wird nicht gefordert, dass er unital zu sein hat.

2.1.17 Bemerkung. Jeder surjektive Endomorphismus und jeder surjektive Antiendomorphismus eines nichtassoziativen Ringes mit Eins ist unital, das heißt, hat die Eins als Fixpunkt.

2.1.18 Definition. Sei A ein nichtassoziativer Ring und $a \in A$. Die *durch a bestimmte Links-Multiplikation* in A ist die Abbildung $L_a: x \mapsto ax$ für alle $x \in A$. Analog ist die *durch a bestimmte Rechts-Multiplikation* in A die Abbildung $R_a: x \mapsto xa$ für alle $x \in A$.

2.1.19 Innere direkte Summe (DIVINSKY [72](1965), Seite 25). Sei R ein nichtassoziativer Ring. Ist $R = R_1 + \dots + R_n$ für Unterringe $R_1, \dots, R_n$ von R mit $R_i \cap (\sum_{j \neq i} R_j) = \{0\}$ für alle $i \in \{1, \dots, n\}$, so heißt R eine *direkte Summe* (genauer *innere direkte Summe*) der Unterringe $R_1, \dots, R_n$, in Zeichen $R = R_1 \dot{+} \dots \dot{+} R_n$ oder $R = \sum_{i=1}^n \oplus R_i$.

Dabei sind alle R_i , $i \in \{1, \dots, n\}$, genau dann Ideale von R , wenn $R_i \cdot R_j = \{0\}$ für alle $i \neq j$ gilt. Gilt nämlich letzteres, so sind wegen $R \cdot R_i = (R_1 + \dots + R_n) \cdot R_i \subseteq R_1 \cdot R_i + \dots + R_n \cdot R_i = R_i \cdot R_i \subseteq R_i$ und analog $R_i \cdot R \subseteq R_i$ die R_i , $i = 1, \dots, n$, Ideale von R . Sind andererseits die $\mathfrak{a}_i := R_i$, $i = 1, \dots, n$, Ideale von R , so ist wegen $\mathfrak{a}_i \cdot \mathfrak{a}_j \subseteq \mathfrak{a}_i \cap \mathfrak{a}_j$ auch $\mathfrak{a}_i \cdot \mathfrak{a}_j = 0$ für alle $i \neq j$. Da für zwei Ideale $\mathfrak{a}$ und $\mathfrak{b}$ eines Ringes genau dann $\mathfrak{a}\mathfrak{b} = \{0\}$ ist, wenn $\mathfrak{a} \cdot \mathfrak{b} = \{0\}$ ist, sind also die Unterringe $R_1, \dots, R_n$ der direkten Summe $R = R_1 \dot{+} \dots \dot{+} R_n$ genau dann allesamt Ideale, wenn sie multiplikativ separiert sind; in diesem Fall sagt man, R sei die (innere) direkte Summe der Ideale $\mathfrak{a}_1, \dots, \mathfrak{a}_n$.

$R = R_1 + \dots + R_n$ ist genau dann eine innere direkte Summe von den Unterringen $R_1, \dots, R_n$, wenn die 0 sich eindeutig als eine Summe von Elementen der $R_1, \dots, R_n$ darstellen lässt.

Die eingangs erwähnte Durchschnittsbedingung $R_i \cap (\sum_{j \neq i} R_j) = \{0\}$ für alle $i \in \{1, \dots, n\}$ ist offensichtlich zwar hinreichend, aber im Allgemeinen nicht notwendig dafür, dass $R_i \cap R_j = \{0\}$ für alle $i \neq j$ gilt.

Ist $\mathfrak{a}_\nu$, $\nu \in I$, eine Familie von Linksidealen von R mit $R = \sum_{\nu \in I} \mathfrak{a}_\nu$ und gilt $\mathfrak{a}_\nu \cap (\sum_{\mu \in I \setminus \{\nu\}} \mathfrak{a}_\mu) = \{0\}$ für alle $\nu \in I$, so sagt man entsprechend, dass R die *innere direkte Summe* der $\mathfrak{a}_\nu$, $\nu \in I$, sei, in Zeichen $R = \sum_{\nu \in I}^\oplus \mathfrak{a}_\nu$. Rechtsideale und Ideale entsprechend.

2.1.20 Bemerkungen. (a) Sei R ein Ring mit Eins. Zusammen mit der Ringmultiplikation ist $G(R)$ eine Gruppe. Da ein invertierbares Element eines Ringes mit Eins auch *Einheit* genannt wird, heißt $G(R)$ auch die *Einheitengruppe* von R .

(b) Sei R ein Ring. Zusammen mit dem Quasi-Produkt ist $G^q(R)$ eine Gruppe mit 0 als neutralem Element.

(c) Sei R ein Ring mit Eins. Dann ist die Abbildung $G(R) \rightarrow G^q(R)$, $r \mapsto e - r$ ein Gruppenisomorphismus. (Siehe auch Satz 2.1.41.)

(d) (RICKART [246](1960), Seite 2). Sei R ein Ring und bezeichne $\circ$ und $\bullet$ die beiden in Definition 2.1.9 erklärten Quasi-Produkte auf R . Dann gelten für alle $x, y, z \in R$ die

folgenden Gleichungen:

$$(1 - x)(1 - y) = 1 - (x \circ y), \quad (2.1)$$

$$(1 + x)(1 + y) = 1 + (x \bullet y), \quad (2.2)$$

$$(x \circ y)(1 - x) = (1 - x)(y \circ x), \quad (2.3)$$

$$(x \bullet y)(1 + x) = (1 + x)(y \bullet x), \quad (2.4)$$

$$x \circ (y + z) = x \circ y + x \circ z - x, \quad (2.5)$$

$$x \bullet (y + z) = x \bullet y + x \bullet z - x. \quad (2.6)$$

(e) Ist $\mathfrak{a}$ ein Rechtsideal von R und $x \in \mathfrak{a}$ rechts-quasi-invertierbar, so ist wegen $x \circ y = 0$, $x + y - xy = 0$, $y = xy - x$ auch die Rechts-quasi-Inverse von x in $\mathfrak{a}$.

2.1.21 Bezeichnungen. Der Begriff *positiv* schließt, wenn nicht anderes explizit erwähnt ist, die Möglichkeit »gleich Null« immer mit ein. Soll diese Möglichkeit explizit ausgeschlossen sein, wird der Begriff *echt positiv* (im Englischen *strictly positiv*) verwendet.

2.1.22 Definition. Ein nullteilerfreier, kommutativer Ring R mit $R \neq \{0\}$ heißt *Integritätsbereich* (im Englischen *integral domain* oder auch *domain of integrity*; siehe aber auch die Bemerkung in Definition 2.1.9(e)).

Ein *Hauptidealring* (im Englischen *principal ideal ring*) ist ein kommutativer Ring $R \neq \{0\}$ mit Eins, in dem jedes Ideal ein Hauptideal ist. Ein Integritätsbereich mit Eins, in dem jedes Ideal ein Hauptideal ist, heißt *Hauptidealbereich* (im Englischen *principal ideal domain*). Man beachte, dass in der deutschen Literatur mit einem *Hauptidealring* durchaus ein Hauptidealbereich gemeint sein kann (so etwa bei ARTIN [10](1993), VAN DER WAERDEN [285](1971), HORNFECK [141](1976), MEYBERG [214](1980) und SCHULZE [258](2006)), aber im Englischen ein *principal ideal ring* eher den hier definierten Hauptidealring meinen kann.

Sei R ein Integritätsbereich mit Eins. Ein nicht invertierbares $p \in R^\times$ heißt *Primelement* (in R), wenn für alle $a, b \in R$ aus $p|ab$ stets $p|a$ oder $p|b$ folgt.

2.1.23 Definition. Sei $(R, +, \cdot)$ ein nichtassoziativer Ring mit Eins, $R \neq \{0\}$. Dann heißt R ein *nichtassoziativer Schiefkörper*, wenn die beiden Gleichungen $ax = b$ und $ya = b$ für alle $a \in R^\times$, $b \in R$ eindeutige (!) Lösungen $x, y \in A$ haben. Entsprechend heißt ein Ring $(R, +, \cdot)$ *Schiefkörper* (im Englischen *skew field* oder *sfield*), wenn $(R^\times, \cdot)$ eine Gruppe ist. Ein kommutativer Schiefkörper heißt *Körper*. Anstelle des Wortes Schiefkörper kann man synonym auch die Worte *Divisorenring* oder *Divisionsring* (im Englischen *division ring*) verwenden.

2.1.24 Zahlen. Der Ring der ganzen Zahlen wird mit $\mathbb{Z}$ bezeichnet. Der Körper der reellen beziehungsweise komplexen Zahlen wird mit $\mathbb{R}$ beziehungsweise $\mathbb{C}$ bezeichnet. $\mathbb{R}^+$ bezeichnet die Menge aller positiven reellen Zahlen. $\mathbb{K}$ steht im Allgemeinen für $\mathbb{R}$ oder $\mathbb{C}$, falls nicht etwas anderes ausdrücklich vereinbart ist.

Die *erweiterte reelle Zahlengerade* (im Englischen *extended real line*) $\overline{\mathbb{R}} := \mathbb{R} \cup \{-\infty, \infty\}$ wird hier wie bei FREMLIN [99](2004), 112Ba und Abschnitt 135, aufgefasst. Insbesondere ist also $\overline{\mathbb{R}}$ mit einer algebraischen, einer Ordnungs- und einer topologischen Struktur versehen, per dem Areatangens Hyperbolicus eine zwei-Punkt-Kompaktifizierung von $\mathbb{R}$ und es gilt $-\infty \cdot 0 = \infty \cdot 0 = 0 \cdot (-\infty) = 0 \cdot \infty = 0$. Die Ausdrücke $\infty - \infty$ und $-\infty + \infty$ sind *nicht* definiert. Man bemerke, dass in $\overline{\mathbb{R}}$ die beiden Gleichungen $\sup \emptyset = -\infty$ und $\inf \emptyset = \infty$ gelten, ohne dass dies hier explizit definiert werden müsste. Die Elemente von $\overline{\mathbb{R}}$ heißen *erweiterte reelle Zahlen*.

2.1.25 Limes inferior und Limes superior (DUNFORD und SCHWARTZ [78](1958), Seite 4; MEGGINSON [211](1998), Seite 217). Sei $(x_\nu)_{\nu \in I}$ ein Netz von reellen Zahlen.

Für jedes $\nu \in I$ setze $i_\nu := \inf\{x_\mu : \nu \leq \mu\}$ und $s_\nu := \sup\{x_\mu : \nu \leq \mu\}$, wobei das Infimum und das Supremum in der erweiterten reellen Zahlengerade genommen werden. Für $\nu \leq \mu$ gilt dann $i_\nu \leq i_\mu$ und $s_\nu \geq s_\mu$. Somit sind $(i_\nu)_{\nu \in I}$ und $(s_\nu)_{\nu \in I}$ zwei in $\overline{\mathbb{R}}$ konvergente Netze; deren Grenzwerte werden wie folgt benannt: $\lim_\nu i_\nu$ heißt der *Limes inferior* des Netzes (x_ν) , in Zeichen $\liminf_\nu(x_\nu)$, und $\lim_\nu s_\nu$ heißt der *Limes superior* des Netzes (x_ν) , in Zeichen $\limsup_\nu(x_\nu)$. Eine Liste von Rechenregeln findet man bei MEGGINSON [211](1998), Seite 221, Übung 2.55.

Sei A eine unendliche Teilmenge von $\mathbb{R}$. Dann definiert man den *Limes superior* der Menge A , in Zeichen $\limsup A$, als das in $\overline{\mathbb{R}}$ genommene Infimum aller $x \in \mathbb{R}$, für die es jeweils nur eine endliche Menge von Elementen $a \in A$ mit $x < a$ gibt: $\limsup A = \inf\{x \in \overline{\mathbb{R}} : a \leq x \text{ für fast alle } a \in A\}$; der entsprechende Ausdruck für den *Limes inferior* von A ist definiert durch $\liminf A = \sup\{x \in \overline{\mathbb{R}} : x \leq a \text{ für fast alle } a \in A\}$.

2.1.26 Metrischer Raum (PITTS [241](1972)). Sei X eine Menge. Sei $d: X \times X \rightarrow \mathbb{R}$ eine Abbildung. Weist die Abbildung d die beiden Eigenschaften

- (a) $d(x, y) = 0 \Leftrightarrow x = y$ und
- (b) $d(x, z) \leq d(y, x) + d(y, z)$

für alle $x, y, z \in X$ auf, so heißt das Paar (X, d) ein *metrischer Raum* und die Abbildung d die *Metrik* (im Französischen *distance*) von (X, d) . Die Abbildung d ist genau dann eine Metrik auf X , wenn die folgenden drei Bedingungen für alle $x, y, z \in X$ erfüllt sind: (i) $d(x, y) \geq 0$, wobei Gleichheit nur für $x = y$ vorliegt, (ii) $d(x, y) = d(y, x)$ und (iii) $d(x, z) \leq d(x, y) + d(y, z)$.

Sei (X, d) ein metrischer Raum. Für jedes $x_0 \in X$ und $r > 0$ setzt man $U(x_0; r) := \{x \in X : d(x_0, x) < r\}$. Für jedes $x_0 \in X$ und $r \geq 0$ setzt man $B(x_0; r) := \{x \in X : d(x_0, x) \leq r\}$.

Die Menge aller $U(x; r)$ bilden die Basis einer Topologie von X . Wenn nicht etwas anderes vereinbart ist, ist jeder metrische Raum stets als mit dieser Topologie ausgestattet aufzufassen.

Seien A und B zwei Teilmengen von X . Dann nennt man die erweiterte reelle Zahl $\text{dist}(A, B) := \inf\{d(x, y) : x \in A, y \in B\}$ den *Abstand* der beiden Mengen. Ist A einelementig, etwa $A = \{x\}$, so schreibt man $\text{dist}(x, B)$ anstatt $\text{dist}(\{x\}, B)$; B entsprechend.

Jeder metrische Raum erfüllt die Trennungsaxiome T_n für $n = 1, 2, 3, 3a, 4$.

2.1.27 Epigraph (VAN TIEL [277](1984)). Sei X eine Menge und $f: X \rightarrow \overline{\mathbb{R}}$ eine Abbildung. Dann heißt die Menge

$$\text{epi}(f) := \{(x, y) \in X \times \mathbb{R} : y \geq f(x)\}$$

der *Epigraph* von f .

2.1.28 Ringe mit Operatorenbereich (DIVINSKY und SULIŃSKI [73](1965); KUROSCHE [193](1964), Seite 172; HAZEWINKELE [129](1988-2001)).

Sei $(R, +, \cdot)$ ein nichtassoziativer Ring. Eine Menge K heißt ein *Operatorenbereich* für R (im Englischen *operator domain*), wenn jedes Element $k \in K$ als ein Gruppenendomorphismus auf $(R, +)$ wirkt mit

$$k(r \cdot s) = k(r) \cdot s = r \cdot k(s) \quad \text{für alle } r, s \in R.$$

Die Elemente von K heißen dann *Operatoren*. Ist K ein Operatorenbereich für R , so heißt R ein *nichtassoziativer Ring mit Operatoren*, genauer ein *nichtassoziativer Ring mit Operatorenbereich* K und nennt ihn auch einen *nichtassoziativen K -Operator-Ring* oder einfach nur einen *nichtassoziativen K -Ring*. Ist dabei R ein Ring, so gelten die entsprechenden Begriffe unter Weglassung des Wortes *nichtassoziativ*. Jeder nichtassoziative Ring ist ein nichtassoziativer $\mathbb{Z}$ -Ring.

Ein *Ring mit einem Ring von Operatoren* ist ein Ring R mit einem Operatorenbereich K , wobei K ein Ring ist und für alle $k_1, k_2 \in K, r \in R$ die folgenden beiden Gleichungen gelten:

$$(k_1 + k_2)(r) = k_1(r) + k_2(r) \quad \text{und} \quad (k_1 k_2)(r) = k_1(k_2(r)).$$

Für nichtassoziative Ringe R entsprechend.

Wenn nicht ausdrücklich etwas anderes vereinbart ist, wird im vorliegenden Text ein nichtassoziativer Ring mit Operatorenbereich K mit $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$ immer stillschweigend als ein nichtassoziativer Ring mit einem Ring von Operatoren verstanden.

Sei R ein nichtassoziativer Ring mit einem Operatorenbereich K . Eine Teilmenge M von R heißt *K -zulässig* (im Englischen *K -admissible*), wenn die Inklusion $\bigcup_{k \in K} k(M) \subseteq M$ gilt. Falls R eine Eins hat, so ist dann wegen $k(\mathfrak{a}) = k(e \cdot \mathfrak{a}) = k(e) \cdot \mathfrak{a} \subseteq R \cdot \mathfrak{a} \subseteq \mathfrak{a}$ jedes Ideal $\mathfrak{a}$ von R K -zulässig.

Man hat also neben der algebraischen Struktur der nichtassoziativen Ringe die algebraische Struktur der nichtassoziativen Ringe mit einem Operatorenbereich und in dieser Struktur definiert man wie folgt den Begriff des Ideals:

Ist R ein nichtassoziativer Ring mit einem Operatorenbereich K , so heißt ein $\mathfrak{a} \subseteq R$ genau dann bezüglich der algebraischen Struktur der nichtassoziativen Ringe mit einem Operatorenbereich ein *Ideal* (im weiteren Verlauf meist in solch einem Zusammenhang als ein *Ideal des nichtassoziativen K -Ringes* angesprochen), wenn es bezüglich der algebraischen Struktur der nichtassoziativen Ringe ein Ideal des nichtassoziativen Ringes R ist, das K -zulässig ist. Die einseitigen Ideale und Unterringe werden entsprechend definiert.

2.1.29 Definition. Der *Kommutator* auf einem nichtassoziativen Ring R mit Operatorenbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$, wird mit $[\cdot, \cdot]$ bezeichnet; also $[r, s] = rs - sr$ für alle $r, s \in R$.

Analog wie der Kommutator als Maß dafür dienen kann, wie weit zwei Elemente davon entfernt sind, zu kommutieren, definiert man als ein gewisses Maß für die Assoziativität:

2.1.30 Definition (SCHAFFER [256](1966), Seite 5). Der *Assoziator* auf einem nichtassoziativen Ring R mit Operatorenbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$, ist die Abbildung

$$(\cdot, \cdot, \cdot) : R \times R \times R \rightarrow R, (x, y, z) \mapsto (xy)z - x(yz).$$

2.1.31 Definition. Sei $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$. Ein *alternativer Ring mit Operatorenbereich K* (auch *alternativer K -Ring* genannt) ist ein nichtassoziativer Ring R mit Operatorenbereich K mit

$$(x, x, y) = (y, x, x) = 0 \quad \text{für alle } x, y \in R.$$

2.1.32 Bemerkungen. Sei $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$. Dann ist jeder K -Ring ein alternativer K -Ring. Man kann also die alternativen K -Ringe als eine geringe Verallgemeinerung der K -Ringe betrachten, in dem Sinne, dass das Assoziativgesetz abgeschwächt worden ist.

Die Definition eines alternativen Ringes mit Operatorenbereich K , lässt sich auch mit den Links- und Rechts-Multiplikationen von 2.1.18 formulieren: Ein nichtassoziativer Ring R mit Operatorenbereich K ist genau dann ein alternativer Ring mit Operatorenbereich K , wenn gilt:

$$L_{x^2} = (L_x)^2 \quad \text{und} \quad R_{x^2} = (R_x)^2 \quad \text{für alle } x \in R.$$

Mit KUROSCH [193](1964), Seite 205, hat man die Bemerkung: Ein nichtassoziativer Ring mit Operatorenbereich K , ist genau dann ein Ring mit Operatorenbereich K , wenn sämtliche aus *drei* Elementen erzeugten nichtassoziativen Unterringe mit Operatorenbereich K assoziativ sind. Gemäß KUROSCH, ebd., und SCHAFFER [256](1966), Seite 29, gilt der sogenannte Satz von Artin: Ein nichtassoziativer Ring mit Operatorenbereich K ist

genau dann ein alternativer Ring mit Operatorenbereich K , wenn sämtliche aus *zwei* Elementen erzeugten nichtassoziativen Unterringe mit Operatorenbereich K assoziativ sind.

Der Assoziator ist in gewissen Sinne alternierend: Sei R ein nichtassoziativer Ring mit Operatorenbereich K und π eine Permutation auf $\{1, 2, 3\}$. Dann gilt für alle $x_1, x_2, x_3 \in R$ die Gleichung $(x_{\pi(1)}, x_{\pi(2)}, x_{\pi(3)}) = (\operatorname{sgn} \pi) (x_1, x_2, x_3)$.

2.1.33 Peirce-Zerlegung I (ALBERT [6](1961), Seite 24; DICKSON [63](1923), Seite 98; JACOBSON [150](1964), Seite 48; MCCOY [209](1964), Seite 24; PALMER [231](1994), Seite 667).

Sei R ein Ring und $a \in R$. Dann gilt $(a)_\ell = \mathbb{Z}a + Ra$ und $(a)_r = \mathbb{Z}a + aR$. Ra ist ein Linksideal und aR ist ein Rechtsideal. Dies trifft im Allgemeinen für nichtassoziative Ringe nicht zu. aRa ist zwar im Allgemeinen kein beidseitiges Ideal, aber ein Unterring von R .

Sei p ein idempotentes Element von R . Dann gelten die folgenden Gleichungen:

$$pR = \{x \in R : px = x\}, \quad (2.7)$$

$$(1 - p)R = \{x \in R : px = 0\}, \quad (2.8)$$

$$Rp = \{x \in R : xp = x\}, \quad (2.9)$$

$$R(1 - p) = \{x \in R : xp = 0\}. \quad (2.10)$$

Beweis. Zum Beispiel Gleichung (2.9): Sei $x \in Rp$, also $x = yp$ für ein $y \in R$. Dann ist $x = yp = ypp = xp$; das zeigt „ $\subseteq$ “. „ $\supseteq$ “ ist klar. Die anderen Gleichungen gehen genauso. $\square$

Betrachtet man also zum Beispiel die in Definition 2.1.18 erklärte durch p bestimmte Links-Multiplikation L_p im Ring R , aufgefasst als ein Endomorphismus der dem Ring zugrunde liegenden additiven Gruppe, so ist das Rechtsideal $L_p(R) = pR$ gleich der Fixpunktmenge von L_p und das Rechtsideal $(1 - p)R$ gleich dem Kern von L_p .

Weiterhin gilt

$$p_1Rp_2 = p_1R \cap Rp_2 \quad \text{für alle } p_1, p_2 \in \{p, 1 - p\}, \quad (2.11)$$

das heißt,

$$pRp = \{x \in R : px = x, \quad xp = x\}, \quad (2.12)$$

$$pRp = \{x \in R : pxp = x\}, \quad (2.13)$$

$$(1 - p)Rp = \{x \in R : px = 0, \quad xp = x\}, \quad (2.14)$$

$$pR(1 - p) = \{x \in R : px = x, \quad xp = 0\}, \quad (2.15)$$

$$(1 - p)R(1 - p) = \{x \in R : px = 0, \quad xp = 0\}. \quad (2.16)$$

Der Gleichung (2.16) entnimmt man, dass genau die Elemente des Unterringes $J := (1 - p)R(1 - p)$ die zu p orthogonalen Elemente sind und diese wiederum sind — wie man leicht nachrechnet — genau die zu allen Elementen von pRp orthogonalen Elemente. Insbesondere bemerke man, dass $pJ = Jp = \{0\}$ gilt. Dieser Unterring J heißt auch der *durch p zerstörte Teil* von R (DICKSON, ebd.). pRp ist ein Unterring mit p als Eins, aber im Allgemeinen kein einseitiges Ideal von R ; er wird im Englischen manchmal auch der zum Idempotent p assoziierte *corner ring* bezeichnet, siehe LAM [194](2001), Seite 309. Die R zugrunde liegende additive Gruppe $(R, +)$ ist gleich den folgenden inneren direkten Summen:

$$R = pR \dot{+} (1 - p)R, \quad (2.17)$$

$$R = Rp \dot{+} R(1 - p), \quad (2.18)$$

$$R = pRp \dot{+} pR(1 - p) \dot{+} (1 - p)Rp \dot{+} (1 - p)R(1 - p). \quad (2.19)$$

Die Zerlegungen (2.17), (2.18) und (2.19) heißen jeweils die *rechte*, *linke* und *beidseitige Peirce-Zerlegung* von R bezüglich des idempotenten Elementes p . Sie ist nach dem amerikanischen Algebraiker BENJAMIN PEIRCE benannt. Der Name *Peirce* wird in etwa so ausgesprochen, wie man im Englischen die Buchstabenkombination *purss* aussprechen würde; der Anfang also ähnlich dem englischen Wort *purse*.

Setzt man

$$A := \{x \in R : \exists r \in R : x = pr - rp\},$$

so gilt

$$pA = pR(1 - p), \quad Ap = (1 - p)Rp, \quad pAp = \{0\}.$$

Zum Anschluss siehe auch 2.1.46 und 2.7.7.

2.1.34 Moduln. (WAHRIG [286](1982). Für das vom lateinischen Wort *modulus* für *Maß*, *Maßstab* kommende Wort *Modul* gibt es in der deutschen Sprache sowohl *das Modul* mit dem Plural *die Module* als auch *der Modul* mit dem Plural *die Moduln*; letzterer wird auf der ersten Silbe betont und ist der im Folgenden definierte in der Algebra übliche Begriff.

(BOURBAKI [34](1974), II.1.1, Seite 191; BEHRENS [23](1975), Seite 18). Sei R ein Ring. Ein *Linksmodul* über R (auch *R -Linksmodul*) ist eine abelsche Gruppe $(M, +)$, die mit einer Verknüpfung $\varphi: R \times M \rightarrow M$ versehen ist, die für alle $r, s \in R$ und $m, n \in M$ die folgenden drei Gleichungen erfüllt:

- (i) $\varphi(r, m + n) = \varphi(r, m) + \varphi(r, n)$,
- (ii) $\varphi(r + s, m) = \varphi(r, m) + \varphi(s, m)$,
- (iii) $\varphi(r, \varphi(s, m)) = \varphi(rs, m)$.

Ein *Rechtsmodul* über R (auch *R -Rechtsmodul*) wird bis auf (iii) genauso definiert; (iii) lautet dann:

- (iii)' $\varphi(r, \varphi(s, m)) = \varphi(sr, m)$.

Solange nicht etwas anderes vereinbart ist, gelte fortan die folgende Vereinbarung: Ist M ein Links- oder ein Rechtsmodul, so ist M auch stets aufzufassen als die punktierte Menge M mit dem neutralen Element von $(M, +)$ als Basispunkt. Insofern von einem R -Links- oder einem R -Rechtsmodul M die Rede ist, werden bisweilen die Elemente des Ringes R auch *Skalare* genannt und R auch der *Skalarenring* von M .

Besitzt der Ring R eine Eins und ist M ein R -Links- oder ein R -Rechtsmodul, so heißt M *unital*, falls für alle $m \in M$ die Gleichung $\varphi(e, m) = m$ erfüllt ist. Als Sprechweise wird festgesetzt: Ist von einem unitalen Modul über einen Ring R die Rede, so ist dabei automatisch R als ein Ring mit Eins vorausgesetzt.

Ist M ein Linksmodul über R , so schreibt man anstelle von $\varphi(r, m)$ üblicherweise den Ausdruck rm ; ist M ein Rechtsmodul über R , dann anstelle von $\varphi(r, m)$ den Ausdruck mr .

Mit dem nicht mit einem der Adverbien *links* oder *rechts* in Komposition stehenden Wort *Modul* (im Englischen *module*, Plural *modules*) ist ab hier immer das Wort *Linksmodul* gemeint. Nichtsdestotrotz ist es für manche Formulierungen dem Verständnis dienlicher, explizit das Wort *Linksmodul* zu verwenden, zum Beispiel im Zusammenhang mit Aussagen über Idealen. Jedenfalls sind Aussagen über *Moduln* als Aussagen über *Linksmoduln* aufzufassen und wenn nicht etwas anderes ausdrücklich erwähnt wird, gilt jeweils auch die Aussage, mutatis mutandis, über *Rechtsmoduln*.

Ein *Unterm modul* eines Moduls $(M, +, \varphi)$ über R ist eine additive Untergruppe von $(M, +)$, die bezüglich der Abbildung φ abgeschlossen ist. Jeder Ring R kann kanonisch als ein Modul über sich aufgefasst werden. Dann sind die *Unterlinksmoduln* genau gleich den *Linksideal*en. Man beachte aber (GROVE [122](1983), Seite 126): Ist S ein Unterring von R , dann ist zwar R stets ein Modul über S , aber selbst wenn R und S beide Ringe

mit Eins sind, muss dieser Modul nicht unital sein, da die beiden Einsen verschieden sein können, siehe 2.1.11.

(ROWEN [250](1988), Seiten 6, 7, 22; JACOBSON [152](1980), Seite 101). Sei M ein R -Linksmodul. Die Menge aller Unterlinksmoduln von M versehe man mit der partiellen Ordnung $\leq$,

$$M_1 \leq M_2 :\Leftrightarrow M_1 \text{ ist ein Unterlinksmodul von } M_2. \quad (2.20)$$

Es liegt dann ein Verband vor. Erfüllt $(\mathcal{P}(M), \leq)$ die DCC, so heißt M *linksartinsch* (im Englischen *left Artinian*). Rechtsmoduln entsprechend. Erfüllt $(\mathcal{P}(M), \leq)$ die ACC, so heißt M *linksnoethersch* (im Englischen *left Noetherian*). Rechtsmoduln entsprechend. Die in 1.1.13 präsentierten Äquivalenzen sind im Rahmen von kommutativen, noetherschen Ringen bei VAN DER WAERDEN [284](1967) als sogenannte *Teilerkettensätze* aufgeführt.

Ein Linksmodul $M \neq \{0\}$ heißt *einfach* (genauer *linkseinfach*, im Englischen *(left-)simple*), falls nur $\{0\}$ und M selbst Unterlinksmoduln von M sind. Dementsprechend heißt ein Ring $R \neq \{0\}$ *linkseinfach*, falls R , aufgefasst als Linksmodul über sich, linkseinfach ist, also außer $\{0\}$ und R kein Linksideal enthält. *Rechtseinfach* entsprechend. Ein Ring $R \neq \{0\}$ heißt *einfach* (im Englischen *simple*), wenn er kein von $\{0\}$ und R verschiedenes beidseitiges Ideal enthält. Direkt über die Formulierung mit einseitigen oder beidseitigen Idealen ist die Definition eines linkseinfachen, rechtseinfachen oder einfachen nichtassoziativen Ringes mit Operatorenbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$, entsprechend. Diese Definition von einfachen Ringen findet sich zum Beispiel in DIVINSKY [72](1965), GOODEARL [116](1976), JACOBSON [152](1974), MCCOY [209](1964) und ROWEN [250](1988). Bei dieser Definition ist zwar ein einfacher Ring R immer ungleich $\{0\}$, aber es kann durchaus $R \cdot R = \{0\}$ vorkommen; letztere Gleichung ist dann äquivalent zu $R^2 \neq R$. Es gibt aber auch Literaturstellen, wie zum Beispiel BEHRENS [23](1975), JACOBSON [150](1956) und SCHAFER [256](1966), bei der in der Definition noch zusätzlich die Bedingung $R^2 \neq \{0\}$ gefordert wird.

Sei M ein Linksmodul über einen Ring R , S eine Teilmenge von R und U eine Teilmenge von M . Dann ist mit $S \cdot U$ die Menge $\{su \in M : s \in S, u \in U\}$ gemeint. Analog wie in 2.1.13 meint SU die Menge aller Elemente der Form $s_1u_1 + \dots + s_nu_n$ mit $s_k \in S$, $u_k \in U$, $k \in \{1, \dots, n\}$, $n \in \mathbb{N}^\times$. Ist einer der beiden Faktoren S oder U einelementig, etwa $S = \{s\}$ (bzw. $U = \{u\}$), so schreibt man statt $S \cdot U$ bzw. SU einfach nur $s \cdot U$ bzw. sU (bzw. $S \cdot u$ bzw. Su). Für Rechtsmoduln alles entsprechend mit $U \cdot S$ bzw. US .

2.1.35 Definition (KLINGENBERG und KLEIN [186](1972), §22). Sei $\Psi: R \rightarrow S$ ein Ringhomomorphismus zwischen zwei Ringen R und S . Sei $(X, +, \cdot)$ ein R -Linksmodul und sei $(Y, +, \cdot)$ ein S -Linksmodul. Eine Abbildung $T: X \rightarrow Y$ heißt Ψ -*semilinear*, wenn gilt: T ist ein Gruppenhomomorphismus von $(X, +)$ nach $(Y, +)$ und

$$T(rx) = \Psi(r)T(x) \quad \text{für alle } r \in R, x \in X.$$

In gewissen Situationen kann es sinnvoll sein, Ψ als unital zu fordern.

Im Fall, dass $R = S$ gilt und Ψ die identische Abbildung auf R ist, heißt eine Ψ -semilineare Abbildung $X \rightarrow Y$ ein *R -Linksmodulhomomorphismus* oder *R -linear*, wobei man je nach vorliegendem Kontext das Ringsymbol R oder das Wort *Linksmodul* oder *Modul* weglässt.

Gilt $R = S = \mathbb{C}$ und ist Ψ die komplexe Konjugation in $\mathbb{C}$, dann wird anstelle von Ψ -semilinear die Bezeichnung *konjugiert-linear* verwendet.

Bei allgemeinen Betrachtungen mit $R = S = \mathbb{K}$ meint *semilinear*, dass die Abbildung Ψ im Fall von $\mathbb{K} = \mathbb{C}$ die konjugiert-komplexe Konjugation in $\mathbb{C}$ und im Fall von $\mathbb{K} = \mathbb{R}$ die identische Abbildung in $\mathbb{R}$ ist. Des Weiteren bezeichnet $\bar{z}$ für $z \in \mathbb{K}$ im Fall von $\mathbb{K} = \mathbb{R}$ die Zahl z selbst und im Fall von $\mathbb{K} = \mathbb{C}$ die zu z konjugiert-komplexe Zahl.

2.1.36 Definition. Sei R ein Ring, seien M_ν , $\nu \in I$, I eine Indexmenge, Linksmoduln über R . Das kartesische Produkt $\prod_{\nu \in I} M_\nu$, versehen mit komponentenweiser Addition und Skalarmultiplikation, ist ein Linksmodul über R und heißt das *direkte Produkt* der M_ν , $\nu \in I$, und wird wieder mit $\prod_{\nu \in I} M_\nu$ bezeichnet. Wie bereits in 1.1.2 definiert, gilt: Sind alle M_ν , $\nu \in I$, gleich, so wird für $I = \{1, \dots, n\}$ mit $n \in \mathbb{N}^\times$ wie üblich anstelle von $\prod_{\nu \in I} M_\nu$ mit $M := M_1$ auch die Bezeichnung M^n verwendet. Aber im Hinblick auf die Bezeichnungsweise in Definition 2.2.18 wird für den konkreten Fall $n = 1$ nicht M^1 für $\prod_{\nu \in \{1\}} M_\nu$, sondern immer einfach nur M geschrieben.

Die Teilmenge der I -Tupel $(x_\nu)_{\nu \in I}$ des kartesischen Produktes $\prod_{\nu \in I} M_\nu$ mit $x_\nu = 0$ für fast alle $\nu \in I$ ist ein Linksmodul über R und heißt die *direkte Summe* (genauer die *äußere direkte Summe*) der M_ν , $\nu \in I$, in Zeichen $\bigoplus_{\nu \in I} M_\nu$.

2.1.37 Definition mit Bemerkungen. Sei R ein Ring und seien M und N zwei Linksmoduln über R . Sei S eine Teilmenge von M . Dann wird mit $\text{span}(S)$ der kleinste Untermodul von M bezeichnet, der S enthält; dieser Untermodul heißt der von der Menge S erzeugte (oder auch *aufgespannte*) Untermodul von M . Für nichtleeres S gilt

$$\text{span}(S) = RS + \mathbb{Z}S;$$

also $\text{span}(S) = RS$, falls M dabei unital ist. Falls R ein Schiefkörper ist, so nennt man $\text{span}(S)$ auch die *lineare Hülle* von S . Falls $\text{span}(S) = M$ gilt, heißt S ein *Erzeugendensystem* von M .

Existiert ein $\mathbb{K}$ -Isomorphismus von M auf N , so wird dies durch $M \simeq N$ symbolisiert.

Sei $n \in \mathbb{N}^\times$. Sind $S_1, \dots, S_n$ Teilmengen von M , so ist der Ausdruck $S_1 + \dots + S_n$ definiert; er wird auch mit $\sum_{i=1}^n S_i$ bezeichnet. Ist S_ν , $\nu \in I$, eine Familie von Teilmengen von M , so setzt man $\sum_{\nu \in I} S_\nu := \bigcup \{ \sum_{i \in F} S_i : F \subseteq I \text{ endlich, } F \neq \emptyset \}$. $\sum_{\nu \in I} S_\nu$ wird die *algebraische Summe* der S_ν , $\nu \in I$, genannt. Sind alle $U_\nu := S_\nu$, $\nu \in I$, Untermoduln von M , so bezeichne man für alle $\nu \in I$ mit M_ν den Linksmodul U_ν und betrachte dann die Abbildung $\varphi: \bigoplus_{\nu \in I} M_\nu \rightarrow M$, $(x_\nu)_{\nu \in I} \mapsto \sum_{\nu \in I} x_\nu$. Dann ist offensichtlich $\sum_{\nu \in I} U_\nu = \text{ran}(\varphi)$ und φ eine R -lineare Abbildung; und falls φ dann auch noch injektiv ist, induziert φ eine Isomorphie $\bigoplus_{\nu \in I} M_\nu \simeq \sum_{\nu \in I} U_\nu$. Man sagt dann, die Untermoduln U_ν , $\nu \in I$, seien *R -linear unabhängig* und die Summe $\sum_{\nu \in I} U_\nu$ der Untermoduln U_ν , $\nu \in I$, sei eine *direkte algebraische Summe* (genauer eine *innere direkte algebraische Summe*) und dieser Untermodul von M wird mit $\sum_{\nu \in I}^\oplus U_\nu$ oder auch mit $\sum_{\nu \in I} \oplus U_\nu$ bezeichnet. Falls keine Verwechslung mit der äußeren direkten Summe der M_ν , $\nu \in I$, zu befürchten ist, wird anstelle von $\sum_{\nu \in I}^\oplus U_\nu$ auch die Bezeichnung $\bigoplus_{\nu \in I} U_\nu$ verwendet. Ist die Indexmenge I endlich und nicht leer, etwa $I = \{1, \dots, n\}$ für ein $n \in \mathbb{N}^\times$, so wird anstelle von $\sum_{i \in I}^\oplus U_i$ auch $U_1 \dot{+} \dots \dot{+} U_n$ geschrieben, was wiederum, wenn keine Verwechslung möglich ist, auch als $U_1 \oplus \dots \oplus U_n$ geschrieben wird.

2.1.38 Torsion. Sei M ein Linksmodul über einen Ring R . Betrachte das Annihilatorsystem $(R, M; \cdot; M)$. Wenn es in R trennend ist — mit anderen Worten, wenn $M_\perp = \{0\}$ gilt —, heißt M *treu*. Ein $x \in M$ heißt ein *Torsionselement*, wenn $\{x\}_\perp \neq \{0\}$ gilt. Hat M kein von Null verschiedenes Torsionselement, so heißt M *torsionsfrei* (und R ist dann auch notwendigerweise nullteilerfrei).

Falls R eine Eins hat, dann ist M wegen $x = ex + (x - ex)$, $x \in M$, gleich der inneren direkten Summe $M = (eM) \dot{+} R^\perp$; falls also M unital ist, folgt $R^\perp = \{0\}$, das heißt, $(R, M; \cdot; M)$ ist in M trennend.

2.1.39 Freie Moduln (GOLDHABER und EHRLICH [115](1970), Seite 156). Sei M ein Linksmodul über einen Ring R , I eine Menge und α eine Abbildung $I \rightarrow M$. Dann heißt $M := (M, \alpha)$ ein *freier Linksmodul über R (auf I)*, wenn die folgende Bedingung erfüllt ist: Zu jedem Linksmodul N über R und zu jeder Abbildung $\gamma: I \rightarrow N$, gibt es einen eindeutig bestimmten R -Modulhomomorphismus $\mu: M \rightarrow N$ mit $\mu \circ \alpha = \gamma$.

Als Folgerungen dieser Definition hat man für einen beliebigen Ring R (siehe ebd.):

(a) Ist (M, α) ein freier R -Linksmodul auf einer Menge I , dann ist α injektiv und $\text{span}(\alpha(I)) = M$.

(b) Sind (M, α) und (N, β) freie R -Linksmoduln auf einer Menge I , so gibt es einen eindeutig bestimmten R -Isomorphismus $\mu: M \rightarrow N$ mit $\mu \circ \alpha = \beta$.

(c) Jedes isomorphe Bild eines freien R -Linksmoduls ist frei. Genauer: Sei (M, α) ein freier R -Linksmodul auf einer Menge I , β eine Abbildung von I nach einem R -Linksmodul N und $\mu: M \rightarrow N$ ein R -Isomorphismus mit $\mu \circ \alpha = \beta$. Dann ist (N, β) auch ein freier R -Linksmodul auf I .

Des Weiteren gilt (siehe ebd.): Für jede Menge I und jeden Ring R existiert ein freier R -Linksmodul auf I . Folglich ist jeder R -Linksmodul ein homomorphes Bild eines freien R -Linksmoduls.

Sei nun M ein *unitaler* Linksmodul über einen Ring R , I eine Menge und α eine Abbildung $I \rightarrow M$. Dann ist $M = (M, \alpha)$ genau dann ein freier R -Linksmodul auf I , wenn sowohl $M = \sum_{\nu \in I}^{\oplus} R\alpha(\nu)$ gilt als auch $\alpha(\nu)$ für keines der $\nu \in I$ ein Torsionselement ist.

Nun gilt (siehe ebd., Seite 155): Sei R ein Ring und M ein R -Linksmodul mit $M = \text{span}(\{x\})$ für ein $x \in M$. Dann ist die Abbildung $\mu_x: R \rightarrow M, r \mapsto rx$, mit R aufgefasst als ein Linksmodul über sich, ein R -Homomorphismus. Falls x kein Torsionselement ist, ist μ_x injektiv. Falls M unital ist, ist μ_x surjektiv. Also: Ist M ein unitaler Linksmodul über einen Ring R , wobei M von einem Element erzeugt wird, welches kein Torsionselement ist, so ist M gleich Rx und isomorph zu R als einem R -Linksmodul.

Es folgt (siehe ebd.): Ein unitaler Linksmodul über einen Ring R ist genau dann frei, wenn es eine Menge I gibt, so dass $M = \sum_{\nu \in I}^{\oplus} R\nu$ ist, wobei für jedes $\nu \in I$ $R\nu$ ein zu R , aufgefasst als ein Linksmodul über sich, isomorphes Untermodul von M ist.

Man kann also sagen: Ein unitaler Linksmodul über einen Ring R ist genau dann frei, wenn er isomorph zu einer äußeren direkten Summe von Kopien von R ist.

Man beachte: Hat man ein Linksmodul über einen Ring R vorliegen, wobei R zwar eine Eins hat, der Linksmodul aber nicht unital ist, so kann der Linksmodul zwar frei sein, muss aber nicht mehr notwendig isomorph zu einer äußeren direkten Summe von Kopien von R sein. Ein konkretes Beispiel für diese Situation findet man bei HUNGERFORD [143](1980), Seite 188, Übung 2.

(GOLDHABER und EHRLICH, ebd., Seite 158)). Ist (M, α) ein freier Linksmodul über einen Ring R auf einer Menge I und $\varepsilon: \alpha(I) \hookrightarrow M$ die kanonische Einbettung von $\alpha(I)$ in M , so ist (M, ε) ein freier R -Linksmodul auf $\alpha(I)$.

Beweis. Nach Voraussetzung gibt es zu jeder Abbildung $\gamma: I \rightarrow N$, N ein R -Linksmodul, einen eindeutig bestimmten R -Homomorphismus $\mu = \mu_\gamma: M \rightarrow N$. Bezeichne $\upharpoonright \alpha$ die Restriktion $I \rightarrow \alpha(I)$ von α auf ihrem Bild. Dann gilt: $\mu_\gamma \circ \alpha = \mu_\gamma \circ \varepsilon \circ (\upharpoonright \alpha) = \gamma$, also $\mu_\gamma \circ \varepsilon = \gamma \circ (\upharpoonright \alpha)^{-1}$ und da zu jeder Abbildung $\gamma': \alpha(I) \rightarrow N$ per $\gamma_{\gamma'} := \gamma' \circ (\upharpoonright \alpha)$ eine Abbildung $I \rightarrow N$ mit $\gamma_{\gamma'} \circ (\upharpoonright \alpha)^{-1} = \gamma'$ existiert, gilt $\mu_{\gamma_{\gamma'}} \circ \varepsilon = \gamma'$. Somit ist $\mu_{\gamma' \circ (\upharpoonright \alpha)}$ der zu jeder Abbildung $\gamma': \alpha(I) \rightarrow N$ eindeutig bestimmte R -Modulhomomorphismus $M \rightarrow N$. $\square$

Folglich ist M ein freier R -Linksmodul auf einer ihn selbst erzeugenden Teilmenge von sich. Somit definiert man:

Sei M ein Linksmodul über einen Ring R . Eine Teilmenge B von M heißt eine *Basis* für M , wenn (M, ε) ein freier R -Linksmodul auf B ist, wobei ε die kanonische Einbettung von B in M ist. Eine Teilmenge S von M heißt *frei* für M , wenn S eine Basis für $\text{span}(S)$ ist. Eine Teilmenge S von M heißt *linear unabhängig*, wenn für jedes $n \in \mathbb{N}^\times$ für beliebige, aber paarweise verschiedene $x_1, \dots, x_n \in S$ für alle $r_1, \dots, r_n \in R$ aus $r_1x_1 + \dots + r_nx_n = 0$ stets $r_1 = \dots = r_n = 0$ folgt.

Man bemerke, dass die leere Menge eine Basis für $M = \{0\}$ ist und gemäß Definition 2.1.1 linear unabhängig ist. Offensichtlich kann eine linear unabhängige Menge kein Torsionselement enthalten. Ist B eine Basis eines R -Linksmoduls M , so gilt $\text{span}(B) = M$. Somit gilt:

Eine (notwendig nicht leere) Teilmenge B von M ist genau dann eine Basis eines R -Linksmoduls M mit $M \neq \{0\}$, wenn B eine freie, M erzeugende Menge ist.

Sei M ein unitaler Linksmodul über einen Ring R . Dann ist eine Teilmenge S von M genau dann frei, wenn sie linear unabhängig ist.

Beweis. Setze $N := \text{span}(S)$ und sei ε die kanonische Einbettung von S in N . Dann ist S genau dann frei für N , wenn (N, ε) ein freier R -Linksmodul auf S ist. Da mit M auch N unital ist, ist gemäß der ersten oben angegebenen äquivalenten Formulierung eines freien unitalen Linksmoduls dies genau dann der Fall, wenn $N = \sum_{s \in S}^{\oplus} R\varepsilon(s) = \sum_{s \in S}^{\oplus} Rs$, wobei jedes $s \in S$ kein Torsionselement ist. Die Existenz solch einer Zerlegung von N ist nun aber äquivalent zu der linearen Unabhängigkeit von S . $\square$

Gemäß HUNGERFORD [143](1980), Seite 181, hat man, ähnlich wie für Ringe, die nicht notwendig eine Eins haben, die folgende Charakterisierung im Falle von unitalen Linksmoduln:

Ein unitaler R -Linksmodul M mit $M \neq \{0\}$ ist genau dann frei, wenn eine Menge I und eine Abbildung $\alpha: I \rightarrow M$ existiert, so dass zu jedem unitalen R -Linksmodul N und zu jeder Abbildung $\gamma: I \rightarrow N$ ein eindeutig bestimmter R -Modulhomomorphismus $\mu: M \rightarrow N$ mit $\mu \circ \alpha = \gamma$ existiert.

Sei M ein freier Linksmodul über einen Ring R . Falls alle Basen von M dieselbe Kardinalität aufweisen, heißt M *dimensioniert* (im Englischen *dimensional*) und in diesem Fall heißt die besagte Kardinalität die *Dimension* von M . Ein Ring R mit der Eigenschaft, dass jeder freie R -Linksmodul M — der im Fall, dass R eine Eins hat, unital sein soll — dimensioniert ist, heißt ein *dimensionierter Ring*.

Jeder unitale freie Linksmodul, der eine unendlichdimensionale Basis hat, ist dimensioniert. Es existieren freie Linksmoduln, die endliche Basen verschiedener Kardinalität haben.

Jeder Ring mit Eins, der ein homomorphes Bild hat, welches ein Schiefkörper ist, ist dimensioniert. Da jeder von $\{0\}$ verschiedene Ring mit Eins ein maximales Linksideal (zur Definition siehe 2.1.50) hat, hat man als Spezialfall, dass ein Ring mit Eins dimensioniert ist, falls er ein homomorphes Bild hat, welches ein von $\{0\}$ verschiedener kommutativer Ring ist.

2.1.40 Bemerkung und Definition. Sei R ein nichtassoziativer Ring. Dann ist die Menge $R \times \mathbb{Z}$, versehen mit komponentenweise definierter Addition

$$(r_1, z_1) + (r_2, z_2) := (r_1 + r_2, z_1 + z_2) \quad \text{für alle } r_1, r_2 \in R, z_1, z_2 \in \mathbb{Z},$$

additivem neutralen Element $0 := (0, 0)$ und Multiplikation

$$(r_1, z_1)(r_2, z_2) := (r_1 r_2 + z_1 r_2 + z_2 r_1, z_1 z_2) \quad \text{für alle } r_1, r_2 \in R, z_1, z_2 \in \mathbb{Z}$$

ein nichtassoziativer Ring mit Eins $(0, 1)$, der mit R^1 bezeichnet wird.

Die Abbildung $R \rightarrow R^1, r \mapsto (r, 0)$ ist ein Ringmonomorphismus. Ist R assoziativ, so ist auch R^1 assoziativ. Jedes Linksideal (bzw. Rechtsideal) von R ist auch ein Linksideal (bzw. Rechtsideal) von R^1 .

Siehe auch 2.3.19.

2.1.41 Satz. Sei R ein nichtassoziativer Ring ohne Eins. Dann ist $r \in R$ genau dann quasi-invertierbar, wenn $e - r \in R^1$ invertierbar ist.

2.1.42 Definition. Sei R ein nichtassoziativer Ring mit Operatorenbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$. Ein Linksideal $\mathfrak{a}$ von R heißt *nilpotent*, wenn es ein $n \in \mathbb{N}^\times$ gibt, so dass für alle $a_1, \dots, a_n \in \mathfrak{a}$ das Produkt $a_1 \cdots a_n$ für jede Assoziation, sprich Beklammerung, gleich Null ist. Ein $r \in R$ heißt *nilpotent*, wenn ein $n \in \mathbb{N}^\times$ existiert, so dass für mindestens eine Beklammerung $r^n = 0$ gilt (BEHRENS [22](1954), Seite 442). Ein Linksideal $\mathfrak{a}$ von R heißt *nil* (oder genauer auch ein *nils Linksideal*), wenn jedes $a \in \mathfrak{a}$ nilpotent ist. Entsprechend für Rechtsideal und Ideal. Der K -Ring R heißt *halbprim* (im Englischen *semiprime*), wenn R kein nilpotentes Ideal $\mathfrak{a} \neq \{0\}$ enthält. Siehe auch die Anmerkung in 2.1.44! Ein Ideal $\mathfrak{a}$ des K -Ringes R heißt *halbprim* (im Englischen *semiprime*), falls der K -Ring $R/\mathfrak{a}$ halbprim ist (PALMER [231](1994), Seite 482).

2.1.43 Primideale und Primringe (MCCOY [209](1964), Seite 62). Sei R ein Ring mit Operatorenbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$, und $\mathfrak{p}$ ein Ideal von R . Die folgenden Aussagen sind äquivalent:

- (i) Für alle Ideale $\mathfrak{a}$ und $\mathfrak{b}$ von R mit $\mathfrak{a}\mathfrak{b} \subseteq \mathfrak{p}$ gilt $\mathfrak{a} \subseteq \mathfrak{p}$ oder $\mathfrak{b} \subseteq \mathfrak{p}$.
- (ii) Für alle $a, b \in R$ mit $aRb \subseteq \mathfrak{p}$ gilt $a \in \mathfrak{p}$ oder $b \in \mathfrak{p}$.
- (ii') Für alle $a, b \in R \setminus \mathfrak{p}$ existiert ein $x \in R$ mit $axb \in R \setminus \mathfrak{p}$.
- (iii) Für alle Hauptideale (a) und (b) von R mit $(a)(b) \subseteq \mathfrak{p}$ gilt $a \in \mathfrak{p}$ oder $b \in \mathfrak{p}$.
- (iv) Für alle Linksideale $\mathfrak{a}$ und $\mathfrak{b}$ von R mit $\mathfrak{a}\mathfrak{b} \subseteq \mathfrak{p}$ gilt $\mathfrak{a} \subseteq \mathfrak{p}$ oder $\mathfrak{b} \subseteq \mathfrak{p}$.
- (v) Für alle Rechtsideale $\mathfrak{a}$ und $\mathfrak{b}$ von R mit $\mathfrak{a}\mathfrak{b} \subseteq \mathfrak{p}$ gilt $\mathfrak{a} \subseteq \mathfrak{p}$ oder $\mathfrak{b} \subseteq \mathfrak{p}$.

Beweis. (i) $\Rightarrow$ (ii): Gelte (i) und sei $aRb \subseteq \mathfrak{p}$. Insbesondere ist dann $aRRb \subseteq \mathfrak{p}$ und da $\mathfrak{p}$ ein Ideal ist, auch $RaRRbR \subseteq \mathfrak{p}$. Nach Voraussetzung somit $RaR \subseteq \mathfrak{p}$ oder $RbR \subseteq \mathfrak{p}$. Liege der Fall $RaR \subseteq \mathfrak{p}$ vor. Setze $\mathfrak{c} := (a)$. Dann gilt $\mathfrak{c}\mathfrak{c} \subseteq R\mathfrak{c}R \subseteq RaR \subseteq \mathfrak{p}$. Nach Voraussetzung also $\mathfrak{c} \subseteq \mathfrak{p}$, also $a \in \mathfrak{p}$. Hätte der andere Fall vorgelegen, also $RbR \subseteq \mathfrak{p}$, so wäre analog $b \in \mathfrak{p}$ gefolgt.

(ii) $\Leftrightarrow$ (ii'): Klar.

(ii) $\Rightarrow$ (iii): Gelte (ii) und sei $(a)(b) \subseteq \mathfrak{p}$. Dann gilt wegen $aRb \subseteq aR(b) \subseteq a(b) \subseteq (a)(b)$ die Inklusion $aRb \subseteq \mathfrak{p}$, also nach Voraussetzung $a \in \mathfrak{p}$ oder $b \in \mathfrak{p}$.

(iii) $\Rightarrow$ (iv): Gelte (iii) und seien $\mathfrak{a}$ und $\mathfrak{b}$ Linksideale von R mit $\mathfrak{a}\mathfrak{b} \subseteq \mathfrak{p}$. Liege der Fall vor, dass $\mathfrak{a} \not\subseteq \mathfrak{p}$. Dann gibt es ein $a \in \mathfrak{a} \setminus \mathfrak{p}$. Sei $b \in \mathfrak{b}$. $(a)(b)$ ist eine Teilmenge der Menge aller endlichen Summen von Termen der Form $r_1 a r_2 b r_3$ mit $r_1, r_2, r_3 \in R \cup K$. Da $\mathfrak{a}$ und $\mathfrak{b}$ Linksideale sind, sind die Elemente des Produkts $\mathfrak{a}\mathfrak{b}$ endliche Summen von Termen der Form $r_1 a r_2 b$ mit $r_1, r_2 \in R \cup K$. Also $(a)(b) \subseteq \mathfrak{a}\mathfrak{b} + \mathfrak{a}\mathfrak{b}R \subseteq \mathfrak{p}$. Da $a \notin \mathfrak{p}$, muss nach Voraussetzung $b \in \mathfrak{p}$ gelten, im vorliegenden Fall also $\mathfrak{b} \subseteq \mathfrak{p}$. Hätte der andere Fall vorgelegen, so wäre analog $\mathfrak{a} \subseteq \mathfrak{p}$ gefolgt.

(iii) $\Rightarrow$ (v): Geht analog wie (iii) $\Rightarrow$ (iv).

(iv) $\Rightarrow$ (i) und (v) $\Rightarrow$ (i) sind klar. □

Liegt eine (und damit alle) der vorstehenden äquivalenten Aussagen vor, so heißt $\mathfrak{p}$ ein *Primideal* (im Englischen *prime ideal*). Liegt der spezielle Fall vor, dass $\mathfrak{p} := \{0\}$ ein Primideal von R ist, so heißt R ein *Prim-K-Ring* (im Englischen *prime K-ring*). Mit dieser Bezeichnung hat man zusätzlich die zu jeder der Aussagen (i) bis (v) äquivalente Aussage:

- (vi) $R/\mathfrak{p}$ ist ein Prim-K-Ring.

Man bemerke, dass ein Prim-K-Ring keine von $\{0\}$ verschiedenen nilpotenten Ideale besitzen kann, und dass jeder einfache, nicht nilpotente K-Ring ein Prim-K-Ring ist.

(GRAY [120](1970), Seite 13). Ist speziell R ein kommutativer Ring mit Eins, so ist ein Ideal $\mathfrak{p}$ von R genau dann ein Primideal, wenn $ab \in \mathfrak{p}$ immer $a \in \mathfrak{p}$ oder $b \in \mathfrak{p}$

impliziert. Dass diese Äquivalenz im nichtkommutativen Fall nicht stimmen muss, sieht man zum Beispiel am Ring der 2×2 -Matrizen über $\mathbb{K}$: Dort ist $\{0\}$ ein Primideal, aber es gibt von der Nullmatrix verschiedene Matrizen, deren Produkt die Nullmatrix ist (DIVINSKY [72](1965), Seite 60).

Man beachte, dass nach der hier gegebenen Definition jeder K-Ring R ein Primideal von R ist. Diese Eigenschaft wird in den entsprechenden Definitionen eines Primideals in den folgenden Literaturstellen zugelassen: BEHRENS [23](1975), Seite 81, DIVINSKY [72](1965), Seite 59, HORNFECK [141](1976), Seite 149, JACOBSON [150](1956), Seite 194, MCCOY [209](1964), Seite 64 und ROWEN [250](1988), Seite 163.

Im Zahlenring $\mathbb{Z}$ gilt für jedes $p \in \mathbb{Z}^\times \setminus \{-1, 1\}$, dass das Hauptideal (p) genau dann ein Primideal von $\mathbb{Z}$ ist, wenn p eine Primzahl ist.

Ein gewisser Makel der oben gegebenen Definition des Primideals ist somit, dass zwar $(1) = \mathbb{Z}$ ein Primideal ist, aber 1 keine Primzahl ist. Damit ist nachzuvollziehen, dass bei einigen Autoren, wie zum Beispiel ARTIN [10](1993), DALES [56](2000), GRAY [120](1970), MEYBERG [214](1980), PALMER [231](1994), PIERCE [239](1982) und SCHULZE [258](2006), der Ring R als Primideal ausgeschlossen wird, also selbst kein Primideal sein kann. Trotzdem ist die gegebene Definition, wo R immer ein Primideal ist, keine Einschränkung; man braucht bei den zuletzt genannten Autoren nur an entsprechender Stelle den Zusatz „ $\neq R$ “ einfügen. So gilt zum Beispiel, siehe SCHULZE [258](2006), Seite 79, der folgende Satz:

Ist R ein Integritätsbereich mit Eins und $p \in R^\times$ ein nicht invertierbares Element von R , so ist p genau dann ein Primelement, wenn (p) ein *von R verschiedenes* Primideal ist.

(GRAY [120](1970), Seite 14). Sei R ein kommutativer Ring und $\mathfrak{a}$ ein Ideal von R . Dann wird als *Radikal von $\mathfrak{a}$* , in Zeichen $\sqrt{\mathfrak{a}}$, die Menge $\{x \in R : \exists n \in \mathbb{N}^\times : x^n \in \mathfrak{a}\}$ bezeichnet; es ist ein Ideal von R mit $\mathfrak{a} \subseteq \sqrt{\mathfrak{a}}$.

2.1.44 Halbprim (MCCOY [209](1964), Seite 66; PALMER [231](1994), Seite 483). Sei R ein Ring mit Operatorbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$, und $\mathfrak{p}$ ein Ideal in R . Dann sind äquivalent:

- (i) Für alle Ideale $\mathfrak{a}$ von R mit $\mathfrak{a}^2 \subseteq \mathfrak{p}$ gilt $\mathfrak{a} \subseteq \mathfrak{p}$.
- (ii) Für alle $a \in R$ mit $aRa \subseteq \mathfrak{p}$ gilt $a \in \mathfrak{p}$.
- (ii') Für alle $a \in R \setminus \mathfrak{p}$ existiert ein $x \in R$ mit $axa \in R \setminus \mathfrak{p}$.
- (iii) Für alle Hauptideale (a) von R mit $(a)^2 \subseteq \mathfrak{p}$ gilt $a \in \mathfrak{p}$.
- (iv) Für alle einseitigen Ideale $\mathfrak{a}$ von R mit $\mathfrak{a}^2 \subseteq \mathfrak{p}$ gilt $\mathfrak{a} \subseteq \mathfrak{p}$.
- (v) Für alle einseitigen Ideale $\mathfrak{a}$ von R für die ein $n \in \mathbb{N}^\times$ existiert mit $\mathfrak{a}^n \subseteq \mathfrak{p}$ gilt $\mathfrak{a} \subseteq \mathfrak{p}$.
- (vi) $\mathfrak{p}$ ist ein halbprimales Ideal von R . (Definition 2.1.42)

Beweis. Bis auf (i) $\Rightarrow$ (v) $\Rightarrow$ (vi) alles analog wie in 2.1.43.

(i) $\Rightarrow$ (v): Sei $\mathfrak{a}$ ein Linksideal mit $\mathfrak{a}^n \subseteq \mathfrak{p}$. Dann ist $\mathfrak{a}(R^1)$ ein Ideal mit $(\mathfrak{a}R^1)^n \subseteq \mathfrak{a}(R^1\mathfrak{a})^{n-1}R^1 \subseteq \mathfrak{a}^n R^1 \subseteq \mathfrak{p}R^1 = \mathfrak{p}$. Sei $m \in \mathbb{N}^\times$ die kleinste Zahl mit $(\mathfrak{a}R^1)^m \subseteq \mathfrak{p}$. Angenommen, es wäre $m > 1$, so würde $((\mathfrak{a}R^1)^{m-1})^2 \subseteq \mathfrak{p}$ gelten und nach (a) würde $(\mathfrak{a}R^1)^{m-1} \subseteq \mathfrak{p}$ folgen, Widerspruch zur Minimalität. Also ist $m = 1$ und $\mathfrak{a} \subseteq \mathfrak{a}R^1 \subseteq \mathfrak{p}$.

(v) $\Rightarrow$ (vi): Gelte (v), dann ist zu zeigen, dass $R/\mathfrak{a}$ kein nilpotentes Ideal $\mathfrak{b} \neq \{0\}$ enthält. Sei $\mathfrak{b}^n = \{0\}$. Als eine Folge des Homomorphiesatzes für Ringe lässt sich jedes Ideal von $R/\mathfrak{a}$ eindeutig in der Form $\mathfrak{c}/\mathfrak{a}$ schreiben, wobei $\mathfrak{c}$ ein Ideal von R ist, das $\mathfrak{a}$ umfasst (MCCOY [209](1964), Seite 46, Theorem 2.45(iv)). Es ist hier also $\mathfrak{b} = \mathfrak{c}/\mathfrak{a}$ für ein eindeutig bestimmtes Ideal $\mathfrak{c}$ von R mit $\mathfrak{a} \subseteq \mathfrak{c}$. Somit gilt in $R/\mathfrak{a}$ die Gleichung $(\mathfrak{c}/\mathfrak{a})^n = \{\mathfrak{a}\}$, also $\mathfrak{c}^n \subseteq \mathfrak{a}$. Nach Voraussetzung folgt $\mathfrak{c} \subseteq \mathfrak{a}$, das heißt, $\mathfrak{c}/\mathfrak{a}$ ist bereits gleich $\mathfrak{a}$ in $R/\mathfrak{a}$, mit anderen Worten ist $\mathfrak{b} = \{0\}$. $\square$

Da für $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$ in jedem halbprimen K -Ring das Ideal $\{0\}$ halbprim ist, bemerke man, dass wegen (v) jeder halbprime K -Ring keine nilpotenten, einseitigen Ideale $\mathfrak{a}$ mit $\mathfrak{a} \neq \{0\}$ enthalten kann.

2.1.45 Ringordnungen I. Sei R ein Ring. Auf der Menge der idempotenten Elemente von R wird durch

$$\begin{aligned} x \leq_r y &: \Leftrightarrow yx = x, \\ x \leq_\ell y &: \Leftrightarrow xy = x \end{aligned}$$

jeweils eine Präordnung und durch

$$x \leq_{\ell r} y : \Leftrightarrow xy = yx = x$$

eine partielle Ordnung definiert. Im vorliegenden Abschnitt wird anstelle von $\leq_{\ell r}$ einfach nur das Symbol $\leq$ verwendet.

Beweis. Wegen $x \leq_r x \Leftrightarrow x^2 = x$ ist $\leq_r$ reflexiv. Seien x, y, z idempotente Elemente aus R mit $x \leq_r y$ und $y \leq_r z$; also $yx = x$ und $zy = y$. Somit $zx = zyx = yx = x$, das heißt, $x \leq_r z$ und $\leq_r$ ist transitiv. „ $\leq_\ell$ “ und „ $\leq$ “ gehen analog. $\square$

Es gelten die beiden Äquivalenzen $x \leq_r y \Leftrightarrow xR \subseteq yR$ und $x \leq_\ell y \Leftrightarrow Rx \subseteq Ry$.

Beweis. Seien x, y idempotente Elemente aus R mit $x \leq_r y$. Sei $a \in R$. Dann ist $xa = yxa \in yR$; das zeigt $xR \subseteq yR$. Gelte nun $xR \subseteq yR$ mit $x, y \in R$ idempotent. Dann gilt insbesondere $xx = yb$ für ein $b \in R$. Somit: $yx = yyb = yb = x$, das heißt, $x \leq_r y$. „ $\leq_\ell$ “ geht analog. $\square$

Somit ist zum Beispiel für $x, y \in \text{Idem}(R)$ genau dann $Rx = Ry$, wenn sowohl $xy = x$ als auch $yx = y$ gilt. Dies sieht man auch direkt mit Gleichung (2.9) aus 2.1.33. Auf $\text{Idem}(R)$ ist dadurch eine Äquivalenzrelation $\equiv_\ell$ erklärt. (Und auch auf der Menge der weiter unten in 2.1.47 definierten schiefminimal idempotenten Elemente.) Die p enthaltende Äquivalenzklasse wird mit $[p]_\ell$ bezeichnet. Die Äquivalenzrelation $\equiv_r$ definiert sich entsprechend.

Seien $x, y \in \text{Idem}(R)$. Dann gilt die Äquivalenz

$$x \leq y \Leftrightarrow y - x \in \text{Idem}(R) \text{ und } x \perp y - x, \quad (2.21)$$

mit anderen Worten also

$$x \leq y \Leftrightarrow \exists z \in x^\perp \cap \text{Idem}(R) : y = x + z.$$

Beweis. Sei $x \leq y$, also $xy = yx = x$ und somit $(y - x)^2 = y - yx - xy + x = y - x$; des Weiteren $x(y - x) = xy - x = 0 = yx - x = (y - x)x$. Sei nun umgekehrt $y - x$ idempotent und orthogonal zu x . Dann gilt $xy = xy + x(x - y) = x(y + x - y) = x$ und $yx = (x - y)x + yx = x$. $\square$

(JACOBSON [150](1964), Seite 49). Wenn $p_1, \dots, p_n$ paarweise orthogonale, idempotente Elemente von R sind, dann ist auch $p := p_1 + \dots + p_n$ idempotent. Es ist $p_i \leq p$ für alle $i \in \{1, \dots, n\}$. Ist q ein idempotentes Element von R , so gilt dann genau dann $p_i \leq q$ für alle $i \in \{1, \dots, n\}$, wenn $p \leq q$ gilt.

Seien p und q zwei kommutative, idempotente Elemente von R . Sei $\circ$ das Quasi-Produkt aus 2.1.9(d). Dann sind auch $\ell := pq$ und $u := p \circ q$ idempotent. Bezüglich $\leq$ ist ℓ gleich dem Infimum $\inf\{p, q\}$ und u ist gleich dem Supremum $\sup\{p, q\}$. Allgemeiner gilt: Sind $p_1, \dots, p_n$ kommutative Idempotente, so sind $p_1 p_2 \dots p_n$ und $p_1 \circ p_2 \circ \dots \circ p_n$ idempotent mit $p_1 p_2 \dots p_n = \inf\{p_1, \dots, p_n\}$ und $p_1 \circ p_2 \circ \dots \circ p_n = \sup\{p_1, \dots, p_n\}$.

Seien $p_1, \dots, p_n$ Idempotente mit $p_i p_j = 0$ für alle $i, j \in \{1, \dots, n\}$ mit $i > j$. Dann ist $p_1 \circ p_2 = \sup\{p_1, p_2\}, \dots, p_1 \circ \dots \circ p_n = \sup\{p_1, \dots, p_n\}$.

2.1.46 Peirce-Zerlegung II (JACOBSON [150](1964), Seite 49). Sei R ein Ring und P eine endliche Menge von paarweise orthogonalen Idempotenten aus R , etwa $P = \{p_1, \dots, p_n\}$. Setze $p := p_1 + \dots + p_n$. Dann gelten die folgenden inneren direkten Summenzerlegungen von R :

$$R = p_1 R + \dots + p_n R + (1 - p)R, \quad (2.22)$$

$$R = R p_1 + \dots + R p_n + R(1 - p), \quad (2.23)$$

$$R = \sum_{i,j} \oplus p_i R p_j + p R(1 - p) + (1 - p) R p + (1 - p) R(1 - p). \quad (2.24)$$

Der Fall $n = 1$ wurde in 2.1.33 behandelt. Falls R eine Eins hat und diese gleich p ist, vereinfachen sich die drei vorstehenden Gleichungen offenbar zu

$$R = \sum_{i=1}^n \oplus p_i R, \quad R = \sum_{i=1}^n \oplus R p_i, \quad \text{und} \quad R = \sum_{i,j} \oplus p_i R p_j.$$

Die Zerlegungen (2.22), (2.23) und (2.24) heißen jeweils die *rechte*, *linke* und *beidseitige* (verallgemeinerte) *Peirce-Zerlegung* von R bezüglich der endlichen Menge $P = \{p_1, \dots, p_n\}$ von paarweise orthogonalen Idempotenten von R . Es gilt der folgende Satz (ebd., Seite 50):

Sei R ein Ring mit Eins. Bezüglich der R zugrunde liegenden additiven Gruppe $(R, +)$ sei $R = \mathfrak{a}_1 + \dots + \mathfrak{a}_n$ eine innere direkte Summe von Rechtsidealen $\mathfrak{a}_1, \dots, \mathfrak{a}_n$. Dann gibt es eine endliche Menge $P = \{p_1, \dots, p_n\}$ von paarweise orthogonalen Idempotenten von R mit $\mathfrak{a}_i = p_i R$ für alle $i \in \{1, \dots, n\}$.

Beweis. Sei $a \in R$. Es ist $a = a_1 + \dots + a_n$ für eindeutig bestimmte $a_i \in \mathfrak{a}_i$, insbesondere ist $e = r_1 + \dots + r_n$ für eindeutig bestimmte $r_i \in \mathfrak{a}_i$. Also $a = ea = r_1 a + \dots + r_n a$ mit $r_i a \in \mathfrak{a}_i$. Somit $a_i = r_i a$. Also $r_i r_i = r_i(0 + \dots + r_i + \dots + 0) = r_i$ und $r_i r_j = r_i(0 + \dots + r_j + \dots + 0) = 0$ für $i \neq j$. Die $r_1, \dots, r_n$ sind also paarweise orthogonale Idempotenten und $\mathfrak{a}_i = r_i R$. $\square$

Siehe weiter bei 2.3.14.

2.1.47 Minimale und schiefminimale Idempotenten.

Sei R ein Ring und gelten die Bezeichnungen der in 2.1.45 definierten Präordnungen und partiellen Ordnung. Sei x ein von Null verschiedenes, idempotentes Element von R . Dann sind die folgenden Aussagen äquivalent:

- (a) Es gibt keine zwei orthogonale $u, v \in \text{Idem}(R)^\times$ mit $x = u + v$; mit anderen Worten gibt es keine orthogonale, nicht triviale Zerlegung von x (im Englischen *orthogonal decomposition*).
- (b) $\text{Idem}(xRx)^\times = \{x\}$, das heißt in Worten, dass x das einzige von Null verschiedene idempotente Element in xRx ist.
- (c) x ist ein minimales Element in $(\text{Idem}(R)^\times, \leq)$.
- (c') Es gibt kein $y \in \text{Idem}(R)^\times$ mit $y \neq x$ und $y \leq x$.

Beweis. (b) $\Leftrightarrow$ (c') (DICKSON [64](1927), Seite 105, Hilfssatz). „ $\Rightarrow$ “: Per Kontraposition; sei also $y \in \text{Idem}(R)^\times$ mit $y \neq x$ und $y \leq x$. Dann ist $y = xy = yx = yx \in \text{Idem}(xRx)^\times \setminus \{x\}$. „ $\Leftarrow$ “: Per Kontraposition; sei also $y \in \text{Idem}(xRx)^\times \setminus \{x\}$, etwa $y = xrx$ für ein $r \in R$. Dann ist $xy = xrx = xrx = y = xrx = xrx = yx$, also $y \leq x$.

(a) $\Leftrightarrow$ (c') Wegen der in 2.1.45 aufgeführten Äquivalenz (2.21) im Grunde genommen klar. Nichtsdestotrotz sei hier der elementare Beweis, wie er in DIVINSKY [72](1965), Seite 29, Lemma 21, zu finden ist, präsentiert: „ $\Rightarrow$ “: Gelte (a). Angenommen, es existiert ein $u \in \text{Idem}(R)^\times$ mit $u \neq x$ und $u \leq x$. Setze $v := x - u$. Dann gilt $v^2 = (x - u)(x - u) =$

$x - xu - ux + u = x - u - u + u = x - u = v$, also $v \in \text{Idem}(R)^\times$. Des Weiteren $uv = u(x - u) = ux - u = u - u = 0$ und $vu = (x - u)u = xu - u = u - u = 0$, Widerspruch. „ $\Leftarrow$ “: Per Kontraposition; sei also $x = y + v$ mit zwei orthogonalen $y, v \in \text{Idem}(R)^\times$. Also $y \neq x$, $xy = (y + v)y = y$ und $yx = y(y + v) = y$, das heißt $y \leq x$.

Die Äquivalenz von (c) und (c') ist klar wegen 1.1.12. $\square$

Ist ein Element aus $\text{Idem}(R)$ von Null verschieden und erfüllt es eine — und damit alle — der vorstehenden äquivalenten Aussagen, so heißt es *minimal idempotent*. Diese Bezeichnung steht in Einklang mit der üblichen weiter unten in 2.1.50 gegebenen Definition von minimalen Idealen. Eine andere Bezeichnung für ein minimal idempotentes Element ist *primitiv idempotent* oder wie bei NAIMARK [221](1972), Seite 164, auch *irreduzibel idempotent*; diese beiden Bezeichnungen werden besonders dann verwendet, wenn die Eigenschaft (a) betont werden soll.

Ein zentral idempotentes $p \in R^\times$ heißt *halbprimitiv*, falls es keine zwei orthogonale zentral idempotente $u, v \in R^\times$ mit $x = u + v$ gibt. Ein zentral idempotentes $p \in R^\times$ ist genau dann halbprimitiv, wenn es das einzige zentral idempotente Element in $R^\times$ ist, welches im Ideal pR enthalten ist (mit der Symbolik von 2.1.45 also genau dann, wenn kein $q \in R'^\times$ mit $q <_r p$ existiert).

Beweis. (DIVINSKY [72](1965), Seite 26). Vorab sei an die für idempotente Elemente eines Ringes gültige Gleichung (2.7) aus 2.1.33 erinnert.

Sei p halbprimitiv. Angenommen, es existiert ein zentral idempotentes $q \in R^\times$, $q \neq p$, mit $q = pq$ (also $q <_r p$). Setze $d := p - q$. Dann ist $d^2 = (p - q)^2 = p^2 - pq - qp + q^2 = p - 2q + q = p - q = d$, also $d \in R^\times$ idempotent. Und für alle $r \in R$ ist $dr = (p - q)r = pr - qr = rp - rq = r(p - q) = rd$, also d zentral. Des Weiteren ist $qd = q(p - q) = qp - qq = pq - q = q - q = 0$ und $p = p - q + q = q + d$ im Widerspruch zur Halbprimitivität von p .

Sei nun p nicht halbprimitiv. Dann gibt es also zwei orthogonale zentral idempotente $u, v \in R^\times$ mit $p = u + v$. Dann ist $u \neq p$ und $pu = (u + v)u = u^2 + vu = u$, also $p \in pR$. $\square$

Ein idempotentes Element x von R heißt *schiefminimal idempotent* (also im Englischen etwa *skew minimal idempotent*), wenn xRx ein Schiefkörper ist. Man bemerke, dass jedes schiefminimal idempotente Element x von R notwendigerweise von 0 verschieden und wegen $x = xxx$ die Eins des Schiefkörpers xRx ist. Siehe auch 2.1.33.

Sei nun x ein schiefminimal idempotentes Element von R und y ein idempotentes Element von R . Ist $y \leq_r x$, so folgt $x \leq y$ oder $yx = 0$. Ist $y \leq_\ell x$, so folgt $x \leq y$ oder $yx = 0$. Ist $y \leq x$, so folgt $x \leq y$ oder $y = 0$.

Beweis. Sei $x \in R$ schiefminimal idempotent und $y \in R$ idempotent mit $y \leq_r x$, das heißt $xy = y$. Es gilt $xyx = yyx = yx$ und $xyx = yx$, also ist yx ein idempotentes Element in dem Schiefkörper xRx . Ist $yx \neq 0$, dann folgen aus $xyx = yx$ die Gleichungen $yx(yx)(yx)^{-1} = x(yx)(yx)^{-1}$, $yxx = xx$, $yx = x$, also $x \leq_r y$. „ $\leq_\ell$ “ geht analog. $\square$

Somit ist ein schiefminimal idempotentes Element x immer auch ein minimales Element in der Präordnung „ $\leq_r$ “ all der idempotenten Elemente y , die $yx \neq 0$ erfüllen, insbesondere also in Präordnung „ $\leq_r$ “ all der idempotenten Elemente y , die nicht orthogonal zu x sind. Für ein Beispiel, dass die Umkehrung im Allgemeinen aber nicht gilt, siehe 2.4.18.

Insbesondere folgt, dass jedes schiefminimal idempotente Element von R auch ein minimal idempotentes Element von R ist.

2.1.48 Ring-Radikale (DIVINSKY [72](1965)). Eine übliche Methode zur Klassifizierung von Ringen ist das Abwerfen (im Englischen *to discard*) eines bestimmten Teiles der Struktur, so dass nur der im gewissen Sinne sich gut verhaltende Teil zurückbleibt. Der abgeworfene Teil wird *Radikal* genannt. Je nachdem nach welchen Kriterien dabei ausgerangiert wird, erhält man verschiedene Radikale.

Sei $\mathcal{S}$ eine Eigenschaft, die ein Ring aufweisen kann. Ein Ring R , der die Eigenschaft $\mathcal{S}$ hat, wird dann als ein $\mathcal{S}$ -Ring bezeichnet. Die Klasse aller $\mathcal{S}$ -Ringe wird mit $\sigma(\mathcal{S})$ bezeichnet. Ein Ring, der kein Ideal $\mathfrak{a} \neq \{0\}$ mit $\mathfrak{a} \in \sigma(\mathcal{S})$ enthält, wird $\mathcal{S}$ -*halbeinfach* (im Englischen *$\mathcal{S}$ -semi-simple*) genannt.

Ob eine Eigenschaft $\mathcal{S}$ geeignet ist, als ein Kriterium zu dienen, was gewissermaßen Ausschuss eines Ringes sein soll, wird mit der folgenden Definition festgelegt:

Eine Eigenschaft $\mathcal{S}$ heißt eine *Radikal-Eigenschaft* (nach KUROSCHEW, 1953), wenn die folgenden drei Bedingungen erfüllt sind:

(i) (*vernünftig*). Das Bild eines Homomorphismus eines $\mathcal{S}$ -Ringes ist wieder ein $\mathcal{S}$ -Ring.

(ii) (*eindämmbar*). Jeder Ring R enthält ein Ideal $\mathfrak{X}(R) \in \sigma(\mathcal{S})$, das jedes Ideal $\mathfrak{a} \in \sigma(\mathcal{S})$ von R enthält.

(iii) (*abwerfbar*). Der Quotientenring $R/\mathfrak{X}(R)$ ist $\mathcal{S}$ -halbeinfach, das heißt, $\mathfrak{X}(R/\mathfrak{X}(R)) = \{0\}$.

Das Ideal $\mathfrak{X}(R)$ eines Ringes R heißt das $\mathcal{S}$ -*Radikal* von R und wird genauer auch mit $\mathcal{S}\text{-}\mathfrak{X}(R)$ bezeichnet; anschaulich misst es, wie weit ein Ring davon entfernt ist, halbeinfach zu sein. Ist $R = \mathfrak{X}(R)$, so heißt R ein $\mathcal{S}$ -*Radikalring*; man sagt auch, R sei $\mathcal{S}$ -*radikal*.

Ist $\mathcal{S}$ eine Radikal-Eigenschaft, so ist wegen (ii) der Ring $\{0\}$ ein $\mathcal{S}$ -Ring. Daher kann man die $\mathcal{S}$ -halbeinfachen Ringe als die Ringe auffassen, deren $\mathcal{S}$ -Radikal gleich dem Ideal $\{0\}$ ist. Jeder $\mathcal{S}$ -Ring ist gleich seinem $\mathcal{S}$ -Radikal.

Der so für Ringe erhaltene Radikalbegriff definiert sich ganz analog auch für nichtassoziative Ringe.

Um festzustellen, ob eine Eigenschaft eine Radikal-Eigenschaft ist, ist die folgende Äquivalenz nützlich: Eine Ring-Eigenschaft $\mathcal{S}$ ist genau dann eine Radikal-Eigenschaft, wenn sie einerseits die oben aufgeführte Bedingung (i) erfüllt und andererseits gilt:

(iv) Wenn jedes Bild ungleich $\{0\}$ eines Homomorphismus eines Ringes R ein Ideal $\mathfrak{a} \neq \{0\}$ mit $\mathfrak{a} \in \sigma(\mathcal{S})$ enthält, dann ist R ein $\mathcal{S}$ -Ring.

Um für eine Radikal-Eigenschaft $\mathcal{S}$ festzustellen, ob ein Ring ein $\mathcal{S}$ -Radikal ist, ist die folgende Äquivalenz nützlich: Ein Ring R ist genau dann ein $\mathcal{S}$ -Radikal, wenn R nicht homomorph auf einen von $\{0\}$ verschiedenen $\mathcal{S}$ -halbeinfachen Ring abgebildet werden kann.

Die Eigenschaft eines Ringes, nil zu sein, ist eine Radikal-Eigenschaft. Dagegen ist nilpotent keine Radikal-Eigenschaft, da im Allgemeinen (ii) nicht erfüllt ist: Ist L_ν , $\nu \in I$, die Familie aller nilpotenten Linksideale eines Ringes R , R_ν , $\nu \in J$, die Familie aller nilpotenten Rechtsideale von R und N_ν , $\nu \in K$, die Familie aller nilpotenten Ideale von R , so gilt mit $N_\ell := \sum_{\nu \in I} L_\nu$, $N_r := \sum_{\nu \in J} R_\nu$ und $N := \sum_{\nu \in K} N_\nu$ die Gleichungskette $N_\ell = N_r = N$; das Ideal N ist nun zwar nil, aber im Allgemeinen nicht nilpotent (GRAY [120](1970), Seite 28).

Ist R ein kommutativer Ring, so ist sein nil-Radikal gleich $\sqrt{\{0\}}$, dem Radikal des Ideals $\{0\}$ (siehe 2.1.43); mit anderen Worten ist dann das nil-Radikal gleich der Menge aller Wurzeln von Null. Von da rührt auch der Begriff Radikal für den zu ignorierenden Teil eines Ringes her und nicht etwa von einem aufmüpfigen (im Englischen *obstreperous*) Verhalten.

2.1.49 K -Ring-Radikale (DIVINSKY und SULIŃSKI [73](1965)). Indem man in 2.1.48 den Terminus *Ring* durch den Terminus *nichtassoziativer K -Ring* austauscht, sind für eine Eigenschaft $\mathcal{S}$, die ein nichtassoziativer Ring mit einem Operatorenbereich K haben

kann, die Begriffe *nichtassoziativer $\mathcal{S}$ - K -Ring*, $\sigma(\mathcal{S})$, *$\mathcal{S}$ -halbeinfach*, *Radikal-Eigenschaft* und *$\mathcal{S}$ -Radikal* für nichtassoziative K -Ringe definiert. Insbesondere ist entsprechend für K -Ringe der Terminus *$\mathcal{S}$ - K -Ring* definiert.

Es gilt das folgende Theorem:

Sei R ein nichtassoziativer Ring mit einem Operatorenbereich K und $\mathcal{S}$ eine Radikal-Eigenschaft des nichtassoziativen Ringes R . Dann existiert das $\mathcal{S}$ -Radikal des nichtassoziativen K -Ringes R und stimmt überein mit dem $\mathcal{S}$ -Radikal des nichtassoziativen Ringes R .

2.1.50 Ringordnungen II (DIVINSKY [72](1965)). Sei R ein Ring. Für R , aufgefasst als Linksmodul über sich selbst, sei $\leq$ die auf der Potenzmenge von R in 2.1.34 definierte partielle Ordnung (2.20). Mit anderen Worten bezeichnet $\leq$ die durch

$$\mathfrak{a} \leq \mathfrak{b} :\Leftrightarrow \mathfrak{a} \subseteq \mathfrak{b}, \quad (2.25)$$

$\mathfrak{a}, \mathfrak{b}$ Linksideale von R , definierte partielle Ordnung auf der Potenzmenge von R . Der Ring R heißt dann *linksartinsch*, wenn der R -Linksmodul R linksartinsch ist, also mit anderen Worten, wenn $(\mathcal{P}(R), \leq)$ artinsch ist. Der Begriff *rechtsartinsch* definiert sich entsprechend mit Rechtsidealen $\mathfrak{a}, \mathfrak{b}$ von R in (2.25). Und für den Begriff *artinsch* wird (2.25) mit $\mathfrak{a}, \mathfrak{b}$ als Ideale von R formuliert. Für die Begriffe *linksnoethersch*, *rechtsnoethersch* und *noethersch* verfähre man entsprechend. Für nichtassoziative Ringe verwende man entsprechend direkt nur (2.25). Für nichtassoziative Ringe mit Operatorenbereich $\mathbb{K}$ definieren sich die Begriffe ganz entsprechend. Für einen linksartinschen nichtassoziativen Ring mit Operatorenbereich sagt man auch, er erfülle die DCC auf Linksidealen. Die anderen Begriffe entsprechend. Mit diesen Definitionen ist klar, dass ein links- oder rechtsartinscher nichtassoziativer Ring artinsch ist, und dass ein links- oder rechtsnoetherscher nichtassoziativer Ring noethersch ist. Entsprechendes gilt für nichtassoziative Ringe mit Operatorenbereich K , $K = \mathbb{Z}$ oder $\mathbb{K}$. Der Ring $\left\{ \begin{pmatrix} q & 0 \\ x & y \end{pmatrix} : q \in \mathbb{Q}, x, y \in \mathbb{R} \right\}$ ist ein linksartinscher Ring, der nicht rechtsartinsch ist (GRAY [120](1970), Seite 13). Der Ring $\left\{ \begin{pmatrix} z & r \\ 0 & q \end{pmatrix} : z \in \mathbb{Z}, r, q \in \mathbb{Q} \right\}$ ist ein rechtsnoetherscher Ring, der nicht linksnoethersch ist (GOODEARL und WARFIELD [118](2004), Seite 6). In diesem Zusammenhang sei hier erwähnt, dass es Autoren gibt, die einen Ring als artinsch bezeichnen, wenn er sowohl links- als auch rechtsartinsch ist; und für noethersch entsprechend.

(McCOY [209](1964), Seite 34). Sei R ein nichtassoziativer Ring mit einem Operatorenbereich K . Ein minimales (bzw. maximales) Element in $(\mathcal{P}(R) \setminus \{\{0\}, R\}, \leq)$ heißt ein *minimales* (bzw. *maximales*) Linksideal von R . Rechtsideale und Ideale entsprechend. Ein Linksideal $\mathfrak{a} \neq \{0\}$ von R ist genau dann ein minimales Linksideal von R , wenn für alle $x \in \mathfrak{a}$ die Gleichung $(x)_\ell = \mathfrak{a}$ gilt. Ein Linksideal $\mathfrak{a} \neq R$ von R ist genau dann ein maximales Linksideal von R , wenn R für alle $x \in R \setminus \mathfrak{a}$ gleich $(\mathfrak{a} \cup \{x\})_\ell$ (also gleich $\mathfrak{a} + (x)_\ell$) ist.

Mit dem Lemma von Zorn zeigt man (GRAY [120](1970), Seite 7) (oder auch wie McCOY [209](1964), Seite 54, mit dem Maximal-Prinzip): Ist R ein Ring mit Eins und $\mathfrak{a} \neq R$ ein Ideal von R , so existiert ein maximales Ideal $\mathfrak{m}$ von R mit $\mathfrak{a} \subseteq \mathfrak{m}$.

(DIVINSKY [72](1965), Seiten 22 und 51-54). Ist R ein linksartinscher Ring, so ist jedes Linksideal, das nil ist, auch nilpotent; R enthält dann ein idempotentes Element $p \neq 0$. Folglich ist das nil-Radikal — das ja als ein Ideal insbesondere ein Linksideal ist — eines linksartinschen Ringes nilpotent; genauer ist in einem linksartinschen Ring R das nil-Radikal gleich der Summe aller nilpotenten Ideale von R (mit der Bezeichnung von 2.1.48 also gleich N und damit, siehe ebd., auch gleich N_ℓ und gleich N_r). Ist R ein linksnoetherscher Ring, so ist sowohl jedes Linksideal, das nil ist, als auch jedes Rechtsideal, das nil ist, auch nilpotent.

(DIVINSKY [72](1965), Seite 135). Ist $\mathfrak{m}$ ein minimales Linksideal eines Ringes R , dann ist $\mathfrak{m}$ entweder ein nicht trivialer, einfacher Ring oder es ist $\mathfrak{m}^2 = \{0\}$.

2.1.51. Sei R ein Ring mit Operatorenbereich $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$. Sei $p \in R^\times$ ein idempotentes Element. Ist dann Rp ein minimales Linksideal von R , so ist p schiefminimal idempotent.

Beweis. (PALMER [231](1994), Seite 668, Satz 8.2.2). Liege der Fall vor, dass $p \in \text{Idem}(R)$ mit $\mathfrak{m} := Rp$ ein minimales Linksideal von R . Dann ist $p = pp$ ein Element von $\mathfrak{m}$. Betrachte den Unterring pRp . Sei $x \in (pRp)^\times$, etwa $x = pap$ für ein $a \in R^\times$. Dann ist offensichtlich $xp = px = x$, das heißt, p ist eine Eins des Unterringes pRp . Des Weiteren gilt $\{ppxp\} \subseteq Rpxp \subseteq R\mathfrak{m} \subseteq \mathfrak{m}$, wegen der Minimalität also $\mathfrak{m} = Rpxp$. Es gibt also ein $b \in R$ mit $p = bpxp$, also wegen $p = pp$ mit $p = (pbp)(pxp)$. Somit hat also jedes Element $x \in (pRp)^\times$ eine Linksinverse und damit erfüllt $(pRp)^\times$ die Bedingungen einer multiplikativen Gruppe, das heißt, pRp ist ein Schiefkörper und p ist schiefminimal idempotent. $\square$

Da jedes schiefminimal idempotente Element auch minimal idempotent ist, gilt folglich auch: Ist Rp für ein idempotentes $p \in R^\times$ ein minimales Linksideal von R , so ist p minimal idempotent. Diese Aussage kann auch direkt gezeigt werden:

Beweis. Sei p nicht primitiv; es existieren also zwei orthogonale Idempotente $u, v \in R^\times$ mit $p = u + v$. Wegen $up = u(u + v) = u$ ist dann $Au = Aup \subseteq Ap$. Wegen $p = p^2$ ist $p \in Ap$. Es ist aber $p \notin Au$, denn: Wäre $p = xu$ für ein $x \in A$, dann $p v = x u v = 0$, aber $p v = (u + v)v = v$ und man hätte $v = 0$, Widerspruch. Also $\{0\} \subsetneq Au \subsetneq Ap$, das heißt, Ap ist kein minimales Linksideal. $\square$

2.1.52 (PALMER [231](1994), Seite 668, Satz 8.2.2). Sei $\mathfrak{m}$ ein minimales Linksideal eines Ringes R mit Operatorenbereich $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$. Ist $\mathfrak{m}^2 \neq \{0\}$ (also zum Beispiel, wenn R ein halbprimärer K -Ring ist), so existiert ein Idempotent $p \in R$ mit $\mathfrak{m} = Rp$ (womit also insbesondere $p \in \mathfrak{m}$ ist); jedes solche p ist schiefminimal idempotent. Ist umgekehrt der K -Ring R halbprim und $p \in R$ schiefminimal idempotent in R , so ist Rp ein minimales Linksideal (das offensichtlich p enthält). Entsprechendes gilt für Rechtsideale.

Beweis. Sei $\mathfrak{m}$ ein minimales Linksideal mit $\mathfrak{m}^2 \neq \{0\}$. Es existiert also ein $a \in \mathfrak{m}^\times$ mit $\mathfrak{m}a \neq \{0\}$. $\mathfrak{m}a$ ist ein Linksideal mit $\mathfrak{m}a \subseteq \mathfrak{m}$. Wegen der Minimalität von $\mathfrak{m}$ gilt $\mathfrak{m}a = \mathfrak{m}$. Da $a \in \mathfrak{m}^\times$, existiert also ein $p \in \mathfrak{m}^\times$ mit $pa = a$. $\mathfrak{m}(p - 1)$ ist ein Linksideal, das von $\mathfrak{m}$ umfasst wird. Wegen der Minimalität ist entweder $\mathfrak{m}(p - 1) = \mathfrak{m}$ oder $\mathfrak{m}(p - 1) = \{0\}$. Die erste Gleichung ist aber unmöglich, denn sie würde die Existenz eines $b \in \mathfrak{m}$ implizieren mit $b(p - 1) = p$, und diese Gleichung wiederum würde den Widerspruch $a = pa = b(p - 1)a = b(pa - a) = bpa - ba = ba - ba = 0$ implizieren. Also gilt $\mathfrak{m}(p - 1) = \{0\}$. Somit auch $p(p - 1) = 0$, $p^2 - p = 0$ und p ist idempotent. Da $p \in \mathfrak{m}$ und $\mathfrak{m}$ ein Linksideal ist, gilt $Rp \subseteq \mathfrak{m}$. Insbesondere ist $p = pp \in Rp$ und wegen der Minimalität also $\mathfrak{m} = Rp$.

Dass jedes solche p schiefminimal idempotent ist, wurde in 2.1.51 gezeigt.

Sei nun umgekehrt der K -Ring R halbprim und $p \in R$ schiefminimal idempotent. Sei $\mathfrak{a} \subseteq Rp$ ein Linksideal mit $\mathfrak{a} \neq \{0\}$. Nach Voraussetzung und wegen der Bemerkung in 2.1.44 ist $\mathfrak{a}^2 \neq \{0\}$. Also ist wegen $\mathfrak{a}^2 \subseteq Rpa$ auch $pa \neq \{0\}$. Es existiert also ein $x \in \mathfrak{a}$, etwa $x = ap$ für ein $a \in R$, mit $px \neq 0$, also $pap \neq 0$. Da pRp ein Schiefkörper ist mit p als Eins, existiert ein $b \in R$ mit $pbpap = p$; also $p = (pbp)x$. Da $\mathfrak{a}$ ein Linksideal ist, folgt $p \in \mathfrak{a}$ und $Rp \subseteq \mathfrak{a}$. Und es gilt $\mathfrak{a} = Rp$. $\square$

In halbprimen Ringen R mit Operatorenbereich $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$ ist also ein Idempotent $p \in R^\times$ genau dann schiefminimal idempotent, wenn Rp ein minimales Linksideal von R

ist. Somit ist dann die zwischen den in 2.1.45 durch $\equiv_\ell$ definierten Äquivalenzklassen aus schiefminimal idempotenten Elementen und den minimalen Linksidealen erklärte Abbildung $[p]_\ell \mapsto Rp$ bijektiv.

2.1.53 Nil-halbeinfacher linksartinscher Ring (DIVINSKY [72](1965)). Sei R ein nil-halbeinfacher linksartinscher Ring.

(ebd., Seite 23). Für jedes Linksideal $\mathfrak{a} \neq \{0\}$ von R existiert ein $p \in \text{Idem}(R)^\times$ mit $\mathfrak{a} = Rp$. Für jedes beidseitige Ideal $\mathfrak{a} \neq \{0\}$ von R existiert genau ein $p \in \text{Idem}(R)^\times$ mit $\mathfrak{a} = Rp = pR$; dieses p ist eine Eins von $\mathfrak{a}$; insbesondere hat also R selbst eine Eins. Ein $p \in \text{Idem}(R)^\times$ ist genau dann zentral in R , wenn es die Eins eines Ideals von R ist. Ist $\mathfrak{a}$ ein Ideal von R , so ist jedes Linksideal (bzw. Rechtsideal, bzw. Ideal) von $\mathfrak{a}$ ein Linksideal (bzw. Rechtsideal, bzw. Ideal) von R . Folglich erfüllt jedes Ideal von R auch die DCC auf Linksideale, und jedes Ideal von R ist ebenfalls nil-halbeinfach.

(ebd., Seite 27). Ein zentral idempotentes $p \in R^\times$ ist genau dann halbprimitiv, wenn das Ideal Rp einfach ist. Jedes zentral idempotente, nicht halbprimitive $p \in R^\times$ ist eine Summe von endlich vielen paarweise orthogonalen halbprimitiven Elementen aus $\text{Idem}(R)^\times$. Es gilt das folgende *Strukturtheorem*: Im Fall von $R \neq \{0\}$ hat R nur eine endliche Anzahl von einfachen Idealen; diese Ideale sind linksartinsche, nicht nilpotente Ringe und R ist die innere direkte Summe dieser Ideale.

(ebd., Seiten 29-32). Ein $p \in \text{Idem}(R)^\times$ ist genau dann minimal idempotent, wenn Rp ein minimales Linksideal von R ist und dies ist genau dann der Fall, wenn p schiefminimal idempotent ist. Jedes idempotente, nicht minimal idempotente $p \in R^\times$ kann als eine endliche Summe von paarweise orthogonalen minimal idempotenten Elementen aus $R^\times$ ausgedrückt werden.

(ebd., Seite 37). Sei R ein nil-halbeinfacher Ring. Dann ist R genau dann linksartinsch, wenn er rechtsartinsch ist. Dies gilt im Allgemeinen nicht für nicht nil-halbeinfache Ringe.

(ebd., Seite 39). Sei R ein linksartinscher Ring und N sein nil-Radikal. Dann ist R/N ein nil-halbeinfacher linksartinscher Ring.

2.1.54 (DIVINSKY [72](1965), Seite 32, 55-56). Mittels des sogenannten Jordan-Hölder-Theorems aus der Gruppentheorie zeigt man das folgende *Strukturtheorem*: Ein Ring R ist genau dann einfach, linksartinsch und nicht nilpotent, wenn es ein $n \in \mathbb{N}^\times$ und einen Schiefkörper D derart gibt, dass der Ring R isomorph ist zu der Menge aller $n \times n$ -Matrizen mit Elementen aus D .

Das Analogon dieses Strukturtheorems für linksnoethersche Ringe gilt nicht: Es gibt einen einfachen, linksnoetherschen, nicht nilpotenten Ring, der kein Matrizenring über einem Schiefkörper ist. Bezüglich gültiger Struktursätze über linksnoethersche, (halb)Primringe siehe DIVINSKY [72](1965), Seite 88.

Man bemerke, dass jeder einfache, linksartinsche, nicht nilpotente Ring ein nil-halbeinfacher linksartinscher Ring ist.

Allgemeiner gilt das *Wedderburn-Artin-Theorem*:

Ist R ein linksartinscher Ring mit $R \neq \{0\}$, so ist R genau dann nil-halbeinfach, wenn R isomorph zu einer inneren direkten Summe von endlich vielen Matrizenringen über Schiefkörper ist.

Beweis. Das eben genannte Strukturtheorem ist ein Baustein für beide Beweisrichtungen. Zusammen mit dem Strukturtheorem in 2.1.53 liefert es den „genau-dann“-Teil. Und zusammen mit der Tatsache, dass die innere Summe von endlich vielen einfachen, linksartinschen Ringen nil-halbeinfach ist, siehe GRAY [120](1970), Seite 40, Theorem 17, folgt mit ihm der „wenn“-Teil. $\square$

2.1.55 (DIVINSKY [72](1965), Seite 47). Sei R ein Ring und $\leq$ die in 2.1.50 definierte partielle Ordnung (2.25) auf $\mathcal{P}(R)$. Dann ist R genau dann linksnoethersch, wenn jedes Linksideal von R endlich erzeugt ist.

Beweis. Sei R linksnoethersch und $\mathfrak{a}$ ein Linksideal von R . Sei $x_1 \in \mathfrak{a}$. Dann ist $(x_1)_\ell \subseteq \mathfrak{a}$ und falls diese Inklusion echt ist, gibt es ein $x_2 \in \mathfrak{a} \setminus (x_1)_\ell$. Dann ist $(x_1, x_2)_\ell \subseteq \mathfrak{a}$. Gilt hierbei wieder keine Gleichheit, so existiert ein $x_3 \in \mathfrak{a} \setminus (x_1, x_2)_\ell$. So fortfahrend erhält man eine aufsteigende Kette $(x_1)_\ell \subseteq (x_1, x_2)_\ell \subseteq (x_1, x_2, x_3)_\ell \subseteq \dots$, die nach Voraussetzung endlich sein muss, das heißt, $\mathfrak{a} = (x_1, \dots, x_n)_\ell$.

Sei nun R ein Ring, in dem jedes Linksideal endlich erzeugt ist. Sei $K = (\mathfrak{a}_\nu)_{\nu \in I}$ eine Kette in $\mathcal{P}(R)$. Setze $\mathfrak{b} := \bigcup_{\nu \in I} \mathfrak{a}_\nu$. Dann ist $\mathfrak{b}$ ein Linksideal von R , also nach Voraussetzung endlich erzeugt, etwa von $x_1, \dots, x_n$. Also $x_1, \dots, x_n \in \mathfrak{b}$ und daher $x_i \in \mathfrak{a}_{\nu_i}$ für gewisse $\nu_i \in I$, $i = 1, \dots, n$. Setze $m := \max\{\nu_1, \dots, \nu_n\}$. Dann also $x_i \in \mathfrak{a}_m$ für alle $i \in \{1, \dots, n\}$ und wegen $\mathfrak{b} = (x_1, \dots, x_n)$ also $\mathfrak{b} \subseteq \mathfrak{a}_m$. Folglich $\mathfrak{a}_m = \mathfrak{b}$, $\mathfrak{a}_m = \mathfrak{a}_{m+1} = \dots$ und R ist linksnoethersch. $\square$

2.1.56 (DIVINSKY [72](1965), Seiten 47, 48; GRAY [120](1970), Seiten 13, 31). Im Allgemeinen sind die beiden Eigenschaften artinsch und noethersch voneinander unabhängig. Es gibt artinsche Ringe, die nicht noethersch sind. Wieder mittels des Jordan-Hölder-Theorems gilt aber zumindest das Theorem, dass jeder linksartinsche Ring mit Linkseins linksnoethersch ist. Insbesondere ist somit jeder nil-halbeinfache linksartinsche Ring linksnoethersch. Jeder Hauptidealring ist noethersch, aber zum Beispiel ist $\mathbb{Z}$ nicht artinsch. Für ein Beispiel eines nicht noetherschen Ringes siehe GRAY, ebd., Seite 13.

2.1.57 (VAN DER WAERDEN [284](1967), Seite 122). Sei R ein kommutativer Ring mit Eins. Ist R noethersch, so ist auch der Ring $R[x_1, \dots, x_n]$ über endlich vielen unabhängigen Unbestimmten $x_1, \dots, x_n$ noethersch. Insbesondere sind somit die Ringe $\mathbb{Z}[x_1, \dots, x_n]$ und $K[x_1, \dots, x_n]$, K ein Körper, noethersch. Bezüglich des Falles, dass R ein Körper ist, siehe auch COX, LITTLE und O'SHEA [54](1997), Seite 74, Theorem 2.5.4.

2.1.58 (DIVINSKY [72](1965), Seiten 40, 44, 61, 88). Ist R ein Ring, so ist der Durchschnitt $\mathfrak{P}$ aller Primideale von R ein Radikal. Ist dieses Radikal gleich dem ganzen Ring, so ist dieser Ring ein nil-Radikalring. Ansonsten stimmt dieses Radikal in linksartinschen Ringen mit dem nil-Radikal überein. In Untersuchungen in linksnoetherschen Ringen nimmt dieses Radikal die Rolle des bei linksartinschen Ringen betrachteten nil-Radikals ein.

Ist R ein linksnoetherscher Ring und $\mathfrak{P}$ das eben erwähnte Radikal, so ist $R/\mathfrak{P}$ ein Unterring eines Ringes, der isomorph zu einer endlichen Summe von Matrizenringen über Schiefkörpern ist. Ist R ein linksnoetherscher Primring, so ist R ein Unterring eines Ringes, der für ein $n \in \mathbb{N}^\times$ isomorph ist zu dem $n \times n$ -Matrizenring über einen Schiefkörper.

2.1.59 Jacobson-Radikal (BEHRENS [23](1975), Seite 74; DIVINSKY [72] (1965), Seiten 40, 44, 93-95, 107, 108, 120; JACOBSON [150](1956), Seiten 5, 7, 9, 10, 38; MCCOY [209](1964), Seiten 112, 115).

Sei R ein Ring und $x \in R$. Man betrachte das Rechtsideal $(1 - x)R$. Ein Rechtsideal $\mathfrak{a}$ von R wird *modular* genannt, falls ein $y \in R$ mit $(1 - y)R \subseteq \mathfrak{a}$ existiert; solch ein y heißt dann eine *Linkseins modulo $\mathfrak{a}$* . Diese Bezeichnung ist klar, wenn man bemerkt, dass ein Rechtsideal $\mathfrak{a}$ von R genau dann modular mit x als Linkseins modulo $\mathfrak{a}$ ist, wenn für alle $r \in R$ die Kongruenz $r \equiv xr \pmod{\mathfrak{a}}$ gilt, also genau dann, wenn $x + \mathfrak{a}$ eine Linkseins in $R/\mathfrak{a}$ ist. Somit ist ein Ideal $\mathfrak{a}$ genau dann modular, wenn der Restklassenring $R/\mathfrak{a}$ eine Eins hat. Hat R eine Rechtseins, so ist jedes Rechtsideal modular. Offenbar ist $(1 - x)R$ das kleinste modulare Rechtsideal mit x als Linkseins modulo $(1 - x)R$. Mit anderen Worten enthält jedes modulare Rechtsideal $\mathfrak{a}$ mit x als Linkseins modulo $\mathfrak{a}$ das Rechtsideal $(1 - x)R$.

Nun kann es ja auch vorkommen, dass $(1 - x)R = R$ ist. In diesem Fall ist dann insbesondere $-x \in (1 - x)R$, also $-x = (1 - x)y$ für ein $y \in R$, also $-x = y - xy$,

$x + y - xy = 0, x \circ y = 0$, wobei $\circ$ das in 2.1.9 definierte Quasi-Produkt ist. Ist umgekehrt x rechts-quasi-invertierbar, gilt also $x \circ y = 0$ für ein $y \in R$, so folgt mit den gleichen Gleichungen wieder $-x \in (1 - x)R$, also $x \in (1 - x)R$. Sei nun $r \in R$ beliebig. Dann ist $xr = (1 - x)(-y)r$ auch in $(1 - x)R$ enthalten. Es ist $r - xr = (1 - x)r \in (1 - x)R$, also $r = (r - xr) + xr \in (1 - x)R$. Somit gilt also: x ist genau dann rechts-quasi-invertierbar, wenn es ein Element von $(1 - x)R$ ist, und dies wiederum ist genau dann der Fall, wenn $(1 - x)R = R$ ist. Analog ist ein $x \in R$ genau dann links-quasi-invertierbar, wenn es ein Element von $R(1 - x)$ ist, und dies wiederum nur genau dann, wenn $R = R(1 - x)$ gilt. (Möchte man das Quasi-Produkt $\bullet$ verwenden, so braucht man nur alle Aussagen auf dem Rechtsideal $(1 + x)R$ anstelle des Rechtsideals $(1 - x)R$ aufbauend formulieren.) Insbesondere folgt, falls R eine Eins hat, dass ein $r \in R$ genau dann rechts-invertierbar ist, wenn $rR = R$ gilt, also genau dann, wenn r in keinem echten Rechtsideal von R enthalten ist.

Sei $x \in R$ ein nicht rechts-quasi-invertierbares Element von R . Dann existiert ein maximales Rechtsideal $\mathfrak{a}$ von R , das modular ist und die Eigenschaft hat, dass x eine Linkseins modulo $\mathfrak{a}$ mit $x \notin \mathfrak{a}$ ist.

Beweis. Nach Voraussetzung ist $x \notin (1 - x)R$. Würde ein Rechtsideal $\mathfrak{a}$ von R sowohl $(1 - x)R$ als auch $\{x\}$ umfassen, so wäre $(1 - x)R + xR$ eine Teilmenge von $\mathfrak{a}$ und daher jedes $r \in R$ zumindest als $r = r - xr + xr$ auch in $\mathfrak{a}$ enthalten, das heißt, $\mathfrak{a}$ wäre ganz R .

Wenn man alle Rechtsideale betrachtet, die zwar $(1 - x)R$ umfassen, aber nicht das Element x enthalten, und diese per Mengeninklusion wie in (2.25) partiell ordnet, so erhält man mit dem Lemma von Zorn ein Rechtsideal, welches bezüglich sowohl der Ausschließung von x als auch der Umfassung von $(1 - x)R$ maximal ist. Dieses Rechtsideal ist ein maximales Rechtsideal von R ; denn, gäbe es ein Rechtsideal $\mathfrak{b} \neq R$ mit $\mathfrak{a} \subsetneq \mathfrak{b}$, so müsste $x \in \mathfrak{b}$ gelten, dann wäre aber, wie schon bemerkt, $\mathfrak{b} = R$, Widerspruch. $\square$

Sei $\mathfrak{a}$ ein Linksideal oder ein Rechtsideal von R . Dann nennt man $\mathfrak{a}$ *links-quasi-regulär* (im Englischen *left quasi-regular*), falls jedes Element von $\mathfrak{a}$ links-quasi-invertierbar ist. Entsprechend nennt man $\mathfrak{a}$ *rechts-quasi-regulär* (im Englischen *right quasi-regular*), falls jedes Element von $\mathfrak{a}$ rechts-quasi-invertierbar ist. Und man nennt $\mathfrak{a}$ *quasi-regulär* (im Englischen *quasi-regular*), falls jedes Element von $\mathfrak{a}$ quasi-invertierbar ist. Jedes rechts-quasi-reguläre Rechtsideal $\mathfrak{a}$ von R ist eine Untergruppe von $G^q(R)$, insbesondere also quasi-regulär.

Sei $r \in R$ ein nilpotentes Element, also $r^n = 0$ für ein $n \in \mathbb{N}^\times$. Setze $s := -r - r^2 - \dots - r^{n-1}$. Dann ist s wegen $r \circ s = s \circ r = r^n$ die Quasi-Inverse von r . Somit ist jedes nile Rechtsideal und jedes nile Linksideal quasi-regulär. Die Eigenschaft rechts-quasi-regulär ist eine Radikaleigenschaft. Das zugehörige Radikal ist die Vereinigung aller rechts-quasi-regulären Rechtsideale von R ; es ist ein quasi-reguläres, beidseitiges Ideal und heißt das *Jacobson-Radikal*, in Zeichen $\mathfrak{J}(R)$ oder einfach nur $\mathfrak{J}$.

Wie das folgende Beispiel zeigt, ist nicht unbedingt jedes rechts-quasi-invertierbare Element eines Ringes auch ein Element des Jacobson-Radikals. Betrachte dazu den Zahlenring $\mathbb{Z}$. Nach Bemerkung 2.1.20, Gleichung (2.1), ist ein $x \in \mathbb{Z}$ genau dann rechts-quasi-invertierbar, wenn ein $y \in \mathbb{Z}$ mit $(1 - x)(1 - y) = 1$ existiert. Das heißt, in $\mathbb{Z}$ sind nur die beiden Zahlen 0 und 2 rechts-quasi-invertierbar. Insbesondere ist das aus allen geraden Zahlen bestehende Hauptideal (2) nicht rechts-quasi-regulär. Somit ist $\{0\}$ das Jacobson-Radikal von $\mathbb{Z}$, welches offensichtlich nicht die rechts-quasi-invertierbare 2 enthält.

$\mathfrak{J}(R)$ ist gleich der Menge aller $x \in R$, für die xR rechts-quasi-regulär ist. Und analog ist $\mathfrak{J}(R)$ auch gleich der Menge aller $x \in R$, für die Rx links-quasi-regulär ist.

Ist $R \neq \mathfrak{J}(R)$, so existiert ein $x \in R$, welches nicht rechts-quasi-invertierbar in R ist. Oben wurde bewiesen, dass dann ein modulares, maximales Rechtsideal existiert. Für

jeden Ring R gilt: $\mathfrak{J}(R)$ ist gleich dem Durchschnitt von allen modularen, maximalen Rechtsidealen von R ; dabei sei an 1.1.2 erinnert, wonach hier der Durchschnitt über eine leere Indexmenge gleich ganz R ist. Diese Aussage gilt auch analog für Linksideale. Falls also zum Beispiel R ein kommutativer Ring mit einer Eins ist, so ist $\mathfrak{J}(R)$ gleich dem Durchschnitt von allen maximalen Idealen.

$\mathfrak{J}(R)$ umfasst alle nilen Linksideale und alle nilen Rechtsideale. Grundsätzlich umfasst das Jacobson-Radikal das nil-Radikal. Bei artinschen Ringen stimmt das nil-Radikal mit dem Jacobson-Radikal überein. In linksartinschen Ringen ist das Jacobson-Radikal nilpotent und folglich ist in linksartinschen Ringen auch jedes nile Links- und nile Rechtsideal nilpotent.

Ist $\mathfrak{a}$ ein Ideal von R , so gilt, für $\mathfrak{a}$ aufgefasst als ein Ring, die Gleichung $\mathfrak{J}(\mathfrak{a}) = \mathfrak{a} \cap \mathfrak{J}(R)$. Aber im Allgemeinen muss nicht jeder Unterring eines Jacobson-Radikal-Ringes wieder ein Jacobson-Radikal-Ring sein.

Des Weiteren gilt $\mathfrak{J}(R) = \mathfrak{J}(R^1)$.

Es sei angemerkt, dass jeder einfache, unitale Ring mit Operatorenbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$, Jacobson-halbeinfach ist.

2.1.60 Sockel (PALMER [231](1994), Seite 671). Sei R ein Ring mit Operatorenbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$. Der *Linkssockel* ist die Summe aller minimalen Linksideale von R . *Rechtssockel* entsprechend. Stimmen der Links- und Rechtssockel von R überein, so heißt ihr gemeinsamer Wert der *Sockel* (im Englischen *socle*) von A , in Zeichen $\text{soc}(R)$. Man bemerke, dass jede solcher Summen — für R aufgefasst als ein Ring — eine innere direkte Summe ist. Des Weiteren beachte man, dass im Fall einer leeren Indexmenge die Summe gleich $\{0\}$ ist.

(a) Ist $\mathfrak{m}$ ein minimales Linksideal und $a \in R$, so ist $\mathfrak{m}a$ entweder $\{0\}$ oder ein minimales Linksideal. Minimales Rechtsideal entsprechend. Folglich ist sowohl der Links- als auch der Rechtssockel ein beidseitiges Ideal.

Beweis. Das der Linkssockel ein Linksideal ist, ist klar. Sei $\mathfrak{m}$ ein minimales Linksideal und $a \in R$. Dann ist $\mathfrak{m}a$ ein Linksideal. Falls $\mathfrak{a}$ ein Linksideal mit $\mathfrak{a} \subseteq \mathfrak{m}a$ ist, so ist $\{m \in \mathfrak{m} : ma \in \mathfrak{a}\}$ ein Linksideal, das von $\mathfrak{m}$ umfasst wird und somit gleich $\{0\}$ oder $\mathfrak{m}$ ist; das heißt, $\mathfrak{a}$ ist gleich $\{0\}$ oder gleich $\mathfrak{m}a$, mit anderen Worten ist $\mathfrak{m}a$ entweder $\{0\}$ oder ein minimales Linksideal. Somit ist der Linkssockel abgeschlossen unter der Multiplikation von rechts und er ist ein Ideal. $\square$

(b) Falls der K -Ring R halbprim ist, so stimmen Links- und Rechtssockel überein, das heißt, R hat einen Sockel. Dieser Sockel ist gleich dem Linksideal, Rechtsideal und Ideal, welches von der Menge der schiefminimal idempotenten Elemente von R erzeugt wird.

Beweis. Ist der K -Ring R halbprim, so zeigt 2.1.52, dass der Linkssockel gleich dem Linksideal ist, das von der Menge aller schiefminimal Idempotente erzeugt wird. Entsprechend ist der Rechtssockel das Rechtsideal, das von dieser Menge erzeugt wird. Nach (a) ist also sowohl der Linkssockel als auch der Rechtssockel das Ideal, das von der Menge der schiefminimal idempotenten Elementen erzeugt wird. Somit stimmen sie überein und R hat einen Sockel. $\square$

(c) Bezeichne $\mathfrak{J}(R)$ das Jacobson-Radikal von R . Dann gilt

$$\text{Idem}(R) \cap \mathfrak{J}(R) = \{0\}.$$

Beweis. Sei $p \in \mathfrak{J}(R)$ idempotent. Dann ist p rechts-quasi-invertierbar und es gilt $p = p + (p^q - pp^q) - (p^q - pp^q) = p \circ p^q + pp^q - p^q = pp^q - pp^q = 0$. $\square$

(d) Falls der K -Ring R halbprim ist, so gilt

$$\text{soc}(R) \cap \mathfrak{J}(R) = \{0\}.$$

Beweis. Angenommen, $a \in \text{soc}(R) \cap \mathfrak{J}(R)$ mit $a \neq 0$. Bezeichne M die Menge aller schiefminimal idempotenten Elemente von R . $\text{soc}(R)$ ist ein Ideal, also ist mit $a \in \text{soc}(R)$ für jedes $p \in M$ auch $ap \in \text{soc}(R)$, genauer ist $ap \in Rp$. Da nach (b) $\text{soc}(R) = \sum_{p \in M} \oplus Rp$ gilt, ist somit $a = \sum_{p \in M} ap$, wobei nur für endlich viele $p \in M$ der Ausdruck ap ungleich Null ist. Da $a \neq 0$ ist, muss demnach für ein $p \in M$ die Ungleichung $ap \neq 0$ gelten. Da $\text{soc}(R) \cap \mathfrak{J}(R)$ ein Ideal ist, enthält es dieses von Null verschiedene ap . Insbesondere ist $\text{soc}(R) \cap \mathfrak{J}(R)$ ein Linksideal, das ap enthält. Nun ist Rp ein minimales Linksideal, das ap enthält. Somit ist $Rp \cap (\text{soc}(R) \cap \mathfrak{J}(R)) \subseteq Rp$ auch ein Linksideal, das ap enthält und wegen der Minimalität gilt $Rp = Rp \cap (\text{soc}(R) \cap \mathfrak{J}(R))$. Das heißt insbesondere $Rp \subseteq Rp \cap (\text{soc}(R) \cap \mathfrak{J}(R))$ und damit ist wegen $pp = p$ auch $p \in \mathfrak{J}(R)$, im Widerspruch zu (c). $\square$

2.2 Vektorräume

2.2.1 Definition. Ein unitaler Linksmodul (bzw. Rechtsmodul) X über einen Schiefkörper $\mathbb{D}$ heißt ein *Linksvektorraum* (bzw. *Rechtsvektorraum*) über $\mathbb{D}$; die Elemente von X heißen dann *Vektoren* und die Untermoduln von X *Untervektorräume* oder einfach *Unterräume*. Falls nicht ausdrücklich etwas anderes vereinbart ist, meint ein *Vektorraum* über einen Schiefkörper $\mathbb{D}$ stets einen Linksvektorraum über $\mathbb{D}$.

2.2.2 Definition. Sei X ein Vektorraum über $\mathbb{K}$. Eine Abbildung $p: X \rightarrow \mathbb{R}$ heißt eine *Halbnorm*, wenn die beiden folgenden Aussagen gelten:

- (a) $p(\lambda x) = |\lambda|p(x)$ für alle $\lambda \in \mathbb{K}$, $x \in X$, (absolute Homogenität)
- (b) $p(x + y) \leq p(x) + p(y)$ für alle $x, y \in X$. (Subadditivität)

Gilt zusätzlich

- (c) $p(x) = 0 \Rightarrow x = 0$ für alle $x \in X$, (Treue)

so heißt p eine *Norm*.

2.2.3 Bemerkung. Sei X ein Vektorraum über $\mathbb{K}$ und $p: X \rightarrow \mathbb{R}$ eine Halbnorm. Dann gilt $p(0) = 0$ und für alle $x \in X$ ist $p(x) \geq 0$ (denn: $0 = p(x - x) \leq p(x) + p(-x) = p(x) + p(x) = 2p(x)$); des Weiteren gilt für alle $x, y \in X$ die Ungleichung $|p(x) - p(y)| \leq p(x - y)$ (denn: $p(x) = p(x - y + y) \leq p(x - y) + p(y)$, also $p(x) - p(y) \leq p(x - y)$. Analog auch $p(y) - p(x) \leq p(x - y)$).

Eine Halbnorm kann man daher auch direkt als eine Abbildung $X \rightarrow \mathbb{R}^+$ definieren oder betrachten.

2.2.4. Sei X ein Vektorraum über $\mathbb{K}$ und p eine Halbnorm auf X . Seien $x, y \in X$ mit $p(x + y) = p(x) + p(y)$. Dann gilt für alle $\lambda, \mu \in \mathbb{R}^+$ die Gleichung $p(\lambda x + \mu y) = \lambda p(x) + \mu p(y)$.

Beweis. Seien $\lambda, \mu \in \mathbb{R}^+$, wobei ohne Einschränkung $\lambda \geq \mu$ angenommen werden kann. Dann gilt: $\lambda p(x) + \mu p(y) \geq p(\lambda x + \mu y) = p(\lambda(x + y) - (\lambda - \mu)y) \geq \lambda p(x + y) - (\lambda - \mu)p(y) = \lambda(p(x) + p(y)) - (\lambda - \mu)p(y) = \lambda p(x) + \mu p(y)$. $\square$

2.2.5 Bezeichnungen. Seien X und Y zwei normierte Vektorräume über $\mathbb{K}$ und sei $\|\cdot\|$ die Norm von X . Mit B_X beziehungsweise S_X wird durchgehend die Menge $\{x \in X : \|x\| \leq 1\}$ (*abgeschlossene Einheitskugel*) beziehungsweise $\{x \in X : \|x\| = 1\}$ (*Einheitssphäre*) bezeichnet. B_X wird stets auch als eine punktierte Menge mit Basispunkt 0 aufgefasst. Die Menge aller $\frac{1}{n} \cdot B_X$, $n \in \mathbb{N}^\times$, ist eine Nullumgebungsbasis der sogenannten *Norm-Topologie* von X ; wie üblich, ist jeder normierte Vektorraum über $\mathbb{K}$, falls nicht ausdrücklich etwas anderes vereinbart ist, stets mit dieser Topologie versehen. Als Symbol für die Norm-Topologie wird das Symbol der Norm oder der Buchstabe n verwendet. Jeder normierte Vektorraum X über $\mathbb{K}$ wird auch stets, falls nicht etwas anderes vereinbart ist, per der Metrik $(x, y) \mapsto \|x - y\|$ als ein metrischer Raum aufgefasst.

Sei T eine lineare Abbildung von X nach Y . T heißt ein *Homomorphismus*, wenn T stetig und offen ist. Dementsprechend ist T ein *Isomorphismus*, wenn T auch ein Homöomorphismus ist und dann heißen X und Y isomorph, in Zeichen $X \simeq Y$. Eine Abbildung von X nach Y heißt eine *Isometrie* oder auch *isometrisch*, falls $\|Tx\| = \|x\|$ für alle $x \in X$ gilt. Existiert ein isometrischer Isomorphismus von X auf Y , so heißen X und Y isometrisch isomorph, in Zeichen $X \cong Y$.

2.2.6 Definition. (BOURBAKI [35](1987), II.2.5; SCHAEFER [255](1966); MEGGINSON [211](1998)). Sei X ein Vektorraum über $\mathbb{R}$. Zusammen mit einer Präordnung $\leq$ auf X heißt X , oder genauer $(X, \leq)$, ein *prägeordneter Vektorraum*, wenn die beiden folgenden Verträglichkeitsbedingungen erfüllt sind:

- (a) $x, y, z \in X, x \leq y \Rightarrow x + z \leq y + z$,
- (b) $x, y \in X, x \leq y, \lambda \in \mathbb{R}^{+\times} \Rightarrow \lambda x \leq \lambda y$.

(JAMESON [153](1970)). Für einen prägeordneten Vektorraum X und $a, b \in X$ mit $a \leq b$ heißt die Menge $\{x \in X : a \leq x \leq b\}$ ein *Ordnungsintervall*, in Zeichen $[a, b]$; siehe auch 2.2.43.

Ist die Ordnungsrelation eines prägeordneten Vektorraumes X antisymmetrisch — also eine partielle Ordnung —, dann heißt X ein *geordneter Vektorraum*. Ein geordneter Vektorraum, der zugleich ein Verband ist, heißt ein *Vektorverband* oder auch ein *Riesz'scher Raum*.

2.2.7. FREMLIN [99](2004), 352B, zeigt: Ein geordneter Vektorraum X ist genau dann auch ein Verband, also ein Vektorverband, wenn für alle $x \in X$ $\sup\{0, x\}$ definiert ist.

2.2.8 Definition mit Bemerkungen (BOURBAKI [35](1987); KÖTHE [187] (1966)). Sei X ein Vektorraum über $\mathbb{K}$ und seien A und B Teilmengen von X .

Bezüglich der im Folgenden verwendeten Ausdrücke der Gestalt $K \cdot A$ für $K \subseteq \mathbb{K}$ sei an die Definition am Ende von 2.1.34 erinnert.

A heißt *konvex*, wenn gilt:

$$\forall \lambda \in \mathbb{R}, 0 < \lambda < 1 \quad \forall x, y \in A : \lambda x + (1 - \lambda)y \in A.$$

A heißt *zirkulär* (im Englischen *circular*, im Französischen *cerclé*), wenn gilt:

$$S_{\mathbb{K}} \cdot A \subseteq A \quad \text{und} \quad 0 \in A.$$

A heißt *kreisförmig* (auch *äquilibriert*; im Englischen *balanced* oder auch *circled* oder *equilibrated*; im Französischen *équilibré*), wenn gilt:

$$B_{\mathbb{K}} \cdot A \subseteq A.$$

A heißt *absolutkonvex*, wenn A konvex und kreisförmig ist.

Die *konvexe Hülle* $\text{co}(A)$ von A ist die kleinste konvexe Teilmenge von X , die A enthält; sie ist gleich dem Durchschnitt aller A enthaltenden konvexen Teilmengen von X ; für nicht leeres A ist sie gleich der Menge

$$\left\{ \sum_{\nu=1}^n \lambda_{\nu} x_{\nu} : \lambda_1, \dots, \lambda_n \in \mathbb{R}^+, \sum_{\nu=1}^n \lambda_{\nu} = 1, x_1, \dots, x_n \in A, n \in \mathbb{N}^{\times} \right\}.$$

Die *kreisförmige Hülle* von A ist gleich der Menge

$$B_{\mathbb{K}} \cdot A.$$

Die *absolutkonvexe Hülle* $|\text{co}|(A)$ (auch mit $\Gamma(A)$ oder mit $\text{aco}(A)$ bezeichnet) von A ist der Durchschnitt aller A enthaltenden absolutkonvexen Teilmengen von X ; für nicht leeres A ist sie gleich der Menge

$$\left\{ \sum_{\nu=1}^n \lambda_{\nu} x_{\nu} : \lambda_1, \dots, \lambda_n \in \mathbb{K}, \sum_{\nu=1}^n |\lambda_{\nu}| \leq 1, x_1, \dots, x_n \in A, n \in \mathbb{N}^{\times} \right\}.$$

Ist τ eine Topologie auf X , so wird anstelle von $\text{cl}(\tau; \text{co}(A))$ auch $\overline{\text{co}}^{\tau}(A)$, oder, falls τ klar ist, einfach nur $\overline{\text{co}}(A)$ geschrieben; absolutkonvexe Hülle entsprechend.

Die Abbildung

$$p_A : X \rightarrow [0, \infty], \quad x \mapsto \inf \{ \lambda > 0 : x \in \lambda \cdot A \}$$

heißt das *Minkowski-Funktional* (im Englischen auch *Minkowski gauge*) oder auch die *Distanzfunktion* von A .

Man sagt, A *absorbiere* B , wenn ein $\lambda_0 > 0$ existiert, so dass $\lambda \cdot A \supseteq B$ für alle $\lambda \in \mathbb{K}$ mit $|\lambda| \geq \lambda_0$. Dies liegt genau dann vor, falls gilt:

$$\exists \lambda > 0 : B_{\mathbb{K}}^{\times} \cdot B \subseteq \lambda \cdot A.$$

A heißt *absorbierend* (oder auch *radial* (bei 0)), wenn es jede einelementige Teilmenge von X absorbiert. A ist also genau dann absorbierend, wenn $\forall x \in X \exists \lambda_0 > 0 \forall \lambda \in \mathbb{K}, |\lambda| \geq \lambda_0 : x \in \lambda \cdot A$; dies ist genau dann der Fall, falls gilt:

$$\forall x \in X \exists \lambda > 0 : B_{\mathbb{K}} \cdot x \subseteq \lambda \cdot A.$$

A heißt *ausgeglichen*, falls gilt:

$$\forall x \in X \exists \lambda > 0 : x \in \lambda \cdot A.$$

Dies ist im Englischen *nicht* als *balanced* zu übersetzen, da dies bereits für den Begriff *kreisförmig* verwendet wird.

2.2.9 Bemerkungen. Sei X ein Vektorraum über $\mathbb{K}$ und A eine Teilmenge von X . A ist genau dann absolutkonvex, wenn $\forall \lambda, \mu \in \mathbb{K}, |\lambda| + |\mu| \leq 1 \forall x, y \in A : \lambda x + \mu y \in A$.

Die absolutkonvexe Hülle von A ist gleich der konvexen Hülle der kreisförmigen Hülle von A .

Äquivalent zur Definition in 2.2.8 sagt man: A absorbiere B , wenn ein $\lambda_0 > 0$ existiert, so dass $\lambda \cdot B \subseteq A$ für alle $\lambda \in \mathbb{K}^{\times}$ mit $|\lambda| \leq \lambda_0$.

A ist genau dann ausgeglichen, wenn für alle $x \in X$ $p_A(x) < \infty$ ist.

Aus A absorbierend folgt stets, dass A ausgeglichen ist. Ist A kreisförmig, so gilt auch die Umkehrung.

2.2.10 Bemerkung (HEUSER [134](2002), Satz 161.4). Eine offene Teilmenge eines normierten Vektorraumes über $\mathbb{K}$ ist genau dann wegzusammenhängend, wenn sie zusammenhängend ist.

2.2.11. Die Definitionen der Termini Kegel und Keil sind in der Literatur nicht einheitlich. Die in 2.2.12 erfolgende Definition eines Kegels und eines Keils und auch die sich daran anschließenden beiden Nummern sind den folgenden Literaturstellen entlehnt: BOURBAKI [35](1987), II.2.4-5; SCHAEFER [255](1966), Seite 38; REED und SIMON [245](1980); EDWARDS [88](1965).

2.2.12 Definition. Sei X ein Vektorraum über $\mathbb{K}$. Ein *Kegel* (im Englischen *cone*, im Französischen *un cône*) K in X ist eine Teilmenge von X , für die $\mathbb{R}^{+\times} \cdot K \subseteq K$ gilt. Ein Kegel K in X heißt *punktiert* (im Englischen *pointed*), falls $0 \in K$ gilt, andernfalls *unpunktiert* (im Englischen *non-pointed*, im Französischen *épointé*). Ein *Kegel* K in X heißt *spitz* (im Englischen *salient*, im Französischen *saillant*) oder auch *echt* (im Englischen *proper*), falls $K \cap (-K) \subseteq \{0\}$ gilt. Ein Kegel K in X heißt *erzeugend* (im Englischen *generating*), falls $X = K - K$ gilt. Ein konvexer Kegel wird auch ein *Keil* (im Englischen *wedge*) genannt.

2.2.13 Bemerkung. Sei X ein Vektorraum über $\mathbb{K}$. Ein Kegel K in X ist genau dann konvex, wenn $K + K \subseteq K$ gilt.

2.2.14 Ordnungen und Kegel. Sei X ein Vektorraum über $\mathbb{R}$. Ist X prägeordnet, so ist die Menge $\{x \in X : x \geq 0\}$ ein punktierter, konvexer Kegel in X . Ist umgekehrt K ein punktierter, konvexer Kegel in X , so wird X per $x \leq y :\Leftrightarrow y - x \in K$ zu einem prägeordneten Vektorraum, der genau dann ein geordneter Vektorraum ist, wenn K auch spitz ist. Demgemäß sagt man, dass ein punktierter, konvexer Kegel K in X eine Präordnung auf X induziere und, dass ein punktierter, spitzer, konvexer Kegel K in X eine partielle Ordnung auf X induziere. In beiden Fällen schreibt man (X, K) für den so entstandenen (prä)geordneten Vektorraum über $\mathbb{R}$.

2.2.15. Es sei an die im Zusammenhang von monotonen Netzen in 1.1.15 eingeführte Terminologie erinnert. Sei X ein geordneter Vektorraum, U ein monoton steigend vollständiger Untervektorraum von X und $(x_\nu)_{\nu \in I}$ ein nach unten ordnungsbeschränktes, monoton fallendes Netz mit $x_\nu \in U$ für alle $\nu \in I$. Setze $y_\nu := -x_\nu$ für alle $\nu \in I$. Dann ist $(y_\nu)_{\nu \in I}$ ein nach oben ordnungsbeschränktes, monoton steigendes Netz. Nach Voraussetzung hat es ein in U liegendes Supremum y . Somit hat das Netz $(x_\nu)_{\nu \in I}$ das in U liegende Infimum $-y$. U ist also auch monoton fallend vollständig.

Man erkennt, dass ein Untervektorraum eines geordneten Vektorraumes genau dann monoton steigend vollständig ist, wenn er monoton fallend vollständig ist und definiert dementsprechend:

2.2.16 Definition. Sei X ein geordneter Vektorraum und U ein Untervektorraum von X . U heißt *monoton vollständig*, wenn U monoton steigend vollständig ist.

2.2.17 Bemerkung. Sei X ein geordneter Vektorraum, U ein Untervektorraum von X . KADISON und RINGROSE [167](1992) und TAKESAKI [274](2001), Seite 137, verlangen in ihrer, der monotonen Vollständigkeit entsprechenden Definition, dass nicht jedes nach *oben beschränkte*, monoton steigende Netz, sondern nur, dass jedes *beschränkte*, monoton steigende Netz in U ein Supremum in U hat. Liegt dies letzteres vor, so sei dies im Folgenden bezeichnet als *monoton vollständig nach KRT*, wobei die Buchstaben K, R und T die Anfangsbuchstaben der Namen der eben genannten Autoren seien. Nichtsdestotrotz ist die hier gegebene Definition 2.2.16 äquivalent zu der von Kadison, Ringrose und Takesaki.

Beweis. Dass aus der monotonen Vollständigkeit die monotone Vollständigkeit nach KRT folgt, ist klar. Sei also U monoton vollständig nach KRT. Sei $(x_\nu)_{\nu \in I}$ ein nach oben ordnungsbeschränktes, monoton steigendes Netz in U . Sei $\nu_0 \in I$ beliebig, aber

fest. Betrachte das Netz $(x_\nu)_{\nu \in J}$ mit $J := \{\nu \in I : \nu \geq \nu_0\}$. Dieses Netz ist durch x_{ν_0} nach unten ordnungsbeschränkt, also ordnungsbeschränkt und nach Voraussetzung hat es ein Supremum $x \in U$. Falls $I \setminus J \neq \emptyset$, sei $\mu \in I \setminus J$. Da I eine gerichtete Menge ist, existiert ein $\xi \in I$ mit $\mu \leq \xi$ und $\nu_0 \leq \xi$. Aus der ersten Ungleichung folgt $x_\mu \leq x_\xi$ und aus der zweiten Ungleichung folgt $\xi \in J$; also ist $x_\mu \leq x$, x auch das Supremum des Ausgangsnetzes $(x_\nu)_{\nu \in I}$ und U ist monoton vollständig. $\square$

2.2.18 Bemerkung und Definition. (a) Sei X ein Vektorraum über $\mathbb{K}$. Mit $\mathbb{K}$, aufgefasst als Vektorraum über $\mathbb{K}$, bezeichnet X^1 das direkte Produkt $X \times \mathbb{K}$.

(b) Sei X ein normierter Vektorraum über $\mathbb{K}$. Im Allgemeinen gibt es verschiedene Möglichkeiten die Norm von X auf X^1 fortzusetzen. Zum Beispiel ist auf X^1 per

$$\|(x, \mu)\| := \|x\| + |\mu| \quad \text{für alle } x \in X, \mu \in \mathbb{K}$$

eine Norm definiert, die genau dann vollständig ist, wenn die Norm von X vollständig ist (Definition 2.4.6). Der so erhaltene normierte Vektorraum über $\mathbb{K}$ wird wieder mit X^1 bezeichnet. Die Abbildung $x \mapsto (x, 0)$ ist ein isometrischer Isomorphismus von X auf einem Untervektorraum von X^1 .

2.2.19 Lineare Abbildungen. Seien X und Y zwei Vektorräume über $\mathbb{K}$.

Mit $L(X, Y)$ wird der Vektorraum über $\mathbb{K}$ aller $\mathbb{K}$ -linearen Abbildungen von X nach Y bezeichnet. Speziell heißt $L(X, \mathbb{K})$ der *algebraische Dualraum* von X und wird mit $X^\#$ bezeichnet. Elemente von $X^\#$ werden wie allgemein hin üblich suggestiv meist als $x^\#, y^\#, \dots$ notiert; des Weiteren wird dazu auch hier das Symbol ℓ verwendet. Die per $x \mapsto (x^\# \mapsto x^\#x)$ definierte *kanonische Inklusionsabbildung* von X nach $X^{\#\#}$ wird mit i_X bezeichnet; man nennt sie auch die *kanonische Einbettung* (im Englischen *canonical imbedding*) oder die *Auswertungsabbildung* (im Englischen *evaluation map*) von X in $X^{\#\#}$; genauer spricht man sie auch an als die *kanonische algebraische Inklusionsabbildung*, die anderen Benennungen entsprechend. Ist U ein Unterraum von $X^\#$, so wird mit $i_{X,U}$ die Abbildung

$$X \rightarrow U^\#, \quad x \mapsto i_X(x) \upharpoonright U$$

bezeichnet. Für $x \in X$ schreibt man für das Auswertungsfunktional $i_X(x)$ das Symbol $\hat{x}$. Für $T \in L(X, Y)$ ist die *algebraische duale Abbildung* von T die mit $T^\#$ bezeichnete lineare Abbildung

$$Y^\# \rightarrow X^\#, \quad y^\# \mapsto y^\# \circ T.$$

Es gilt also $\langle x, T^\#y^\# \rangle = \langle Tx, y^\# \rangle$ für alle $x \in X, y^\# \in Y^\#$.

Anstelle von $L(X, X)$ schreibt man $L(X)$.

Soll deutlich gemacht werden, dass der Skalarenkörper der Vektorräume $\mathbb{K}$ ist, so wird zum Beispiel die Bezeichnung $L_{\mathbb{K}}(X)$ anstelle von $L(X)$ verwendet.

Ein $T \in L(X, Y)$ heißt *Rang-Eins* (im Englischen *rank-one*) oder auch *vom Rang-Eins*, falls $T(X)$ eindimensional ist. Ein $T \in L(X, Y)$ heißt *vom endlichen Rang* (im Englischen *finite rank*), falls $T(X)$ endlichdimensional ist. Der Vektorraum über $\mathbb{K}$, der aus allen $T \in L(X, Y)$ vom endlichen Rang besteht, wird mit $F(X, Y)$ bezeichnet. Statt $F(X, X)$ wird $F(X)$ geschrieben.

Ist $T \in L(X, Y)$, so wird als *Kern* von T , in Zeichen $\ker T$, der Untervektorraum $\{x \in X : T(x) = 0\}$ bezeichnet.

Sei $T \in L(X)$ und $\lambda \in \mathbb{K}$. λ heißt *Eigenwert* von T , wenn es ein $x \in X^\times$ mit $Tx = \lambda x$ gibt; solch ein x heißt dann ein *Eigenvektor* von T zum Eigenwert λ . Der Unterraum $\ker(\lambda \cdot \text{Id}_X - T)$ heißt der *Eigenraum* von T und wird mit $\text{Eig}(T, \lambda)$ oder $\text{Eig}_T(\lambda)$ bezeichnet.

2.2.20 Lemma. Sei X ein Vektorraum über $\mathbb{K}$, $\ell \in X^\#$ und seien $p_1, \dots, p_n$ Halbnormen auf X mit $|\ell(x)| \leq p_1(x) + \dots + p_n(x)$ für alle $x \in X$. Dann existieren $\ell_1, \dots, \ell_n \in X^\#$ mit $\ell = \ell_1 + \dots + \ell_n$ und $|\ell_\nu(x)| \leq p_\nu(x)$ für alle $x \in X$ und $\nu = 1, \dots, n$.

Beweis. (STRÄTILÄ und ZSIDÓ [270](1979), Seite 13; KADISON und RINGROSE [167](1991), Band III, Seite 1.) Sei X^n das direkte Produkt von n Kopien von X und D der „diagonale“ Untervektorraum $\{(x, \dots, x) \in X^n : x \in X\}$. Auf X^n definiere man eine Halbnorm p per $p(x_1, \dots, x_n) := p_1(x_1) + \dots + p_n(x_n)$ für alle $(x_1, \dots, x_n) \in X^n$. Auf D definiere man $\ell_D \in D^\#$ per $\ell_D(x, \dots, x) := \ell(x)$ für alle $x \in X$. Nach Voraussetzung gilt $|\ell_D(x, \dots, x)| = |\ell(x)| \leq p_1(x) + \dots + p_n(x) = p(x, \dots, x)$ für alle $x \in X$, das heißt, die Linearform ℓ_D wird auf dem Untervektorraum D von der Halbnorm p majorisiert. Nach dem Satz von Hahn-Banach kann ℓ_D zu einem $\tilde{\ell} \in (X^n)^\#$ mit $|\tilde{\ell}(x_1, \dots, x_n)| \leq p(x_1, \dots, x_n)$, für alle $(x_1, \dots, x_n) \in X^n$, erweitert werden.

Man definiere für $\nu = 1, \dots, n$: $\ell_\nu(x) := \tilde{\ell}(0, \dots, x, \dots, 0)$ für alle $x \in X$, wobei auf der rechten Seite an der ν -ten Position das x steht und an allen anderen Positionen 0 steht. Dann hat man: $\ell_1(x) + \dots + \ell_n(x) = \tilde{\ell}(x, \dots, x) = \ell_D(x, \dots, x) = \ell(x)$ und $|\ell_\nu(x)| = |\tilde{\ell}(0, \dots, x, \dots, 0)| \leq p(0, \dots, x, \dots, 0) = p_1(0) + \dots + p_\nu(x) + \dots + p_n(0) = p_\nu(x)$ für alle $x \in X$ und $\nu = 1, \dots, n$. $\square$

2.2.21 n -lineare Abbildungen. Sei $n \in \mathbb{N}^\times$ und seien $X_1, \dots, X_n, Y$ Vektorräume über den gleichen Körper $\mathbb{K}$. Eine Abbildung $T: X_1 \times \dots \times X_n \rightarrow Y$ heißt n -linear, wenn für jedes $(a_1, \dots, a_n) \in X_1 \times \dots \times X_n$ und jedes $k \in \{1, \dots, n\}$ die durch $X_k \rightarrow Y$, $x \mapsto T(a_1, \dots, a_{k-1}, x, a_{k+1}, \dots, a_n)$ definierte Abbildung linear ist. Der Vektorraum über $\mathbb{K}$ aller n -linearen Abbildungen von $X_1 \times \dots \times X_n$ nach Y wird mit $L(X_1, \dots, X_n; Y)$ bezeichnet. Gilt $X_1 = \dots = X_n$ so wird mit $X := X_1$ anstelle von $L(X_1, \dots, X_n; Y)$ die Bezeichnung $L^n(X, Y)$ verwendet. $L^0(X, Y)$ wird als Y definiert. Ist $X = Y$, so wird anstelle von $L^n(X, Y)$ einfach nur $L^n(X)$ geschrieben. Für Abbildungen $T \in L^n(X, Y)$ und Permutationen $\sigma \in \mathfrak{S}_n$ sei dann σT definiert als $\sigma T(x_1, \dots, x_n) := T(x_{\sigma(1)}, \dots, x_{\sigma(n)})$ für alle $x_1, \dots, x_n \in X$; σ heißt in diesem Zusammenhang auch *Permutationsoperator*. Dann ist auch $\sigma T \in L^n(X, Y)$. Eine Abbildung $T \in L^n(X, Y)$ heißt *symmetrisch*, wenn $\sigma T = T$ für alle $\sigma \in \mathfrak{S}_n$ gilt. 2-lineare Abbildungen werden bevorzugt *bilinear* und 3-lineare Abbildungen werden bevorzugt *trilinear* genannt.

2.2.22 Bilineare Abbildungen. Seien X, Y und Z drei Vektorräume über $\mathbb{K}$. Sei $B \in L(X, Y; Z)$. Für $Z = \mathbb{K}$ wird die Abbildung B eine *Bilinearform* genannt. Für spätere Bezugnahmen definiere man hier die folgenden zwei linearen Abbildungen:

$$\begin{aligned} S: X &\rightarrow L(Y, Z), & x &\mapsto (y \mapsto B(x, y)); \\ R: Y &\rightarrow L(X, Z), & y &\mapsto (x \mapsto B(x, y)). \end{aligned}$$

Es ist also $S(x) = B(x, \cdot)$ für alle $x \in X$, und es ist $R(y) = B(\cdot, y)$ für alle $y \in Y$.

Die Abbildung $\langle \cdot, \cdot \rangle: X \times X^\# \rightarrow \mathbb{K}$ mit $\langle x, x^\# \rangle := x^\#(x)$ für alle $x \in X$, $x^\# \in X^\#$ heißt die *kanonische Bilinearform* von X .

2.2.23 Bezeichnung. Der Buchstabe „S“ kommt von dem lateinischen Wort „sinister“ für „links“. Er wird hier verwendet, da der Buchstabe „L“ hier zu Missverständnissen bezüglich des Wertebereiches führen könnte. Da das Wort „rechts“ auf lateinisch „dexter“ heißt, müsste man hier konsequenterweise die Abbildung R mit D bezeichnen. Der Buchstabe „D“ wird aber im Kapitel 4 bereits in Anlehnung an das Wort „Derivation“ für gewisse Abbildungen verwendet werden; außerdem steht D auch für den Differentialoperator (siehe 4.1.2).

2.2.24 Sesquilineare Abbildungen. Seien X, Y und Z drei Vektorräume über $\mathbb{K}$. Eine Abbildung $B: X \times Y \rightarrow Z$ heißt *sesquilinear*, wenn sie in der ersten Komponente

linear und in der zweiten Komponente semilinear ist. Für $Z = \mathbb{K}$ wird eine sesquilineare Abbildung eine *Sesquilinearform* genannt.

Sei $B: X \times Y \rightarrow \mathbb{K}$ eine Sesquilinearform. Für spätere Bezugnahmen definiere man hier die folgenden zwei semilinearen Abbildungen:

$$\begin{aligned} S: X &\rightarrow L(Y, Z), & x &\mapsto (y \mapsto \overline{B(x, y)}); \\ R: Y &\rightarrow L(X, Z), & y &\mapsto (x \mapsto B(x, y)). \end{aligned}$$

Es ist also $S(x) = \overline{B(x, \cdot)}$ für alle $x \in X$, und es ist $R(y) = B(\cdot, y)$ für alle $y \in Y$.

Im Fall von $X = Y$ heißt die Abbildung B *hermitesch*, wenn für alle $x, y \in X$ die Gleichung $B(x, y) = \overline{B(y, x)}$ gilt.

2.2.25 Definition. Seien X, Y und Z drei Vektorräume über $\mathbb{K}$. Sei $B: X \times Y \rightarrow Z$ eine bilineare oder eine sesquilineare Abbildung. Man betrachte das somit vorliegende Annihilatorsystem $(X, Y; B; Z)$. Man nennt $(X, Y; B; Z)$ ein *Bilinearsystem* (bzw. ein *Sesquilinearsystem*), wenn B bilinear (bzw. sesquilinear) ist. Ist B eine Bilinearform, so heißt $(X, Y; B)$ ein *Dualsystem*.

Seien S und R die Abbildungen aus 2.2.22, falls B bilinear ist und aus 2.2.24, falls B sesquilinear ist. Gemäß 2.1.4 ist die Abbildung B genau dann *nicht rechts ausgeartet*, wenn die Abbildung R injektiv ist. Analog ist die Abbildung B genau dann *nicht links ausgeartet*, wenn die Abbildung S injektiv ist. Dementsprechend ist die Abbildung B *nicht ausgeartet*, wenn sowohl S als auch R injektiv sind. Ist die Abbildung B symmetrisch oder hermitesch, dann gilt $S = R$ und B ist genau dann nicht ausgeartet, wenn R injektiv ist.

Sei $Z = \mathbb{K}$ und sei B nicht rechts ausgeartet. Dann heißt die inverse Abbildung $\text{grad}_R := R^{-1}$ von $\text{ran}(R)$ nach Y die *rechte Gradientenabbildung*. Für $Z = \mathbb{K}$ und B nicht links ausgeartet, ist die *linke Gradientenabbildung* analog als $\text{grad}_S := S^{-1}$ definiert. Ist B sowohl symmetrisch oder hermitesch als auch nicht ausgeartet, so heißt $\text{grad} := \text{grad}_R$ die *Gradientenabbildung*.

Sei $X = Y$ und $Z = \mathbb{K}$. Die Abbildung B heißt *positiv*, wenn für alle $x \in X$ die Ungleichung $B(x, x) \geq 0$ gilt. Die Abbildung B heißt *positiv definit*, wenn für alle $x \in X^\times$ die Ungleichung $B(x, x) > 0$ gilt. Ist B eine positive, hermitesche Sesquilinearform, so heißt B ein *Semiskalarprodukt* auf X . Die Abbildung $X \rightarrow \mathbb{K}$, $x \mapsto B(x, x)^{1/2}$, ist dann eine Halbnorm und wenn nicht etwas anderes vereinbart ist, ist X dann stets als mit dieser Halbnorm ausgestattet aufzufassen; die durch diese Halbnorm induzierte Topologie ist genau dann Hausdorff'sch, wenn diese Halbnorm eine Norm ist. Ist B eine positiv definite, hermitesche Sesquilinearform, so heißt B ein *Skalarprodukt* auf X und X , ausgestattet mit B , ein *Prähilbertraum* über $\mathbb{K}$. Ist $(X, B) := (X, X; B)$ ein Prähilbertraum über $\mathbb{K}$, so ist X , falls nicht ausdrücklich etwas anderes vereinbart ist, stets mit der Norm $\|x\| := B(x, x)^{1/2}$, $x \in X$, versehen. Es sei darauf hingewiesen, dass BOURBAKI [35](1987) unter einem Prähilbertraum einen Vektorraum, versehen mit einem Semiskalarprodukt, versteht.

2.2.26 Rang-Eins Elemente. Sei $(X, F; B; \mathbb{K})$ ein Dualsystem und Y ein Vektorraum über $\mathbb{K}$. Sei $f \in F$ und $y \in Y$. Dann wird mit $y \otimes f$ das durch

$$(y \otimes f)(x) := B(x, f)y \quad \text{für alle } x \in X$$

definierte Element aus $L(X, Y)$ bezeichnet.

Seien X und Y zwei Vektorräume über $\mathbb{K}$. Jedes Rang-Eins Element von $L(X, Y)$ lässt sich als ein $y \otimes x^\#$, $x^\# \in X^{\# \times}$, $y \in Y^\times$, schreiben. Für eine Erörterung bezüglich des Zusammenhanges mit dem Tensorprodukt, siehe PALMER [231](1994), Kapitel 1.10.

2.2.27 Bemerkung und Definition (KÖTHE [187](1966), Seite 74). Sei $(X, Y; B; Z)$ ein Bilinearsystem oder ein Sesquilinearsystem. Typische Beispiele für orthosymmetrische Systeme liegen vor, wenn B symmetrisch bzw. hermitesch oder die kanonische Bilinearform von X ist.

Sei $M \subseteq X$ und $N \subseteq Y$. Dann ist der Rechtsannihilator von M ein Untervektorraum von Y und entsprechend ist der Linksannihilator von N ein Untervektorraum von X . Dementsprechend wird im orthosymmetrischen Fall die Orthogonalmenge von M auch der *Orthogonalraum* von M genannt; Orthogonalmenge von N entsprechend.

Es gilt:

$$M^\perp = (\text{span } M)^\perp. \quad (2.26)$$

Beweis. Sei $y \in M^\perp$ und $x \in \text{span}(M)$. Also $x = a_1 m_1 + \dots + a_n m_n$ für gewisse $a_1, \dots, a_n \in \mathbb{K}$, $m_1, \dots, m_n \in M$. Es ist $B(x, y) = a_1 B(m_1, y) + \dots + a_n B(m_n, y) = 0$, also $y \in (\text{span } M)^\perp$. $\square$

Für $\{0\} \subseteq X$ gilt $\{0\}^\perp = Y$. Trivialerweise gilt zwar $\{0\} \subseteq X^\perp$ und $\{0\} \subseteq Y_\perp$, aber bezüglich des Falles, wann jeweils auch Gleichheit gilt, siehe die Bemerkungen 2.2.28 und 2.2.29.

Liegt der spezielle Fall $X = Y = Z$ vor, so ist X , versehen mit B als eine Verknüpfung auf X , ein nichtassoziativer Ring; er wird mit (X, B) bezeichnet. Ist dabei noch spezieller B bilinear, so ist, die Definition 2.3.1 vorausgreifend, X eine nichtassoziative Algebra über $\mathbb{K}$.

2.2.28 Bemerkung. Seien (X, x_0) und (Z, z_0) zwei punktierte Mengen. Sei A ein Annihilatorsystem $((X, x_0), Y; B; (Z, z_0))$. Dann sind die folgenden Aussagen äquivalent:

- (a) A ist in X trennend.
- (b) Für alle $x \in X^\times$ existiert ein $y \in Y$ mit $B(x, y) \neq z_0$.

Liegt nun der Fall vor, dass $\{x_0\}^\perp = Y$ gilt, so sind die folgenden vier Aussagen zu jeder der beiden vorstehenden Aussagen äquivalent:

- (c) Nur $x = x_0$ erfüllt für alle $y \in Y$ die Gleichung $B(x, y) = z_0$.
- (d) $\bigcap_{y \in Y} B(\cdot, y)^{-1}(z_0) = \{x_0\}$.
- (e) $Y_\perp = \{x_0\}$.
- (f) $\{x_0\} \subseteq X$ ist annihilatorabgeschlossen.

Ist das Annihilatorsystem A sogar ein Bilinearsystem oder ein Sesquilinearsystem, so ist die folgende Aussage zu jeder der vorstehenden Aussagen (a) bis (f) äquivalent:

- (g) Die Abbildung $S: X \rightarrow L(Y, Z)$, $x \mapsto B(x, \cdot)$ ist injektiv.

Beweis. (a) $\Leftrightarrow A$ ist nicht links ausgeartet $\Leftrightarrow \neg(\exists x \in X^\times \forall y \in Y : B(x, y) = z_0) \Leftrightarrow$ (b).

Liege nun der Fall vor, dass für alle $y \in Y$ die Gleichung $B(x_0, y) = z_0$ gelte. Dann ist die Implikation (b) $\Rightarrow$ (c) klar. Gelte (c) und sei $x \in X^\times$. Nun bemerke nur die Äquivalenz $\neg(\forall y \in Y : B(x, y) = z_0) \Leftrightarrow \exists y \in Y : B(x, y) \neq z_0$. Das heißt, es gilt (b).

Das (c) jeweils zu (d) und (e) äquivalent ist, ist klar.

Die Implikation (e) $\Rightarrow$ (f) ist wegen 1.1.19(e) klar. Die umgekehrte Implikation folgt sofort aus der vorausgesetzten Bedingung $\{x_0\}^\perp = Y$.

Zu (g): Man bemerke nur die folgenden Äquivalenzen: S injektiv $\Leftrightarrow \ker S = \{0\}$ und $x \in \ker S \Leftrightarrow S(x) = B(x, \cdot) = (Y \rightarrow Z, y \mapsto 0) \Leftrightarrow \forall y \in Y : B(x, y) = 0$. Mit ihnen sieht man sofort die Äquivalenz von (c) und (g). $\square$

2.2.29 Bemerkung. Seien (Y, y_0) und (Z, z_0) zwei punktierte Mengen. Sei A ein Annihilatorsystem $(X, (Y, y_0); B; (Z, z_0))$. Dann sind die folgenden Aussagen äquivalent:

- (a) A ist in Y trennend.
- (b) Für alle $y \in Y^\times$ existiert ein $x \in X$ mit $B(x, y) \neq z_0$.

Liegt nun der Fall vor, dass $\{y_0\}^\perp = X$ gilt, so sind die folgenden vier Aussagen zu jeder der beiden vorstehenden Aussagen äquivalent:

- (c) Nur $y = y_0$ erfüllt für alle $x \in X$ die Gleichung $B(x, y) = z_0$.
- (d) $\bigcap_{x \in X} B(x, \cdot)^{-1}(z_0) = \{y_0\}$.
- (e) $X^\perp = \{y_0\}$.
- (f) $\{y_0\} \subseteq Y$ ist annihilatorabgeschlossen.

Ist das Annihilatorsystem A sogar ein Bilinearsystem oder ein Sesquilinearsystem, so ist die folgende Aussage zu jeder der vorstehenden Aussagen (a) bis (f) äquivalent:

- (g) Die Abbildung $R: Y \rightarrow L(X, Z)$, $y \mapsto B(\cdot, y)$ ist injektiv.

Beweis. Ganz analog wie in 2.2.28. □

2.2.30 Bemerkung. Ist R surjektiv, dann ordnet also grad_R jedem Funktional ℓ aus $X^\#$ genau das eine Element $y \in Y$ zu, so dass für alle $x \in X$ die Gleichung $\ell(x) = B(x, y)$ gilt. Entsprechendes gilt, wenn B nicht links ausgeartet ist.

2.2.31 Dualsysteme (BOURBAKI [35](1987), II.6.1). Sei $(X, Y; B)$ ein Dualsystem. Andere übliche sprachliche Ausdrucksformen für das Dualsystem $(X, Y; B)$ sind zum Beispiel, dass man sagt, dass die Bilinearform B die Vektorräume X und Y *in Dualität setze*, dass X und Y (bezüglich B) *in Dualität seien*, oder auch, dass das geordnete Paar (X, Y) eine *Paarung* (im Englischen *pairing*) oder ein *duales Paar* (bezüglich B) sei.

Ist $(X, Y; B)$ ein trennendes Dualsystem, dann kann X per S aus 2.2.22 mit einem Unterraum über $\mathbb{K}$ von $Y^\#$ identifiziert werden. Identifiziert man wie üblich Y mit $i_Y(Y)$, so kann dann des Weiteren B identifiziert werden mit der Restriktion der kanonischen Bilinearform von $Y^\#$ auf $X \times Y$:

$$X \times Y \rightarrow \mathbb{K}, \quad (y^\#, y) \mapsto \langle y, y^\# \rangle.$$

Analog kann Y per R aus 2.2.22 mit einem Unterraum über $\mathbb{K}$ von $X^\#$ identifiziert werden. Dann kann man B identifizieren mit der Restriktion der kanonischen Bilinearform von X auf $X \times Y$:

$$X \times Y \rightarrow \mathbb{K}, \quad (x, x^\#) \mapsto \langle x, x^\# \rangle.$$

2.2.32 Bemerkung. HEUSER [133] nennt ein links trennendes Dualsystem ein rechtes Dualsystem und entsprechend ein rechts trennendes Dualsystem ein linkes Dualsystem. Diese Bezeichnung von Heuser wird hier nicht verwendet.

KÖTHE [187](1966) definiert Dualsysteme direkt so, dass sie trennend sind.

2.2.33 Bemerkung und Definition. Sei X eine Menge und F eine Menge von Funktionen auf X . Man sagt, F *trenne die Punkte von X* , wenn es für je zwei verschiedene Punkte $s, t \in X$ ein $f \in F$ gibt mit $f(s) \neq f(t)$. Dementsprechend hat man speziell für Vektorräume: Sei X ein Vektorraum über $\mathbb{K}$ und S eine Teilmenge von $X^\#$. Man sagt, S *trenne die Punkte von X* , oder auch, dass S *total* sei (über X), wenn für jedes $x \in X^\times$ ein $\ell \in S$ existiert mit $\ell(x) \neq 0$. S trennt also genau dann die Punkte von X , wenn

das Dualsystem $(X, \text{span}(S))$ trennend in X ist. Dies ist nach 2.2.28(e) und Gleichung (2.26) aus 2.2.27 äquivalent zu $S_\perp = \{0\}$ bezüglich des Dualsystems $(X, \text{span}(S))$. Wie man leicht sieht, ist dies wiederum dazu äquivalent, dass $S_\perp = \{0\}$ gilt bezüglich eines beliebigen Dualsystems (X, Y) , solange nur $S \subseteq Y \subseteq X^\#$ erfüllt ist.

2.2.34 Bemerkungen (BOURBAKI [35](1987), II.6.1). Sei X ein Vektorraum über $\mathbb{K}$. Dann ist $(X, X^\#)$ ein trennendes Dualsystem. Für jeden Untervektorraum U von $X^\#$ ist das Dualsystem (X, U) trennend in U . Ist X endlichdimensional, so ist der einzige Untervektorraum U von $X^\#$, mit dem (X, U) ein trennendes Dualsystem ist, der Raum $X^\#$ selbst. Ist X unendlichdimensional und ist U ein Untervektorraum von $X^\#$, so kann das Dualsystem (X, U) in X trennend sein, selbst wenn $U \neq X^\#$ ist.

Ist U ein Untervektorraum von $X^\#$ und ist (X, U) ein trennendes Dualsystem, dann ist auch für jeden Untervektorraum V von $X^\#$ mit $U \subseteq V$ das Dualsystem (X, V) trennend.

2.2.35 Definition (TAYLOR und LAY [275](1980), Seiten 122, 125). Sei X ein Vektorraum über $\mathbb{K}$ und M eine Teilmenge von X . Ein Untervektorraum U von X heißt *maximal*, wenn er in der durch die Mengeninklusion prägeordneten Menge aller echten Untervektorräume von X ein maximales Element ist. M heißt ein *affiner Untervektorraum* von X , wenn es ein $x \in X$ und ein Untervektorraum U von X mit $M = x + U$ gibt; ist dabei U ein maximaler Untervektorraum von X , dann heißt M eine *Hyperebene*.

Sei $x^\# \in X^{\#\times}$ und $\lambda \in \mathbb{R}$. Durch die Menge $(\text{Re } x^\#)^{-1}(\lambda)$ (eine Hyperebene in $X_\mathbb{R}$) sind vier sogenannte *Halbräume* (im Englischen *half spaces*) erklärt: $\{x \in X : \text{Re } x^\# x < \lambda\}$, $\{x \in X : \text{Re } x^\# x > \lambda\}$, $\{x \in X : \text{Re } x^\# x \leq \lambda\}$ und $\{x \in X : \text{Re } x^\# x \geq \lambda\}$. Der Durchschnitt einer konvexen Teilmenge von X mit einem Halbraum von X wird *Kalotte* (im Englischen *calotte*) genannt.

2.2.36 (KÖTHE [187](1966), Seite 74). Man betrachte das Dualsystem $(X, X^\#)$ für einen Vektorraum X über $\mathbb{K}$. Jeder Unterraum von X ist orthogonalabgeschlossen. Wenn X unendlichdimensional ist, so gibt es in $X^\#$ stets Unterräume, insbesondere Hyperebenen, die nicht orthogonalabgeschlossen sind. Jeder endlichdimensionale Unterraum U von $X^\#$ ist orthogonalabgeschlossen. Orthogonalabgeschlossene Unterräume von $X^\#$ werden auch als *algebraisch gesättigt* bezeichnet.

Zwei Unterräume U und V von X sind genau dann gleich, wenn $U^\perp = V^\perp$ gilt. Die analoge Aussage für Unterräume von $X^\#$ gilt im Allgemeinen nicht, siehe zum Beispiel TAYLOR und LAY [275](1980), Seite 46.

Ist $X = U_1 \oplus U_2$, so ist $X^\# = U_1^\perp \oplus U_2^\perp$. Ist $X^\# = V_1 \oplus V_2$ mit orthogonalabgeschlossenen V_1 und V_2 , so ist $X = V_1^\perp \oplus V_2^\perp$.

Ist U ein Unterraum von X , so gilt per $y^\# \mapsto (x \mapsto y^\#(x + U))$ die Isomorphie $(X/U)^\# \simeq U^\perp$ und per $x^\# + U^\perp \mapsto x^\# \upharpoonright U$ die Isomorphie $X^\# / U^\perp \simeq U^\#$ (TAYLOR und LAY [275](1980), Seite 48). Für die entsprechende Aussage in topologischen Vektorräumen siehe Satz 2.4.36

2.2.37 (BOURBAKI [35](1987), II.6.5). Sei (X, Y) ein Dualsystem. Sei U ein Untervektorraum von X und V ein Untervektorraum von Y mit $\langle U, V \rangle = \{0\}$. Dann gilt: Das in kanonischer Weise gebildete Dualsystem $(U, Y/V)$ ist genau dann in U trennend, wenn $Y^\perp \cap U = \{0\}$ gilt; es ist genau dann in Y/V trennend, wenn $V = U^\perp$ gilt.

2.2.38 Definition. Seien X und Y zwei Vektorräume über $\mathbb{K}$. Eine Abbildung $Q: X \rightarrow Y$ heißt *quadratisch*, wenn die folgenden beiden Bedingungen gelten:

- (a) $Q(\lambda x) = \lambda^2 Q(x)$ für alle $\lambda \in \mathbb{K}$, $x \in X$,

(b) die Abbildung $B: X \times X \rightarrow Y$,

$$(x, y) \mapsto Q(x + y) - Q(x) - Q(y)$$

ist bilinear.

Die Abbildung B heißt die zu Q assoziierte bilineare Abbildung. Wenn B nicht ausgeartet ist, so heißt Q nicht ausgeartet.

Sei Q eine quadratische Abbildung von X nach Y mit $Y = \mathbb{K}$. Dann wird Q eine *quadratische Form* auf X genannt und die zu Q assoziierte bilineare Abbildung B heißt die *Bilinearform bezüglich Q* .

2.2.39. In der Situation von Definition 2.2.38 ist die Abbildung B symmetrisch. Des Weiteren gilt für alle $x \in X$: $B(x, x) = 2Q(x)$. Man kann also aus B die quadratische Form Q zurückerhalten.

2.2.40 n -homogene Polynome. Sei $n \in \mathbb{N}^\times$ und seien X und Y zwei Vektorräume über $\mathbb{K}$. Eine Abbildung $p: X \rightarrow Y$ heißt ein *n -homogenes Polynom* (oder *homogenes Polynom vom Grad n*), wenn ein $T \in L^n(X, Y)$ existiert mit $p(x) = T(x, \dots, x)$ für alle $x \in X$; in diesem Fall ist $\tilde{p} := \frac{1}{n!} \sum_{\sigma \in \mathfrak{S}_n} \sigma T$ die eindeutig bestimmte, symmetrische Abbildung mit $p(x) = \tilde{p}(x, \dots, x)$ für alle $x \in X$; falls nicht ausdrücklich etwas anderes vereinbart ist, ist die aufgesetzte Tilde im vorliegenden Text immer in diesem Sinne zu verstehen. Der Vektorraum aller n -homogenen Polynome von X nach Y wird mit $P^n(X, Y)$ bezeichnet. $P^0(X, Y)$ ist der Vektorraum aller konstanten Abbildungen von X nach Y . Es ist $P^1(X, Y) = L(X, Y)$.

Speziell mit den Bezeichnungen von Definition 2.2.38 gilt, dass wegen $Q(x) = \frac{1}{2}B(x, x)$ für alle $x \in X$, eine Abbildung $Q: X \rightarrow Y$ genau dann quadratisch ist, wenn Q ein 2-homogenes Polynom ist.

Sei $F \in L^n(X, Y)$, dann wird die Abbildung $\hat{F} \in L(X, Y)$, die durch $\hat{F}(x) := F(x, \dots, x)$, $x \in X$, definiert ist, das zu F assoziierte *n -homogene Polynom* genannt; speziell für $n = 2$ heißt dann $Q := \hat{F}$ die zu F assoziierte *quadratische Abbildung*, bzw., wenn dabei auch noch $Y = \mathbb{K}$ gilt, die zu F assoziierte *quadratische Form*.

2.2.41. Sei $n \in \mathbb{N}^\times$, seien X, Y zwei Vektorräume über $\mathbb{K}$ und $F \in L^n(X, Y)$. Dann ist $F_{\text{sym}} := \frac{1}{n!} \sum_{\sigma \in \mathfrak{S}_n} \sigma F \in L^n(X, Y)$ symmetrisch und mit der Bezeichnung von 2.2.40 gilt $\widehat{F_{\text{sym}}} = \hat{F}$ und die folgende Polarisationsgleichung:

$$F_{\text{sym}}(x_1, \dots, x_n) = \frac{1}{2^n n!} \sum_{t_1 = \pm 1} \dots \sum_{t_n = \pm 1} \frac{\hat{F}(t_1 x_1 + \dots + t_n x_n)}{t_1 \dots t_n}$$

für alle $x_1, \dots, x_n \in X$. Ist also F symmetrisch, so ist F durch $\hat{F}$ eindeutig bestimmt und unter anderem lässt sich dann F aus $\hat{F}$ zurückgewinnen.

Speziell für $n = 2$ und beliebigem $F \in L^2(X, Y)$ gilt mit $Q := \hat{F}$ für alle $x, y \in X$:

$$\begin{aligned} F_{\text{sym}}(x, y) &= \frac{1}{8} (Q(x + y) - Q(-x + y) - Q(x - y) + Q(-x - y)) \\ &= \frac{1}{4} (Q(x + y) - Q(x - y)) \\ &= \frac{1}{2} (F(x, y) + F(y, x)), \end{aligned}$$

also $F_{\text{sym}} = \frac{1}{2}B$, wenn B die zu $Q = \hat{F}$ assoziierte bilineare Abbildung bezeichnet.

2.2.42 Bemerkung (WEIDMANN [287](1976), Seite 13; [288](2000), Seite 21). Seien X und Y zwei Vektorräume über $\mathbb{K}$. Ist $B: X \times X \rightarrow Y$ eine sesquilineare Abbildung, so kann man die Abbildung $K: X \rightarrow Y$, $x \mapsto B(x, x)$ definieren. Dann gilt $K(\lambda x) = |\lambda|^2 K(x)$ für alle $\lambda \in \mathbb{K}$, $x \in X$. Die Abbildung $F: X \times X \rightarrow Y$, $(x, y) \mapsto K(x + y) - K(x) - K(y)$ ist für $\mathbb{K} = \mathbb{R}$ bilinear und für $\mathbb{K} = \mathbb{C}$ im Allgemeinen weder bilinear noch sesquilinear.

Für $\mathbb{K} = \mathbb{C}$ lässt sich die sesquilineare Abbildung B sowohl aus der Abbildung K als auch aus der Abbildung F per Polarisierung zurückgewinnen:

$$\begin{aligned} B(x, y) &= \frac{1}{4} \left(K(x + y) - K(x - y) + iK(x + iy) - iK(x - iy) \right) \\ &= \frac{1}{4} \sum_{\epsilon^4=1} \epsilon K(x + \epsilon y) \\ &= \frac{1}{2} \left(K(x + y) - K(x) - K(y) \right) + \frac{1}{2} i \left(K(x + iy) - K(x) - K(y) \right) \\ &= \frac{1}{2} \left(F(x, y) + iF(x, iy) \right) \quad \text{für alle } x, y \in X. \end{aligned}$$

Dies ist bei $\mathbb{K} = \mathbb{R}$ nicht immer möglich. Zum Beispiel ist auf $\mathbb{R}^2$ für die Sesquilinearform $B: (x, y) \mapsto x^t A y$ mit $A = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$, also $\left(\begin{pmatrix} x_1 \\ x_2 \end{pmatrix}, \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} \right) \mapsto x_1 y_2 - x_2 y_1$, die Abbildung $K: x \mapsto B(x, x)$ identisch Null. (Siehe auch 3.4.9.)

2.2.43 Strecken. Sei X ein Vektorraum über $\mathbb{K}$ und $a, b \in X$. Dann heißt die Menge $\{\lambda a + (1 - \lambda)b \in X : 0 \leq \lambda \leq 1\}$ die *Strecke* (genauer *abgeschlossene Strecke* und noch genauer *algebraisch abgeschlossene Strecke*) mit den Endpunkten a und b , in Zeichen $[a, b]$. Verbreitet dafür sind auch die Begriffe *Intervall* und *Segment*. Analog wie für Intervalle auf der reellen Zahlengerade bezeichne $(a, b] := [a, b] \setminus \{a\}$ und $[a, b) := [a, b] \setminus \{b\}$ die beiden *halboffenen Strecken* und $(a, b) := [a, b] \setminus \{a, b\}$ die *offene Strecke* mit den Endpunkten a und b . (Die genaueren Bezeichnungen entsprechend.)

Man beachte, dass für das in 2.2.6 im Rahmen von geordneten Vektorräumen definierte Ordnungsintervall dieselbe Symbolik verwendet wird. Aus dem jeweiligen Zusammenhang wird immer klar sein, was mit $[a, b]$, etc. gemeint ist.

Sei $A \subseteq X$ und $a \in A$. Dann heißt A *sternförmig bezüglich a* (im Englischen *starshaped relative to a* oder auch *starlike about a*), falls für jedes $x \in A$ die abgeschlossene Strecke $[a, x]$ eine Teilmenge von A ist. Offensichtlich ist somit A genau dann sternförmig bezüglich a , wenn es Punkte $a_\nu \in X$, $\nu \in I$, I eine Indexmenge, gibt, so dass $A = \bigcup_{\nu \in I} [a, a_\nu]$ gilt.

2.2.44 Extremale Teilmengen (GILES [109](1982), Seite 85; RUDIN [251] (1973), Seite 70; TAYLOR [275](1980), Seite 181). Sei X ein Vektorraum über $\mathbb{K}$. Seien A und K zwei Teilmengen von X mit $A \subseteq K$. (Man bemerke, dass hierbei die leere Menge nicht ausgeschlossen wird.) Die Menge A heißt eine *extremale Teilmenge* von K (im Englischen *extreme set*), wenn gilt:

$$\forall a, b \in K \quad \forall \lambda \in \mathbb{R}, 0 < \lambda < 1 : (\lambda a + (1 - \lambda)b \in A \Rightarrow a, b \in A).$$

Diese Bedingung schreibt sich unter Verwendung der in 2.2.43 erklärten offenen Strecken als

$$\forall a, b \in K : ((a, b) \cap A \neq \emptyset \Rightarrow a, b \in A),$$

in Worten also: Jede offene Strecke (a, b) mit Endpunkten $a, b \in K$, die einen Punkt aus A enthält, hat ihre beiden Endpunkte in A liegen. Offensichtlich ist eine weitere äquivalente Formulierung:

$$\forall a, b \in K \quad \forall \lambda \in \mathbb{R}, 0 < \lambda < 1 : ((a \notin A \text{ oder } b \notin A) \Rightarrow \lambda a + (1 - \lambda)b \notin A),$$

geschrieben mit offenen Strecken:

$$\forall a, b \in K \text{ mit } a \notin A \text{ oder } b \notin A : (a, b) \cap A = \emptyset,$$

in Worten: Jede offene Strecke (a, b) mit Endpunkten $a, b \in K$, die nicht beide gleichzeitig in A liegen, enthält keinen Punkt aus A .

Man bemerke die folgende Kompatibilitätseigenschaft: Sei G eine Teilmenge von X , F eine extremale Teilmenge von G und weiter E eine extremale Teilmenge von F . Dann ist auch E eine extremale Teilmenge von G , denn: Betrachte eine in X liegende offene Strecke mit Endpunkten aus G . Ist nun ein Punkt dieser offenen Strecke ein Element von E , so auch von F und nach der an F gemachten Voraussetzung müssen die beiden Endpunkte in F liegen, woraus wiederum nach der an E gemachten Voraussetzung folgt, dass sie auch in E liegen müssen.

Ist eine extremale Teilmenge von K konvex, so wird sie eine *Seite* (im Englischen *face*) von K genannt. Mit diesen Definitionen ist also die leere Menge stets eine Seite von K , und für konvexes K ist auch K selbst eine Seite von K .

Ist A eine einelementige extremale Teilmenge von K , etwa $A = \{x\}$, so heißt x ein *Extremalpunkt* von K (bei manchen Autoren auch *Extrempunkt*; im Englischen analog *extremal point* bzw. *extreme point*). Die Menge der Extremalpunkte von K wird mit $\text{ex}(K)$ bezeichnet. Speziell gilt aufgrund der oben erwähnten Kompatibilitätseigenschaft: Sei G eine Teilmenge von X und F eine extremale Teilmenge von G . Dann gilt $\text{ex } F = F \cap \text{ex } G$.

Ein $x \in K$ heißt ein *reeller Extremalpunkt* (im Englischen *real extremal point*) von K , falls gilt: Ist $y \in X$ und gilt dann für alle $t \in \mathbb{R}$ mit $|t| \leq 1$ die Aussage $x + ty \in K$, so ist $y = 0$. Formuliert man diese Definition eines reellen Extremalpunktes anstatt mit $t \in \mathbb{R}$ mit $t \in \mathbb{C}$, so liegt die Definition eines *komplexen Extremalpunktes* (im Englischen *complex extremal point*) von K vor. Die Menge der reellen Extremalpunkte von K wird mit $\text{ex}_{\mathbb{R}}(K)$, die der komplexen mit $\text{ex}_{\mathbb{C}}(K)$ bezeichnet.

Offensichtlich ist jeder reelle Extremalpunkt von K ein komplexer Extremalpunkt von K , insofern $\mathbb{R} = \mathbb{C}$ vorliegt. Des Weiteren gilt: Jeder Extremalpunkt von K ist ein reeller Extremalpunkt von K , und falls K konvex ist, gilt auch die Umkehrung.

Beweis. Sei x ein Extremalpunkt von K . Also $x \in K$ und für alle $a, b \in K$ und für alle $\lambda \in \mathbb{R}$ mit $0 < \lambda < 1$ gilt die Implikation $x = b + \lambda(a - b) \Rightarrow a = b = x$. Sei $y \in X$ derart, dass für alle $t \in \mathbb{R}$ mit $|t| \leq 1$ die Aussage $x + ty \in K$ gelte. Insbesondere ist also zum Beispiel $z := x - 0.5y \in K$, also $x = z + 0.5y$. Setze $b := z$ und $a := y + b$; also $b \in K$ und $x = b + 0.5(a - b)$. Wegen $a = y + z = y + x - 0.5y = x + 0.5y$ ist auch $a \in K$. Somit ist nach Voraussetzung $x = a = b$, also $y = a - b = 0$, das heißt, x ist ein reeller Extremalpunkt von K .

Sei nun K als konvex vorausgesetzt und x ein reeller Extremalpunkt von K . Sei $\lambda \in \mathbb{R}$, $0 < \lambda < 1$ und $a, b \in K$ mit $\lambda a + (1 - \lambda)b = x$. Also $x = b + \lambda(a - b)$, $x - \lambda(a - b) = b$, $x + \lambda(b - a) = b \in K$ und $a - a + \lambda a + (1 - \lambda)b = x$, $x + (1 - \lambda)a - (1 - \lambda)b = a$, $x + (1 - \lambda)(a - b) = a$, $x + (\lambda - 1)(b - a) = a \in K$. Da K konvex ist, liegt die Strecke $x + [\lambda - 1, \lambda] \cdot (b - a)$ in K . Setze $\mu := \min\{\lambda, |\lambda - 1|\}$. Insbesondere liegt also die Strecke $x + [-\mu, \mu] \cdot (b - a)$ in K . Da nach Voraussetzung an λ das μ von Null verschieden ist, liegt somit auch die Strecke $x + [-1, 1] \cdot \frac{(b-a)}{\mu}$ in K . Nach Voraussetzung muss $\frac{(b-a)}{\mu} = 0$, also $a = b$ gelten. Somit $x = a = b$ und x ist ein Extremalpunkt von K . $\square$

Falls K konvex ist, ist ein Punkt $x \in K$ genau dann ein Extremalpunkt von K , wenn $K \setminus \{x\}$ konvex ist (HOLMES [137](1975)).

Beispiele für Banachräume X mit $\text{ex}(B_X) = \emptyset$ sind der Nullfolgenraum c_0 über $\mathbb{K}$ (siehe 2.4.19) und die Funktionenräume $L_1([0, 1])$ und $L_1(\mathbb{C})([0, 1])$ (MUKHERJEA und POTHoven [216](1986), Seite 110, Übung 6.6.36). Zugleich gilt aber für den zuletzt genannten Funktionenraum auch $\text{ex}_{\mathbb{C}}(B_X) = S_X$ (DINEEN [69](1981), Seite 161).

2.3 Algebren

2.3.1 Definition (SCHAFFER [256](1966)). Eine *nichtassoziative Algebra* A über $\mathbb{K}$ ist ein Vektorraum über $\mathbb{K}$, ausgestattet mit einer $\mathbb{K}$ -bilinearen Multiplikation $A \times A \rightarrow A$. Ein Untervektorraum einer nichtassoziativen Algebra A heißt eine *Unteralgebra* von A , falls er auch ein Unterring von A ist. Eine *Algebra* A über $\mathbb{K}$ ist eine nichtassoziative Algebra über $\mathbb{K}$, die ein Ring ist. Eine nichtassoziative Algebra A über $\mathbb{K}$, $A \neq \{0\}$, heißt eine *nichtassoziative Divisionsalgebra* über $\mathbb{K}$, wenn deren zugrunde liegender nichtassoziativer Ring ein nichtassoziativer Schiefkörper ist. Entsprechend heißt eine Algebra A über $\mathbb{K}$, $A \neq \{0\}$, eine *Divisionsalgebra* über $\mathbb{K}$, wenn deren zugrunde liegender Ring ein Schiefkörper ist.

Existiert auf dem Vektorraum einer (nichtassoziativen) Algebra A über $\mathbb{K}$ eine Halbnorm p , die die submultiplikative Bedingung

$$p(ab) \leq p(a)p(b) \quad \text{für alle } a, b \in A \quad (2.27)$$

erfüllt, dann heißt solch eine Halbnorm eine *Algebrahalbnorm* und solch eine (nichtassoziative) Algebra *halbnormierbar*; ist dabei p eine Norm, so heißt p eine *Algebranorm* und die (nichtassoziative) Algebra *normierbar*. (Bezüglich des Falles, dass für alle $a, b \in A$ die Gleichung $\|ab\| = \|a\| \|b\|$ gilt, siehe 3.3.2.) Eine (nichtassoziative) Algebra A über $\mathbb{K}$ zusammen mit einer Algebranorm auf A heißt eine *normierte (nichtassoziative) Algebra*. Eine *unitale (nichtassoziative) Algebra* über $\mathbb{K}$ ist eine normierte (nichtassoziative) Algebra über $\mathbb{K}$ mit Eins e und $\|e\| = 1$. Eine *(nichtassoziative) Banachalgebra* über $\mathbb{K}$ ist eine normierte (nichtassoziative) Algebra über $\mathbb{K}$, deren Norm vollständig ist (Definition 2.4.6). Als Erweiterung von 2.1.34 heißt eine nichtassoziative normierte Algebra A über $\mathbb{K}$ *einfach*, falls $\{0\}$ und A die einzigen abgeschlossenen Ideale von A sind.

Sei $(A, +, \cdot)$ eine nichtassoziative Algebra über $\mathbb{K}$ und seien M und N zwei nicht leere Teilmengen von A . Dann bezeichnet MN in Fortsetzung der Definition 2.1.13 den Untervektorraum $\text{span}(M \cdot N)$ von A . Entsprechend ist dann M^n für $n \in \mathbb{N}$, $n \geq 2$, definiert — der jeweilige Kontext wird dabei eine Verwechslung mit dem kartesischen Produkt M^n ausschließen.

2.3.2 Bemerkungen. (a) Eine nichtassoziative Algebra A über $\mathbb{K}$ ist also genau ein nichtassoziativer Ring, der auch ein Vektorraum über $\mathbb{K}$ mit derselben Addition ist, so dass gilt

$$\lambda(ab) = (\lambda a)b = a(\lambda b) \quad \text{für alle } \lambda \in \mathbb{K}, a, b \in A.$$

Ein A ist genau dann eine nichtassoziative Algebra über $\mathbb{K}$, wenn A ein nichtassoziativer Ring mit dem Ring $\mathbb{K}$ ist, siehe 2.1.28. Obwohl der Idealbegriff für nichtassoziative Algebren über $\mathbb{K}$ bereits in 2.1.28 enthalten ist, sei er der Deutlichkeit halber im Folgenden explizit ausformuliert:

Sei A eine nichtassoziative Algebra über $\mathbb{K}$. Ein *Links-* bzw. *Rechtsideal* bzw. *beidseitiges Ideal* von A oder im letzten Fall einfach nur *Ideal* von A ist eine Unteralgebra von A , die die entsprechende Idealeigenschaft bezüglich des A zugrunde liegenden nichtassoziativen Ringes aufweist.

Und nochmals mit anderen Worten:

Sei A eine nichtassoziative Algebra über $\mathbb{K}$. Ein *Ring-Linksideal* $\mathfrak{a}$ von A — soll heißen, ein Linksideal des A zugrunde liegenden nichtassoziativen Ringes — ist genau dann ein *Algebra-Linksideal* von A — soll entsprechend ein Linksideal der nichtassoziativen Algebra A meinen —, wenn es bezüglich der Multiplikation des Skalarenkörpers $\mathbb{K}$ abgeschlossen ist, wenn also $\mathbb{K} \cdot \mathfrak{a} \subseteq \mathfrak{a}$ gilt. Für Rechtsideale und Ideale entsprechend.

(b) (NAIMARK [221](1972), Seite 174). Sei A eine normierte nichtassoziative Algebra über $\mathbb{K}$. Dann ist die Multiplikation $A \times A \rightarrow A$, $(x, y) \mapsto xy$ stetig. Denn: Für alle x ,

$y, x_0, y_0 \in A$ gilt: $\|xy - x_0y_0\| = \|xy - xy_0 - x_0y + x_0y_0 + xy_0 - x_0y_0 + x_0y - x_0y_0\| = \|(x-x_0)(y-y_0) + (x-x_0)y_0 + x_0(y-y_0)\| \leq \|x-x_0\|\|y-y_0\| + \|x-x_0\|\|y_0\| + \|x_0\|\|y-y_0\|$.

(c) Auf die Definition 2.4.6 vorgehend gilt: Letztlich äquivalent zu der Definition 2.3.1 kann man eine Banachalgebra auch definieren als eine Algebra mit einer Banachraumtopologie, so dass die Multiplikation $A \times A \rightarrow A$ (getrennt oder gemeinsam) stetig ist, da dann eine zu der vorgegebenen Banachraumnorm äquivalente Algebranorm auf A existiert. Siehe UPMEIER [279](1985), Seiten 22 und 24, PALMER [231](1994), Seite 15, ŻELASKO [297](1973), Seiten 8-11, und LARSEN [196](1973), Seite 23.

(d) Ist p ein schiefminimal idempotentes Element einer Algebra A über $\mathbb{K}$, so ist pAp eine Divisionsalgebra über $\mathbb{K}$. Siehe auch 3.3.3.

(e) Ein Untervektorraum U einer nichtassoziativen Algebra A ist genau dann eine Unter algebra von A , wenn $U \cdot U \subseteq U$ gilt.

2.3.3 Algebra-Radikale. Genauso wie in 2.1.49 für nichtassoziative K -Ringe durchgeführt, erklärt man den Radikal-Begriff für nichtassoziative Algebren über einen Körper $\mathbb{K}$. Da jede nichtassoziative Algebra über einen Körper $\mathbb{K}$ ein nichtassoziativer Ring mit Operatorenbereich $\mathbb{K}$ ist, gilt für diesen nichtassoziativen $\mathbb{K}$ -Ring auch das ebenfalls in 2.1.49 aufgeführte Theorem. Da genau die Ideale einer nichtassoziativen Algebra über einen Körper $\mathbb{K}$ die $\mathbb{K}$ -zulässigen Ideale des der nichtassoziativen Algebra zugrunde liegenden nichtassoziativen Ringes sind, formuliert sich das besagte Theorem für nichtassoziative Algebren über einen Körper $\mathbb{K}$ wie folgt:

Sei A eine nichtassoziative Algebra über einen Körper $\mathbb{K}$ und bezeichne R den A zugrunde liegenden nichtassoziativen Ring. Ist dann $\mathcal{S}$ eine Radikal-Eigenschaft von R , so existiert das $\mathcal{S}$ -Radikal von A und stimmt überein mit dem $\mathcal{S}$ -Radikal von R .

2.3.4 Algebra-Ordnungen. Alle in 2.1.50 erklärten Begriffe für nichtassoziative Ringe, wie etwa artinsch, noethersch, minimales und maximales Ideal, werden für nichtassoziative Algebren über $\mathbb{K}$ ganz entsprechend definiert; dabei ist nur darauf zu achten, dass nur die ein- und beidseitigen Ideale bezüglich der algebraischen Struktur der nichtassoziativen Algebren über $\mathbb{K}$ in Betracht gezogen werden dürfen.

2.3.5. Da die ein- und beidseitigen Ideale von nichtassoziativen Algebren über $\mathbb{K}$ Vektorräume über $\mathbb{K}$ sind, erfüllt jede endlichdimensionale nichtassoziative Algebra über $\mathbb{K}$ sowohl die DCC als auch die ACC auf ein- und beidseitigen Idealen. Das heißt, jede endlichdimensionale nichtassoziative Algebra über $\mathbb{K}$ ist linksartinsch, rechtsartinsch, linksnoethersch und rechtsnoethersch.

2.3.6 (GRAY [120](1970), Seite 39, Theorem 16). Sei $n \in \mathbb{N}^\times$, D ein Schiefkörper und dann R der Ring der $n \times n$ -Matrizen mit Einträgen aus D . Dann ist R ein einfacher, linksartinscher Ring mit $R^2 \neq \{0\}$.

2.3.7 Definition. Eine (Banach-)Lie-Algebra über $\mathbb{K}$ ist eine nichtassoziative (Banach-)Algebra A über $\mathbb{K}$, deren Multiplikation die folgenden beiden Bedingungen erfüllt:

(a) $x^2 = 0$ für alle $x \in A$.

(b) $(xy)z + (yz)x + (zx)y = 0$ für alle $x, y, z \in A$. (Jacobi-Gleichung)

2.3.8. Die Bedingung $x^2 = 0$ für alle $x \in A$ in der Definition 2.3.7 einer Lie-Algebra A zieht die *Anti-Kommutativität*, das heißt $xy = -yx$ für alle $x, y \in A$, des Produktes nach sich: $0 = (x+y)^2 = x^2 + xy + yx + y^2 = xy + yx$, $x, y \in A$. Da hier keine Algebren über Körpern betrachtet werden, deren Charakteristik 2 ist, gilt auch die Umkehrung, das heißt, die Bedingung $x^2 = 0$ für alle $x \in A$ ist äquivalent zur Anti-Kommutativität von A . (JACOBSON [152](1974), Seite 413.)

2.3.9. Jede Algebra A über $\mathbb{K}$, ausgestattet mit dem Kommutator als Produkt, ist eine Lie-Algebra. Das sieht man, indem man die folgenden Gleichungen addiert und dann durch Verwenden der Assoziativität von A sich paarweise Terme aufheben lässt.

$$\begin{aligned} [[x, y], z] &= [xy - yx, z] = (xy)z - (yx)z - z(xy) + z(yx), \\ [[y, z], x] &= [yz - zy, x] = (yz)x - (zy)x - x(yz) + x(zy), \\ [[z, x], y] &= [zx - xz, y] = (zx)y - (xz)y - y(zx) + y(xz). \end{aligned}$$

Diese Lie-Algebra heißt die zu A *assoziierte Lie-Algebra* und wird mit $A^\ominus$ bezeichnet.

(McCRIMMON [210](2004), Seite 11). Zu jeder Lie-Algebra A gibt es eine Algebra B , so dass A ein unter dem Kommutatorprodukt abgeschlossener Unterraum von $B^\ominus$ ist.

2.3.10 Beispiel (HALL [123](2003), Seite 54). Sei $X = \mathbb{R}^3$ versehen mit dem Kreuz-

produkt $\begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \times \begin{pmatrix} y_1 \\ y_2 \\ y_3 \end{pmatrix} = \begin{pmatrix} x_2 y_3 - x_3 y_2 \\ x_3 y_1 - x_1 y_3 \\ x_1 y_2 - x_2 y_1 \end{pmatrix}$. Dann ist X eine Lie-Algebra. Man

beachte, dass sie ein Beispiel für eine Lie-Algebra ist, auf der es kein Produkt $(x, y) \mapsto xy$ gibt mit $x \times y = xy - yx$ für alle $x, y \in X$. X ist nicht assoziativ.

2.3.11. Sei X eine nichtassoziative Algebra über $\mathbb{K}$, deren Produkt antikommutativ ist (siehe 2.3.8), das heißt, es gilt $xy = -yx$ für alle $x, y \in X$. Bezeichne $L_x \in L(X)$ die durch $x \in X$ bestimmte Links-Multiplikation in X (Definition 2.1.18). Dann lässt sich die Jacobi-Gleichung in $L(X)^\ominus$ schreiben als

$$L_{xy} = L_x L_y \quad \text{für alle } x, y \in X.$$

(In $L(X)$ schreibt sie sich also als $L_{xy} = [L_x, L_y]$ für alle $x, y \in X$.)

Beweis. Sei $x, y, z \in X$. Dann gilt in $L(X)$: $(L_{xy} - [L_x, L_y])(z) = (xy)z - (L_x L_y - L_y L_x)(z) = (xy)z - x(yz) + y(xz) = (xy)z + (yz)x + (zx)y$, also die Behauptung. $\square$

Mit der Rechts-Multiplikation in X schreibt sich die Jacobi-Gleichung in $L(X)^\ominus$ fast analog als

$$R_{yx} = R_x R_y \quad \text{für alle } x, y \in X.$$

(In $L(X)$ schreibt sie sich also als $R_{yx} = [R_x, R_y]$ für alle $x, y \in X$.)

Beweis. Sei $x, y, z \in X$. Dann gilt in $L(X)$: $(R_{yx} - [R_x, R_y])(z) = z(yx) - (R_x R_y - R_y R_x)(z) = z(yx) - (zy)x + (zx)y = -z(xy) + (yz)x + (zx)y = (xy)z + (yz)x + (zx)y$, also die Behauptung. $\square$

2.3.12 Definition (UPMEIER [279](1985), Seite 63). Sei A eine Lie-Algebra über $\mathbb{K}$. Sei $S \subseteq \mathbb{C}$. Eine direkte Summen-Zerlegung $A = \bigoplus_{n \in S} A_n$ heißt eine *additive Graduierung* von A , wenn $A_m \cdot A_n \subseteq A_{m+n}$ für alle $m, n \in S$ gilt, wobei $A_n := \{0\}$ für $n \in \mathbb{C} \setminus S$ gesetzt wird.

Sei $S \subseteq \mathbb{C}^\times$. Eine direkte Summen-Zerlegung $A = \bigoplus_{n \in S} A_n$ heißt eine *multiplikative Graduierung* von A , wenn $A_m \cdot A_n \subseteq A_{mn}$ für alle $m, n \in S$ gilt, wobei $A_n := \{0\}$ für $n \in \mathbb{C}^\times \setminus S$ gesetzt wird.

2.3.13 Definition (SCHAFER [256](1966), Seite 27). Eine *alternative Algebra* über $\mathbb{K}$ ist eine nichtassoziative Algebra A über $\mathbb{K}$ mit $(x, x, y) = (y, x, x) = 0$ für alle $x, y \in A$, die Klammer den Assoziator von Definition 2.1.30 bezeichnend.

2.3.14 Peirce-Zerlegung III (SCHAFFER [256](1966), Seite 32). Sei A eine alternative Algebra über $\mathbb{K}$ und $P = \{p_1, \dots, p_n\}$ eine endliche Menge von paarweise orthogonalen Idempotenten aus A . Genau wie für Ringe (siehe 2.1.46) hat man dann die Peirce-Zerlegung von A bezüglich P . Konkret setze man dazu, dabei das Kronecker-Delta verwendend,

$$U_{ij} := \{x \in A : p_k x = \delta_{ki} x, x p_k = \delta_{jk} x \text{ für alle } k \in \{1, \dots, n\}\} \quad (2.28)$$

für alle $i, j \in \{0, 1, \dots, n\}$. Für A , betrachtet als Vektorraum über $\mathbb{K}$, sind dann alle diese U_{ij} Untervektorräume von A und der Vektorraum A ist gleich der inneren direkten Summe

$$A = \sum_{i,j=0}^n \oplus U_{ij}.$$

Manchmal ist es praktisch, die Elemente von U_{ij} mit x_{ij} zu bezeichnen. Insbesondere kann man mittels dieser Bezeichnung jedes $x \in A$ als eine Matrix $(x_{ij})_{i,j=0}^n$ auffassen.

Bezeichnet L_i (bzw. R_i) die durch p_i bestimmte Links-Multiplikation (bzw. Rechts-Multiplikation), so kann man offensichtlich jedes U_{ij} als Durchschnitt von gewissen Eigenräumen ansehen:

$$U_{ij} = \bigcap_{k=1}^n \text{Eig}(L_k, \delta_{ki}) \cap \text{Eig}(R_k, \delta_{jk})$$

für alle $i, j \in \{0, 1, \dots, n\}$.

Die Zerlegung (2.28) kann man als eine Verfeinerung der Peirce-Zerlegung von A bezüglich des Idempotenten $p := p_1 + \dots + p_n$ auffassen: Für alle $x \in A$ gilt

$$\begin{aligned} x &= \sum_{i,j=0}^n x_{ij} \\ &= \sum_{i,j=1}^n p_i x p_j + \sum_{i=1}^n p_i x (1-p) + \sum_{j=1}^n (1-p) x p_j + (1-p) x (1-p), \end{aligned}$$

wobei für alle $i, j \in \{1, \dots, n\}$ gilt:

$$p_i x p_j \in U_{ij}, p_i x (1-p) \in U_{i0}, (1-p) x p_j \in U_{0j} \text{ und } (1-p) x (1-p) \in U_{00};$$

es sei an Gleichung (2.19) aus 2.1.33 erinnert.

Im Folgenden sind einige Eigenschaften der U_{ij} aufgelistet:

$$\begin{aligned} U_{ij} U_{jk} &\subseteq U_{ik} && \text{für alle } i, j, k \in \{0, 1, \dots, n\}, \\ U_{ij} U_{ij} &\subseteq U_{ji} && \text{für alle } i, j \in \{0, 1, \dots, n\}, \\ U_{ij} U_{kl} &\subseteq \{0\} && \text{für alle } i, j, k, \ell \in \{0, 1, \dots, n\} \text{ mit } j \neq k, (i, j) \neq (k, \ell), \\ x_{ij}^2 &= 0 && \text{für alle } x_{ij} \in U_{ij} \text{ und } i, j \in \{0, 1, \dots, n\} \text{ mit } i \neq j. \end{aligned}$$

Liege nun der Fall vor, dass A assoziativ ist, also eine Algebra über $\mathbb{K}$ ist. Dann ist der Assoziator identisch Null und damit entnimmt man dem Beweis in SCHAFFER, ebd., Seite 36, die Gültigkeit der Inklusion

$$U_{ij} U_{kl} \subseteq \delta_{jk} U_{il} \quad \text{für alle } i, j, k, \ell \in \{0, 1, \dots, n\}, \quad (2.29)$$

die sogenannten *assoziativen Peirce-Multiplikationsregeln* (MCCRIMMON [210] (2004), Seite 241). Somit kann man sich die Multiplikation in assoziativem A als eine Matrix-Multiplikation vorstellen: Sind $x = (x_{ij})$ und $y = (y_{ij})$ zwei Elemente von A , so gilt

$$xy = \left(\sum_{k=0}^n x_{ik} y_{kj} \right)_{i,j=0}^n.$$

Es sei auf eine Gleichung hingewiesen, die ähnlich der Inklusion (2.29) ist. Bezeichne dazu $\text{Mat}_{\mathbb{K}}(n)$ die Algebra der $n \times n$ -Matrizen über $\mathbb{K}$. Des Weiteren bezeichne für $i, j = 1, \dots, n$ mit E_{ij} die Matrix aus $\text{Mat}_{\mathbb{K}}(n)$, deren sämtliche Koeffizienten Null sind mit Ausnahme des Koeffizienten in der i -ten Zeile und der j -ten Spalte, der gleich 1 ist. Dann ist $E_{ij}, (i, j) \in \{1, \dots, n\}^2$, eine Basis (die sogenannte Standardbasis) des Vektorraumes $\text{Mat}_{\mathbb{K}}(n)$ über $\mathbb{K}$ und es gilt

$$E_{ij}E_{kl} = \delta_{jk}E_{il} \quad \text{für alle } i, j, k, l \in \{1, \dots, n\}.$$

2.3.15 Definition. Eine *Jordan-Algebra* über $\mathbb{K}$ ist eine kommutative nichtassoziative Algebra A über $\mathbb{K}$ mit

$$(x, y, x^2) = 0 \quad \text{für alle } x, y \in A, \quad (2.30)$$

wobei die Klammer der mit Definition 2.1.30 eingeführte Assoziator ist. Eine Jordan-Algebra, die zugleich eine nichtassoziative Banachalgebra ist, heißt eine *Banach-Jordan-Algebra*. Sei A eine Jordan-Algebra über $\mathbb{K}$ mit Eins. Ein $x \in A$ heißt *invertierbar*, wenn ein $y \in A$ existiert mit $xy = e$ und $x^2y = x$. Eine *Divisions-Jordan-Algebra* über $\mathbb{K}$ (im Englischen *division Jordan algebra*) ist eine von $\{0\}$ verschiedene, unitale Jordan-Algebra über $\mathbb{K}$, deren alle von Null verschiedenen Elemente invertierbar sind.

2.3.16. Benannt sind die Jordan-Algebren nach Pascual Jordan.

Jede Algebra A über $\mathbb{K}$, ausgestattet mit dem Produkt $x \bullet y := \frac{1}{2}(xy + yx)$ für alle $x, y \in A$ (das sogenannte *Anti-Kommutator-Produkt*), ist eine Jordan-Algebra über $\mathbb{K}$ und wird mit $A^{\oplus}$ bezeichnet. (Für allgemeine Betrachtungen, in denen $\mathbb{K}$ ein Körper mit Charakteristik 2 sein kann, lässt man in der Produkt-Definition den Vorfaktor $\frac{1}{2}$ weg.) Siehe auch Definition 3.2.14. Eine nichtassoziative Algebra B über $\mathbb{K}$ heißt eine *spezielle Jordan-Algebra* (im Englischen *special Jordan algebra*) über $\mathbb{K}$, wenn es eine Algebra A über $\mathbb{K}$ gibt, so dass B isomorph zu einer Unteralgebra von $A^{\oplus}$ ist. Eine Jordan-Algebra über $\mathbb{K}$ heißt *exzeptionell* (im Englischen *exceptional Jordan algebra*), wenn sie nicht speziell ist.

In jeder alternativen Algebra A über $\mathbb{K}$ gilt Gleichung (2.30) aus 2.3.15 (BRAUN und KOECHER [38](1966), Seite 209). Für jede alternative Algebra A über $\mathbb{K}$ ist $A^{\oplus}$ eine spezielle Jordan-Algebra über $\mathbb{K}$ (ebd., Seite 215, Satz 2.7).

2.3.17 Peirce-Zerlegung IV (JACOBSON [151](1968), Seite 120). Sei J eine Jordan-Algebra über $\mathbb{K}$ mit Eins. Sei p ein idempotentes Element von J . Bezeichne R_p die durch p bestimmte Rechtsmultiplikation in J . Dann gilt

$$2R_p^3 - 3R_p^2 + R_p = R_p(2R_p - 1)(R_p - 1) = 0.$$

Seien $p_1, \dots, p_n$ paarweise orthogonale Idempotente von J mit $e = p_1 + \dots + p_n$. Bezeichne L_i die durch p_i bestimmte Linksmultiplikation, $i = 1, \dots, n$. Dann hat man die Peirce-Zerlegung

$$J = \sum_{i \leq j} \oplus J_{ij}$$

der nichtassoziativen Algebra J in eine innere direkte Summe von Peirce-Unterräumen, wobei

$$J_{ii} := \text{Eig}(L_i, 1), \quad i = 1, \dots, n,$$

die sogenannten *diagonalen Peirce-Unteralgebren* sind, und

$$J_{ij} = J_{ji} := \text{Eig}\left(L_i, \frac{1}{2}\right) \cap \text{Eig}\left(L_j, \frac{1}{2}\right), \quad i, j = 1, \dots, n \text{ mit } i \neq j,$$

die *nicht-diagonalen* (im Englischen *off-diagonal*) *Peirce-Räume* sind.

Es gelten die folgenden Rechenregeln:

$$\begin{array}{ll} J_{ii}J_{ii} \subseteq J_{ii}, & J_{ij}J_{ii} \subseteq J_{ij}, \\ J_{ij}J_{ij} \subseteq J_{ii} + J_{jj}, & J_{ii}J_{jj} = \{0\}, \\ J_{ij}J_{jk} \subseteq J_{ik}, & J_{ij}J_{kk} = \{0\}, \\ J_{ij}J_{kl} = \{0\}. \end{array}$$

für alle paarweise verschiedenen $i, j, k, l \in \{1, \dots, n\}$.

Sei A eine Algebra über $\mathbb{K}$ und p ein idempotentes Element von A . Setze $J := A^\oplus$ (siehe 2.3.16), $p_1 := p$ und $p_2 := 1 - p$. (Es sei an Gleichung (2.19) aus 2.1.33 erinnert.) Dann gilt (siehe MCCRIMMON [210](2004), Seite 241):

$$J_{11} = pAp, \quad J_{12} = J_{21} = pA(1-p) + (1-p)Ap, \quad J_{22} = (1-p)A(1-p).$$

2.3.18 Approximative Eins (PALMER [231](1994), Seite 520). Sei A eine normierte Algebra über $\mathbb{K}$. Eine *approximative Linkseins* (im Englischen *left approximate identity*) von A ist ein Netz (a_ν) in A , so dass für jedes $x \in A$ das Netz $(\|a_\nu x - x\|)$ gegen Null konvergiert, sprich, das Netz $(a_\nu x)$ in Norm gegen x konvergiert. *Approximative Rechtseins* entsprechend mit (xa_ν) . Eine *approximative Eins* von A ist ein Netz in A , das sowohl eine approximative Links- als auch eine approximative Rechtseins ist. Eine approximative Linkseins (a_ν) von A heißt *beschränkt durch* $\lambda \in \mathbb{R}^+$, wenn $\sup_\nu \|a_\nu\| \leq \lambda$ gilt, und *beschränkt*, wenn sie durch ein $\lambda \in \mathbb{R}^+$ beschränkt ist. Approximative Rechtseins und approximative Eins entsprechend. Man bemerke, dass im Fall von $A \neq \{0\}$ für eine durch ein $\lambda \in \mathbb{R}^+$ beschränkte approximative Linkseins stets $\lambda \geq 1$ gilt; des Weiteren existiert dann für jedes $x \in A$ und jedes $\varepsilon > 0$ ein $y \in A$ mit $\|y\| < \lambda$ mit $\|yx - x\| < \varepsilon$; für approximative Rechtseins und approximative Eins entsprechend.

Beweis. Sei $A \neq \{0\}$ und (a_ν) eine durch ein $\lambda \in \mathbb{R}^+$ beschränkte approximative Linkseins von A . Zuerst bemerke man, dass die Ungleichung $\|a_\nu x - x\| < \varepsilon$ äquivalent ist zu der Ungleichungskette $\|x\| - \varepsilon < \|a_\nu x\| < \|x\| + \varepsilon$. Für beliebig (kleines) $\varepsilon > 0$ und für $x \neq 0$ gilt somit $1 - \frac{\varepsilon}{\|x\|} < \|a_\nu\|$, also $\lambda \geq 1$. Des Weiteren sieht man, dass stets ein μ mit $0 < \mu < 1$ existiert, so dass mit $\|a_\nu x - x\| < \varepsilon$ auch $\|(\mu a_\nu)x - x\| < \varepsilon$ gilt; setze also nur noch $y := \mu a_\nu$. $\square$

2.3.19 Bemerkung und Definition (SCHAFFER [256](1966), Seite 11. Erweiterung einer nichtassoziativen Algebra durch Anknüpfen des Skalarbereiches). Es sei an 2.1.40 und 2.2.18(a) erinnert. Sei A eine nichtassoziative Algebra über $\mathbb{K}$. Durch Ausstatten des Vektorraumes A^1 mit dem Produkt

$$(a, \eta)(b, \mu) := (ab + \eta b + \mu a, \eta\mu) \quad \text{für alle } a, b \in A, \eta, \mu \in \mathbb{K},$$

liegt eine nichtassoziative Algebra über $\mathbb{K}$ mit $(0, 1)$ als Eins vor, die wieder mit A^1 bezeichnet wird und assoziativ ist, wenn A assoziativ ist.

2.3.20 Bemerkungen. (a) (SCHAFFER [256](1966), Seite 11). Die Menge $\{(a, 0) : a \in A\}$ ist ein Ideal von A^1 .

(b) (RICKART [246](1960), Seite 2). Ist A eine normierte Algebra über $\mathbb{K}$, dann ist A^1 gemäß den Bemerkungen und Definitionen 2.2.18 und 2.3.19 eine unitale Algebra über $\mathbb{K}$ und die Abbildung $a \mapsto (a, 0)$ ist eine isometrische Isomorphie von A auf eine Unteralgebra von A^1 .

Für die Bemerkung 3.7.15 wird die folgende Bemerkung gebraucht werden.

2.3.21 Bemerkung. Sei A eine nichtassoziative Algebra über $\mathbb{K}$ mit einer Eins e . Dann gilt $A = \{a \in A^1 : ae = a\} = \{a \in A^1 : ea = a\} = \{a \in A^1 : eae = a\}$. A ist also die Fixpunktmenge der zu der Eins von A bestimmten Links- und Rechts-Multiplikation in A^1 .

Beweis. Es genügt das erste Gleichheitszeichen zu zeigen, da die beiden anderen analog bewiesen werden. „ $\subseteq$ “ ist klar. Zu „ $\supseteq$ “: Sei $(a, \eta) \in A^1$ mit $(a, \eta)e = (a, \eta)$. Also $(a, \eta) = (ae + \eta e + a \cdot 0, \eta \cdot 0) = (ae + \eta e, 0)$, $\eta = 0$. $\square$

2.3.22 Definition (PALMER [231](1994), Seite 19). Sei $\|\cdot\|$ eine Algebranorm auf einer Algebra A über $\mathbb{K}$.

- (a) Gilt $\|a\| = \sup\{\|ab\|, \|ba\| : b \in B_A\}$ für alle $a \in A$, dann heißt $\|\cdot\|$ eine *reguläre Norm*.
- (b) Hat A keine Eins und ist $\|\cdot\|$ regulär, dann sei die mit $\|\cdot\|_R$ bezeichnete Norm auf A^1 wie folgt definiert:

$$\|(a, \lambda)\|_R := \sup\{\|ab + \lambda b\|, \|ba + \lambda b\| : b \in B_A\} \quad \text{für alle } a \in A, \lambda \in \mathbb{K}.$$

2.3.23 Bemerkung (PALMER [231](1994), Seite 19). Betrachtet man in der Situation von Definition 2.3.22 die Abbildung

$$p: A^1 \rightarrow \mathbb{K}, (a, \lambda) \mapsto \sup\{\|ab + \lambda b\| : b \in B_A\}$$

und ist ein $a \in A$ zum Beispiel zwar eine Linkseins, aber keine Rechtseins in A , dann ist $p(a, -1) = 0$, so dass p keine Norm ist.

2.3.24 Satz (PALMER [231](1994), Seite 20). Sei $\|\cdot\|$ eine Algebranorm auf einer Algebra A über $\mathbb{K}$.

- (a) Wenn A eine Eins hat, dann ist $\|\cdot\|$ genau dann regulär, wenn A unital ist.
- (b) Wenn A keine Eins hat, dann ist die Norm von Definition 2.2.18 (b), $(a, \lambda) \mapsto \|a\| + |\lambda|$, $a \in A, \lambda \in \mathbb{K}$, eine reguläre Norm auf A^1 , welche die Norm $\|\cdot\|$ fortsetzt. Für jede reguläre Norm $\|\cdot\|_r$ auf A^1 , welche die Norm $\|\cdot\|$ fortsetzt, gilt:

$$\|(a, \lambda)\|_r \leq \|a\| + |\lambda| \quad \text{für alle } a \in A, \lambda \in \mathbb{K}.$$

- (c) Wenn A keine Eins hat und $\|\cdot\|$ regulär ist, dann ist $\|\cdot\|_R$ eine reguläre Norm auf A^1 , welche die Norm $\|\cdot\|$ fortsetzt. Für jede reguläre Norm $\|\cdot\|_r$ auf A^1 , welche die Norm $\|\cdot\|$ fortsetzt, gilt:

$$\|(a, \lambda)\|_R \leq \|(a, \lambda)\|_r \leq \|a\| + |\lambda| \quad \text{für alle } a \in A, \lambda \in \mathbb{K}.$$

Ist $\mathbb{K} = \mathbb{C}$ und die Norm $\|\cdot\|$ vollständig (Definition 2.4.6), dann gilt zusätzlich:

$$\|a\| + |\lambda| \leq (1 + 2 \exp(1)) \|(a, \lambda)\|_R \quad \text{für alle } a \in A, \lambda \in \mathbb{K}.$$

2.3.25 Annihilator und Kommutante. Sei A eine nichtassoziative Algebra über $\mathbb{K}$. Da der Kommutator bilinear ist, kann man das orthosymmetrische Bilinearsystem $(A, A; [\cdot, \cdot]; A)$ betrachten. Für $M \subseteq A$ gilt dann bezüglich dieses Bilinearsystems $M^\perp = \{x \in A : [m, x] = 0 \text{ für alle } m \in M\} = \{x \in A : mx = xm \text{ für alle } m \in M\}$, das heißt,

$$M^\perp = M'.$$

Dadurch sind insbesondere die Aussagen von 2.1.4 und 2.2.27 direkt auf die Kommutante übertragbar, siehe den folgenden Satz 2.3.26. Ist A eine Algebra, so ist mit der Bezeichnung von 2.1.4 die Algebra $(A; [\cdot, \cdot])$ offensichtlich gerade genau die zu A assoziierte Lie-Algebra $A^\ominus$; salopp formuliert, ist dann genau das, was in A kommutativ ist, in $A^\ominus$ orthogonal.

Nebenbei sei angemerkt, dass das Bilinearsystem $(A, A; [\cdot, \cdot]; A)$ genau dann trennend ist, wenn nur das Nullelement von A mit allen Elementen von A kommutiert; des Weiteren ist es zwar stets positiv, aber für $A \neq \{0\}$ weder symmetrisch noch positiv definit.

2.3.26 Satz (PALMER [231](1994), Seite 9; MATHIEU [208](1998), Seite 192). *Sei A eine Algebra über $\mathbb{K}$. Seien S und T Teilmengen von A ; ebenso seien für alle $\nu \in I$, I eine beliebige Indexmenge, S_ν Teilmengen von A . Dann gilt:*

- (a) S' ist eine Unteralgebra über $\mathbb{K}$ von A .
- (b) $S \subseteq S''$.
- (c) $S \subseteq T \Rightarrow T' \subseteq S'$.
- (d) $S \subseteq S' \Leftrightarrow S$ ist kommutativ $\Leftrightarrow S \subseteq S'' \subseteq S'$.
- (e) $S' = S'''$.
- (f) S ist unter der $\subseteq$ -Ordnung genau dann eine maximal kommutative Teilmenge von A , wenn $S = S'$ gilt. Insbesondere kann S nur dann maximal kommutativ sein, wenn $S = S''$ gilt.
- (g) Jede kommutative Teilmenge von A ist in einer maximal kommutativen Unteralgebra enthalten.
- (h) Hat A eine Eins, dann ist diese in S' enthalten. Insbesondere folgt damit, dass mit jedem invertierbaren $x \in S'$ auch $x^{-1} \in S'$ ist.
- (i) Hat A eine Eins e , dann gilt $S' = (S + \mathbb{K}e)'$.
- (j) $S' \setminus (X') \subseteq A \setminus ((A \setminus S)')$.
- (k) $S' = \bigcap_{s \in S} \{s\}' = \bigcap_{s \in S} \ker(L_s - R_s)$ (siehe Definition 2.1.18),
also $\left(\bigcup_{\nu \in I} S_\nu \right)' = \bigcap_{\nu \in I} (S_\nu')$, speziell $(S \cup T)' = S' \cap T'$.
- (l) $S' = (\text{span } S)'$.
- (m) $S' \cup T' \subseteq (S \cap T)'$.

Beweis. (a) ist klar. Gemäß 2.3.25 sind (b), (c), (e), (j), (k) ein Spezialfall von 2.1.4 und (l) ist ein Spezialfall der Gleichung (2.26) aus 2.2.27.

(d): Erste Äquivalenz: $S \subseteq S' \Leftrightarrow \forall x \in S \forall y \in S : xy = yx \Leftrightarrow S$ ist kommutativ.
Zweite Äquivalenz: $S \subseteq S' \Rightarrow S'' \subseteq S'$.

(h): Für alle $s \in S$ gilt nach Definition der Eins: $s = es = se$.

(i): $a \in (S + \mathbb{K}e)' \Leftrightarrow \forall s \in S, \lambda \in \mathbb{K} : a(s + \lambda e) = (s + \lambda e)a \Leftrightarrow \forall s \in S, \lambda \in \mathbb{K} : as + \lambda a = sa + \lambda a \Leftrightarrow \forall s \in S : as = sa \Leftrightarrow a \in S'$. $\square$

2.3.27. Im Hinblick auf Satz 2.3.26(a) wird die Kommutante S' einer Teilmenge S einer Algebra von manchen Autoren auch die *kommutierende Algebra* von S genannt; man beachte aber 3.2.38.

2.3.28 Bezeichnung. Ist X ein Vektorraum über $\mathbb{K}$, dann ist $L(X)$ mit der Komposition als Multiplikation eine Algebra über $\mathbb{K}$. Wenn nicht ausdrücklich etwas anderes vereinbart ist, ist mit $L(X)$ immer diese Algebra gemeint. Die idempotenten Elemente von $L(X)$ heißen *Projektionen*. Ist $p \in L(X)$ eine Projektion, so wird die mit $p^\perp$ bezeichnete Abbildung $\text{Id}_X - p \in L(X)$ die zu p *komplementäre Projektion* genannt. Ein zentral idempotentes Element von $L(X)$ heißt *zentrale Projektion*.

2.3.29 (NAIMARK [221](1972), Seite 164). Sei X ein Vektorraum über $\mathbb{K}$. Ist $A \subseteq L(X)$ eine von $\{0\}$ verschiedene Algebra, so dass für keinen von $\{0\}$ und X verschiedenen Untervektorraum U von X für alle $T \in A$ die Inklusion $T(U) \subseteq U$ gilt, so ist A Jacobson-halbeinfach. (Naimark nennt solche Algebren *irreduzibel*.)

2.3.30 Definition (JACOBSON [152](1974), Seite 413). Sei A eine nichtassoziative Algebra über $\mathbb{K}$. Eine *Derivation* von A ist ein Element $\delta \in L(A)$ mit $\delta(xy) = (\delta x)y + x(\delta y)$ für alle $x, y \in A$. Die Menge aller Derivationen von A wird mit $\text{Der}(A)$ bezeichnet.

2.3.31. Sei A eine nichtassoziative Algebra über $\mathbb{K}$. $\text{Der}(A)$ ist eine Unter algebra der Lie-Algebra $(L(A))^\ominus$ und heißt die *Derivationen-Algebra von A* (JACOBSON [152](1974), Seite 414 und SCHAFER [256](1966), Seite 4).

Falls A eine Banachalgebra über $\mathbb{K}$ ist, so ist $\text{Der}(A)$ eine abgeschlossene Unter algebra von $(L(A))^\ominus$ und somit eine Banach-Lie-Algebra über $\mathbb{K}$ (UPMEIER [279](1985), Seite 33).

Sei A eine Lie-Algebra über $\mathbb{K}$ und bezeichne $[\cdot, \cdot]$ das Produkt auf A . Mit ad wird die Abbildung $A \rightarrow \text{Der}(A)$, $ad(x)y := [x, y]$ für alle $x, y \in A$, bezeichnet; sie heißt die *adjungierte Abbildung von A* und ist ein Lie-Algebra-Homomorphismus; $ad(x)$ heißt die *Adjungierte von x* . Falls A eine Banach-Lie-Algebra über $\mathbb{K}$ ist, dann ist ad stetig. Insbesondere ist somit für jede Banachalgebra X über $\mathbb{K}$ der Lie-Algebra-Homomorphismus $ad: X^\ominus \rightarrow \text{Der}(X)$ per $ad(x)y := xy - yx$, $x, y \in X$, definiert. (UPMEIER [279](1985), Seite 34)

2.3.32 Positivität. Es gibt verschiedene Konzepte, ein Element einer Algebra als positiv aufzufassen; deren Bezeichnungsweisen sind in der Literatur nicht einheitlich. Um im vorliegenden Text für eine gegebene Menge A — meist eine Algebra — ein Positivitätskonzept zu erklären, wird eine Teilmenge von A definiert, die genau aus den Elementen von A besteht, die als positiv bezüglich des gewählten Positivitätskonzeptes verstanden werden sollen; diese Teilmengen werden hier immer nach dem folgenden Schema bezeichnet:

$$\text{Pos}(\text{Schlüssel}; A).$$

Schlüssel steht dabei für eine abkürzende Bezeichnung des verwendeten Konzeptes von Positivität. Wenn die Menge $\text{Pos}(\text{Schlüssel}; A)$ ein Keil ist, dann wird gemäß den Bemerkungen 2.2.14 anstelle von $y - x \in \text{Pos}(\text{Schlüssel}; A)$ die Schreibweise $x \leq_{\text{Schlüssel}} y$ verwendet. Wenn klar ist, welches Konzept von Positivität gemeint ist, wird der Index *Schlüssel* weggelassen, also einfach nur $x \leq y$ geschrieben.

2.3.33. Sei A eine Algebra über $\mathbb{K}$. Da für die in 2.1.45 definierten Präordnungen für alle idempotenten Elemente p von A die Ungleichungen $0 \leq_\ell p$, $0 \leq_r p$ und $0 \leq p$ gelten, ist das folgende Konzept von Positivität naheliegend:

$$\text{Pos}(\text{idem}; A) := \text{Idem}(A).$$

Für alle $x, y \in \text{Idem}(A)$ gilt die Implikation $x \leq_{\ell r} y \Rightarrow x \leq_{\text{idem}} y$.

Beweis. $x \leq y \Rightarrow xy + yx = 2x \Leftrightarrow -yx - xy + x = -x \Leftrightarrow yy - yx - xy + xx = y - x \Leftrightarrow (y - x)^2 = y - x \Leftrightarrow y - x \text{ idempotent} \Leftrightarrow y - x \in \text{Pos}(\text{idem}; A) \Leftrightarrow x \leq_{\text{idem}} y$. $\square$

2.4 Topologische Vektorräume

2.4.1 Summen I (BOURBAKI [32](1966), III.5; KADISON und RINGROSE [167] (1983), Seite 25). Sei X ein topologischer Raum, der Hausdorff'sch ist und mit einer additiv geschriebenen Verknüpfung versehen ist, die assoziativ und kommutativ ist. Sei $\{a_\nu \in X : \nu \in I\}$ eine Menge von Elementen aus X , die über eine Indexmenge I indiziert sind. Sei $\mathcal{F}$ die per Mengeninklusion $A \leq B \Leftrightarrow A \subseteq B$ prägeordnete Menge aller endlichen Teilmengen von I . Konvergiert dann das Netz $\left(\sum_{\nu \in F} a_\nu \right)_{F \in \mathcal{F}}$, so wird die Menge $\{a_\nu \in X : \nu \in I\}$ *summierbar* genannt und der Grenzwert mit $\sum_{\nu \in I} a_\nu$ oder mit $\sum \{a_\nu \in X : \nu \in I\}$ bezeichnet.

Ist die Menge $\{a_\nu \in X : \nu \in I\}$ summierbar und ist φ eine Bijektion von einer Menge J auf I , so ist auch die Menge $\{a_{\varphi(\mu)} \in X : \mu \in J\}$ summierbar und es gilt $\sum_{\mu \in J} a_{\varphi(\mu)} = \sum_{\nu \in I} a_\nu$.

2.4.2 Definition. Sei $(G, \circ)$ eine Gruppe, die mit einer Topologie versehen ist. G heißt eine *topologische Gruppe*, wenn die Abbildungen $G \times G \rightarrow G$, $(x, y) \mapsto x \circ y$ und $G \rightarrow G$, $x \mapsto x^{-1}$ stetig sind, wobei $G \times G$ die Produkttopologie trägt.

Ein *topologischer Vektorraum* über $\mathbb{K}$ ist ein Vektorraum X über $\mathbb{K}$ mit einer Topologie τ , so dass die Abbildungen $X \times X \rightarrow X$, $(x, y) \mapsto x + y$ und $\mathbb{K} \times X \rightarrow X$, $(\lambda, x) \mapsto \lambda x$ stetig sind, wobei $X \times X$ und $\mathbb{K} \times X$ mit der jeweiligen Produkttopologie versehen sind; genauer schreibt man dann auch (X, τ) anstelle von X , um die Topologie kenntlich zu machen. Eine lineare stetige Abbildung zwischen zwei topologischen Vektorräumen heißt ein *Homomorphismus*, wenn sie offen ist.

Ein topologischer Vektorraum über $\mathbb{K}$ heißt *lokal konvex*, wenn er eine Nullumgebungsbasis hat, die aus konvexen Mengen besteht.

Sei X ein topologischer Vektorraum. Trägt X eine Norm $\|\cdot\|$, dann ist mit $\text{cl}(A)$ immer $\text{cl}(\|\cdot\|; A)$ gemeint, falls nicht ausdrücklich etwas anderes vereinbart ist.

2.4.3 Bemerkungen. (a) Ein topologischer Vektorraum X über $\mathbb{K}$ ist genau dann Hausdorff'sch, wenn für alle $x \in X^\times$ eine Nullumgebung existiert, welche nicht x enthält. KÖTHE [187](1966) definiert topologische Vektorräume direkt als Hausdorffräume.

(b) Jeder topologische Vektorraum $(X, +, \cdot_{\mathbb{K}})$ über $\mathbb{K}$ ist wegen $-x = (-1)x$ für alle $x \in X$ auch eine topologische Gruppe $(X, +)$. Somit kann man einen topologischen Vektorraum über $\mathbb{K}$ auch definieren als einen Vektorraum $(X, +)$ über $\mathbb{K}$, der eine topologische Gruppe ist, so dass die Skalarmultiplikation stetig ist.

(c) Sei X_ν , $\nu \in I$, eine Familie von topologischen Vektorräumen, alle über den gleichen Körper $\mathbb{K}$. Dann ist $X := \prod_{\nu \in I} X_\nu$, versehen mit der Produkttopologie, ein topologischer Vektorraum über $\mathbb{K}$. X ist genau dann Hausdorff'sch, wenn alle X_ν Hausdorff'sch sind.

(d) (ROBERTSON und ROBERTSON [247](1967), Seite 17). Sei X ein topologischer Vektorraum über $\mathbb{K}$, $\mathfrak{B}$ eine Nullumgebungsbasis in X und dann $U \in \mathfrak{B}$. Als eine Folge der in der Definition eines topologischen Vektorraumes geforderten Stetigkeit der Vektoraddition und der Skalarmultiplikation ist U absorbierend, es existiert ein $V \in \mathfrak{B}$ mit $V + V \subseteq U$ und es existiert eine kreisförmige Nullumgebung W in X mit $W \subseteq U$.

2.4.4. Sei X ein topologischer Vektorraum über $\mathbb{K}$. Sei $K \subseteq X$ kompakt und $G \subseteq X$ offen mit $K \subseteq G$. Dann existiert eine Nullumgebung U mit $K + U \subseteq G$.

Beweis. Zu jedem $x \in K$ sei U_x eine Nullumgebung mit $x + U_x \subseteq G$. Für jedes $x \in K$ sei dann V_x eine Nullumgebung mit $V_x + V_x \subseteq U_x$. Die $x + V_x$, $x \in K$, bilden eine offene Überdeckung von K . Es gibt $x_1, \dots, x_n$ aus K , so dass die $x_i + V_{x_i}$, $i = 1, \dots, n$, eine endliche offene Überdeckung von K bilden. $V := V_{x_1} \cap \dots \cap V_{x_n}$ ist eine Nullumgebung

mit $K + V \subseteq G$, denn: Ist $x \in K$, so ist $x \in x_j + V_{x_j}$ für ein $j \in \{1, \dots, n\}$ und es gilt $x + V \subseteq x_j + V_{x_j} + V \subseteq x_j + V_{x_j} + V_{x_j} \subseteq x_j + U_{x_j} \subseteq G$.

Alternativ sei hier noch ein zweiter Beweis angeführt, siehe KELLEY und NAMIOKA [185](1963), Seite 35, 5.2(vi): Angenommen, es gebe keine solche Nullumgebung. Dann gibt es also zu jeder Nullumgebung U ein $x_U \in K$ mit $(x_U + U) \cap (X \setminus G) \neq \emptyset$. Betrachte das Netz $\{x_U, U \text{ Nullumgebung}, \supseteq\}$. Da K kompakt ist, hat dieses Netz einen Häufungspunkt $x_0 \in K$, heißt, für jede Umgebung W von x_0 und für jede Nullumgebung A gibt es eine Nullumgebung B mit $A \supseteq B$ und $x_B \in W$. Sei Z eine Umgebung von x_0 . Z ist von der Gestalt $x_0 + V$, V eine Nullumgebung. Es gibt eine Nullumgebung A mit $A + A \subseteq V$. Zur Umgebung $W := x_0 + A$ existiert eine Nullumgebung B mit $A \supseteq B$ und $x_B \in W$. Somit $x_B + B \subseteq x_B + A \subseteq W + A = x_0 + A + A \subseteq x_0 + V = Z$. Da aber $x_B + B$ in $X \setminus G$ reinschneidet, tut dies auch Z . Somit schneidet jede Umgebung von x_0 in $X \setminus G$ hinein, im Widerspruch dazu, dass x_0 ein innerer Punkt von G ist. $\square$

Allgemeiner gilt sogar, siehe RUDIN [251](1973), Seite 9, Theorem 1.10: Ist $K \subseteq X$ kompakt und $C \subseteq X$ abgeschlossen mit $K \cap C = \emptyset$, so existiert eine Nullumgebung U mit $(K + U) \cap (C + U) = \emptyset$.

2.4.5 Cauchy-Filter (HORVÁTH [142](1966), Seite 128). Sei X ein topologischer Vektorraum über $\mathbb{K}$ und $M \subseteq X$. Ein Filter $\mathcal{F}$ auf M heißt ein *Cauchy-Filter* (auf M), wenn für jede Nullumgebung U in X ein $A \in \mathcal{F}$ mit $A - A \subseteq U$ existiert. Jeder Filter auf M , der zu einem Punkt von M konvergiert, ist ein Cauchy-Filter. M heißt *vollständig*, wenn jeder Cauchy-Filter auf M zu einem Punkt von M konvergiert. Ein Netz $(x_\nu)_{\nu \in I}$ in X heißt ein *Cauchy-Netz*, wenn der zu dem Netz $(x_\nu)_{\nu \in I}$ gehörige Abschnittsfilter ein Cauchy-Filter ist; ist dabei das Netz $(x_\nu)_{\nu \in I}$ eine Folge in X , so heißt $(x_\nu)_{\nu \in I}$ eine *Cauchy-Folge*. M heißt *folgenvollständig*, wenn jede Cauchy-Folge in X konvergiert. Da es für jede Nullumgebung U eine kreisförmige Nullumgebung A in X mit $A + A \subseteq U$, also auch $A - A \subseteq U$, gibt, ist ein Netz $(x_\nu)_{\nu \in I}$ in X genau dann ein Cauchy-Netz, wenn für jede Nullumgebung V in X ein $\gamma \in I$ existiert, so dass aus dem gleichzeitigen Vorliegen von sowohl $\nu \geq \gamma$ als auch $\mu \geq \gamma$ stets $x_\nu - x_\mu \in V$ folgt.

2.4.6 Definition. Ein *Banachraum* über $\mathbb{K}$ ist ein normierter Vektorraum über $\mathbb{K}$, der bezüglich seiner Normtopologie vollständig ist. Man sagt dann auch, dass die Norm *vollständig* sei. Ein *Hilbertraum* über $\mathbb{K}$ ist ein Prähilbertraum über $\mathbb{K}$ (siehe Definition 2.2.25), der bezüglich seiner Normtopologie vollständig ist.

2.4.7 Summen II. (BOURBAKI [32](1966), III.5). Sei X eine additiv geschriebene kommutative topologische Gruppe, die Hausdorff'sch ist. Ist $\{a_\nu \in X : \nu \in I\}$ eine summierbare Menge, so ist die Menge der $\nu \in I$ mit $a_\nu \neq 0$ abzählbar, falls X eine abzählbare Nullumgebungsbasis hat.

Sei $(a_n)_{n \in \mathbb{N}}$ eine Folge in X . Setze $s_n := \sum_{k=0}^n a_k$ für alle $n \in \mathbb{N}$. Die durch die Folge $(a_n)_{n \in \mathbb{N}}$ definierte *Reihe*, in Zeichen $(a_n)_{n \in \mathbb{N}}$ oder $\sum_{n=0}^{\infty} a_n$, ist das Paar der beiden Folgen $(a_n)_{n \in \mathbb{N}}$ und $(s_n)_{n \in \mathbb{N}}$; diese Reihe heißt *konvergent*, wenn die Folge $(s_n)_{n \in \mathbb{N}}$ konvergent ist; und dann heißt der Grenzwert der Folge $(s_n)_{n \in \mathbb{N}}$ die *Summe* der Reihe und wird mit $\sum_{n=0}^{\infty} a_n$ bezeichnet. a_n heißt das *n-te Glied* der Reihe; s_n heißt die *n-te Partialsumme* der Reihe. Man spricht auch von der Reihe $(a_n)_{n \in \mathbb{N}}$ als die Reihe mit dem *allgemeinen Glied* a_n .

Ist die Folge $(a_n)_{n \in \mathbb{N}}$ summierbar — womit natürlich gemeint ist, dass die Menge $\{a_n \in X : n \in \mathbb{N}\}$ summierbar ist —, so ist die Reihe $\sum_{n=0}^{\infty} a_n$ konvergent und es gilt $\sum_{n \in \mathbb{N}} a_n = \sum_{n=0}^{\infty} a_n$. Die Umkehrung stimmt im Allgemeinen nicht.

Eine Reihe $\sum_{n=0}^{\infty} a_n$ heißt *kommutativ konvergent*, wenn für jede Permutation σ von $\mathbb{N}$ die Reihe $\sum_{n=0}^{\infty} a_{\sigma(n)}$ konvergiert. Eine Reihe $\sum_{n=0}^{\infty} a_n$ ist genau dann kommutativ konvergent, wenn die Folge $(a_n)_{n \in \mathbb{N}}$ summierbar ist. Ist dabei X ein Hausdorff'scher

topologischer Vektorraum, so ist es üblich, statt kommutativ konvergent die Bezeichnung *unbedingt konvergent* (im Englischen *unconditionally convergent*) zu verwenden. Ist X ein normierter Vektorraum, so heißt eine Menge $\{a_\nu \in X : \nu \in I\}$ *absolut summierbar*, wenn die Menge $\{\|a_\nu\| \in \mathbb{R} : \nu \in I\}$ in $\mathbb{R}$ summierbar ist; in der Theorie der normierten Vektorräume wird dies durch die Formulierung „die Reihe $\sum_{\nu \in I} a_\nu$ konvergiert absolut“ ausgedrückt.

Speziell für lokal konvexe Vektorräume und $I = \mathbb{N}$ findet man bei DAY [61](1973), Kapitel IV, noch weitere Arten von Reihen-Konvergenz; dort findet sich der Begriff „summierbar“ als *unordered convergent* (im Deutschen also etwa „ungeordnet konvergent“) wieder, und der Begriff „kommutativ konvergent“ als *reordered convergent* (im Deutschen also etwa „umgeordnet konvergent“); ist X ein lokal konvexer Vektorraum, der Hausdorff'sch ist, so heißt eine Menge $\{a_n \in X : n \in \mathbb{N}\}$ *absolut konvergent*, wenn für jede Nullumgebung U in X die Reihe $\sum_{n \in \mathbb{N}} p_U(a_n)$ konvergiert, wobei p_U das Minkowski-Funktional von U ist. Für folgenvollständige lokal konvexe Hausdorffräume X ist jede absolut konvergente Menge $\{a_n \in X : n \in \mathbb{N}\}$ summierbar; ist X endlichdimensional, so gilt auch die Umkehrung. Ist dagegen X ein unendlichdimensionaler Banachraum, so gilt die Umkehrung nicht mehr, da Dvoretzky und Rogers 1950 zeigten: Sei X ein unendlichdimensionaler Banachraum und sei $(a_n)_{n \in \mathbb{N}}$ eine beliebige Folge positiver Zahlen mit $\sum_{n \in \mathbb{N}} a_n^2 < \infty$, dann existiert in X eine summierbare Folge $(x_n)_{n \in \mathbb{N}}$ mit $\|x_n\| = a_n$ für alle $n \in \mathbb{N}$. (Einfaches Beispiel: Setze $X := c_0$, $a_n := \frac{1}{n}$, $x_n := \frac{1}{n}(\delta_{nm})_{m \in \mathbb{N}}$ (Kronecker-Delta), $n \in \mathbb{N}$.) Im dazu gehörigen Beweis werden die x_n , $n \in \mathbb{N}$, als „approximativ orthogonal“ konstruiert. Dvoretzky hat diese beinahe-Orthogonalität weiter verbessert: Sei X ein unendlichdimensionaler normierter Vektorraum, $n \in \mathbb{N}$ und $\varepsilon > 0$. Dann existiert eine bijektive lineare Abbildung T von $\mathbb{R}^n$, versehen mit der euklidischen Norm, auf einen Unterraum $U = U(n, \varepsilon)$ von X mit $\|T\| \|T^{-1}\| < 1 + \varepsilon$; man sagt: Jeder unendlichdimensionale normierte Vektorraum *imitiert* einen Hilbertraum. So imitiert zum Beispiel c_0 jeden normierten Raum und jeder Raum X imitiert seinen Bidualraum X^{**} ; siehe DAY [61](1973), Seite 168.

Ist $X = \mathbb{K}$, $\{a_\nu \in \mathbb{K} : \nu \in I\}$ eine Menge positiver Zahlen und bezeichnet s das Element $\sup \left\{ \sum_{\nu \in F} a_\nu \in \mathbb{K} : F \subseteq I \text{ endlich} \right\}$ aus $[0, \infty]$, so ist die Menge $\{a_\nu \in \mathbb{K} : \nu \in I\}$ genau dann summierbar, wenn $s < \infty$ ist, und in diesem Fall gilt $\sum_{\nu \in I} a_\nu = s$.

Einen einführenden Überblick über die Konvergenz von Reihen in Hausdorff'schen topologischen Vektorräumen geben KAMTHAN und GUPTA [169] (1981), Kapitel 3. Dort findet man unter anderem die folgenden Aussagen: **(a)** Ist X ein metrisierbarer lokal konvexer Raum, in dem jede absolut konvergente Reihe konvergiert, so ist X vollständig (also ein sogenannter Fréchet-Raum) (ebd., Seite 142). **(b)** Ein metrisierbarer vollständiger topologischer Vektorraum ist genau dann lokal konvex (also wieder ein sogenannter Fréchet-Raum), wenn jede absolut konvergente Reihe summierbar ist (ebd., Seite 161). **(c)** Ein metrisierbarer, vollständiger, lokal konvexer Vektorraum (also abermals ein sogenannter Fréchet-Raum) ist genau dann ein sogenannter *nuklearer Raum*, wenn jede summierbare Reihe absolut konvergent ist (ebd., Seite 162).

2.4.8 Satz (BOURBAKI [35](1987), I.2.3). *Sei X ein topologischer Vektorraum über $\mathbb{K}$. Sei U ein abgeschlossener Unterraum von X und F ein endlichdimensionaler Unterraum von X . Dann ist der Unterraum $U + F$ abgeschlossen in X .*

Beweis. Sei $T: X \rightarrow X/U$, $x \mapsto x + U$, die kanonische Quotientenabbildung und sei X/U mit der Quotiententopologie versehen, das heißt, mit der finalen Topologie von X/U bezüglich X und T (GROTEMEYER: Topologie, 1969; RUDIN [251](1973)). Es gilt: $T^{-1}(T(F)) = T^{-1}(\bigcup_{x \in F} (x + U)) = \bigcup_{x \in F} T^{-1}(x + U) = \bigcup_{x \in F} (x + U) = F + U$. Da ein

Unterraum S eines topologischen Raumes Y genau dann in Y abgeschlossen ist, wenn Y/S Hausdorff'sch ist (HOLMES [137](1975), Seite 51, Theorem 9D; ROBERTSON und ROBERTSON [247](1967), Seite 87), ist hier der Quotientenraum X/U Hausdorff'sch. Da jeder endlichdimensionale Unterraum eines Hausdorff'schen topologischen Vektorraumes abgeschlossen ist, ist der endlichdimensionale Unterraum $T(F)$ in X/U abgeschlossen. Somit ist $F + U = T^{-1}(T(F))$ abgeschlossen in X . $\square$

2.4.9 Definition. Seien X und Y zwei topologische Vektorräume über $\mathbb{K}$. Mit $\mathcal{L}(X, Y)$ wird der Vektorraum über $\mathbb{K}$ aller stetigen linearen Abbildungen von X nach Y bezeichnet. $\mathcal{L}(X, \mathbb{K})$ heißt der *stetige Dualraum* von X oder meist einfach nur der *Dualraum* von X und wird mit X^* bezeichnet. Analog wie bei $X^\#$ werden die Elemente von X^* suggestiv meist als $x^*, y^*, \dots$ notiert; gelegentlich wird auch hier das Symbol ℓ für Elemente von X^* verwendet. Die per $x \mapsto (x^* \mapsto x^*(x))$ definierte *kanonische Inklusionsabbildung* von X nach X^{**} wird (wieder) mit i_X bezeichnet; sie kann auch genauer als die *kanonische topologische Inklusionsabbildung* angesprochen werden. Für $x \in X$ schreibt man für das Auswertungsfunktional $i_X(x)$ auch hier das Symbol $\hat{x}$. Ist U ein Unterraum von X^* so wird ebenfalls wieder mit $i_{X,U}$ die Abbildung $X \rightarrow U^*, x \mapsto (\ell \mapsto \ell(x))$ bezeichnet. Aus dem jeweiligen Zusammenhang wird immer klar hervorgehen, ob mit i_X beziehungsweise $i_{X,U}$ die algebraische oder die topologische Variante gemeint ist. So ist, falls nicht etwas anderes vereinbart ist, im Rahmen von topologischen Vektorräumen immer die topologische Variante gemeint.

Für $T \in \mathcal{L}(X, Y)$ ist die *duale Abbildung* von T die mit T^* bezeichnete Astriktion und Restriktion $X^* \upharpoonright T^\# \upharpoonright Y^*$ von $T^\#$ auf X^* und Y^* . $\mathcal{L}(X) := \mathcal{L}(X, X)$ wird als eine Unter algebra von $L(X)$ aufgefasst.

Des Weiteren setzt man $\mathcal{F}(X, Y) := F(X, Y) \cap \mathcal{L}(X, Y)$ und $\mathcal{F}(X) := \mathcal{F}(X, X)$.

Eine Abbildung $T \in L(X, Y)$ heißt *kompakt*, wenn es eine Nullumgebung von X gibt, deren Bild unter T in Y relativ kompakt ist. (Folglich ist dann T auch stetig.) Die Menge aller kompakten Abbildungen von X nach Y wird mit $K(X, Y)$ bezeichnet. Man setzt $K(X) := K(X, X)$.

Sei X ein normierter Vektorraum über $\mathbb{K}$ und U ein abgeschlossener Unterraum von X^* . Dann heißt U *Norm-bestimmend* oder *normierend* oder *duzial* (für X), wenn für alle $x \in X$ die Gleichung $\|x\| = \sup \{|\ell(x)| : \ell \in U\}$ gilt — mit anderen Worten, wenn die Abbildung $i_{X,U} : X \rightarrow U^*$ isometrisch ist.

2.4.10. Ist X ein Banachraum über $\mathbb{K}$, dann ist $K(X)$ ein abgeschlossenes Ideal in $\mathcal{L}(X)$.

Sind X und Y zwei Banachräume über $\mathbb{K}$, so gilt $\mathcal{F}(X, Y) = F(X, Y) \cap K(X, Y)$ (MEGGINSON [211](1998), Seite 320).

Sind X und Y zwei Hausdorff'sche lokal konvexe topologische Vektorräume über $\mathbb{K}$, so gilt $\mathcal{F}(X, Y) \subseteq K(X, Y)$ (SCHAEFER [255](1967), Seite 98).

Ist X ein Banachraum über $\mathbb{K}$ mit $X \neq \{0\}$, dann ist $\mathcal{F}(X)$ ein minimales Ideal von $\mathcal{L}(X)$ (DALES [56](2000), Seite 41; RICKART [246](1960), Seite 278).

Es sei an 2.2.26 erinnert. Sind X und Y zwei normierte Vektorräume über $\mathbb{K}$, so gilt

$$\mathcal{F}(X, Y) = \left\{ \sum_{i=1}^n y_i \otimes x_i^* \in L(X, Y) : n \in \mathbb{N}, y_i \in Y, x_i^* \in X^* \right\}.$$

2.4.11 Definition. Seien X, Y und V drei normierte Vektorräume über $\mathbb{K}$. Ein $T \in \mathcal{L}(X, Y)$ heißt *den Raum V fixierend*, wenn X eine isomorphe Kopie W von V enthält, so dass $T(W) \upharpoonright T \upharpoonright W$ ein Isomorphismus ist.

2.4.12 Stetige n -lineare Abbildungen. In Analogie zu 2.2.21 definiert man weiter: Sei $n \in \mathbb{N}^\times$ und seien $X_1, \dots, X_n, Y$ normierte Vektorräume über den gleichen Körper

$\mathbb{K}$. Dann ist durch

$$\|T\| := \sup_{\max\{\|x_1\|, \dots, \|x_n\|\} \leq 1} \|T(x_1, \dots, x_n)\| \quad \text{für alle } T \in L(X_1, \dots, X_n; Y)$$

eine Abbildung $L(X_1, \dots, X_n; Y) \rightarrow \mathbb{R} \cup \{\infty\}$ definiert. Die Menge aller $T \in L(X_1, \dots, X_n; Y)$ mit $\|T\| < \infty$ bildet einen Unterraum von $L(X_1, \dots, X_n; Y)$, auf dem die soeben definierte Abbildung $\|\cdot\|$ eine Norm ist; versehen mit dieser Norm ist dieser Unterraum also ein normierter Vektorraum über $\mathbb{K}$ — er wird mit $\mathcal{L}(X_1, \dots, X_n; Y)$ bezeichnet; er ist vollständig, wenn Y vollständig ist. Sind alle $X_1, \dots, X_n$ von $\{0\}$ verschieden, so gilt

$$\|T\| = \sup_{x_1, \dots, x_n \neq 0} \frac{\|T(x_1, \dots, x_n)\|}{\|x_1\| \cdots \|x_n\|} \quad \text{für alle } T \in L(X_1, \dots, X_n; Y).$$

Gilt $X_1 = \cdots = X_n$ so wird mit $X := X_1$ anstelle von $\mathcal{L}(X_1, \dots, X_n; Y)$ die Bezeichnung $\mathcal{L}^n(X, Y)$ verwendet. $\mathcal{L}^0(X, Y)$ wird als Y definiert. Ist dann auch noch $X = Y$, so wird anstelle von $\mathcal{L}^n(X, Y)$ einfach nur $\mathcal{L}^n(X)$ geschrieben.

Die Abbildung $\mathcal{L}(X_1, X_2; Y) \rightarrow \mathcal{L}(X_1, \mathcal{L}(X_2, Y))$, $T \mapsto (x_1 \mapsto T(x_1, \cdot))$ mit der Umkehrabbildung $R \mapsto ((x_1, x_2) \mapsto (Rx_1)(x_2))$ ist ein isometrischer Isomorphismus.

2.4.13 Stetige n -homogene Polynome. Als Fortsetzung von 2.2.40 hat man: Sei $n \in \mathbb{N}^\times$ und seien X, Y normierte Vektorräume über $\mathbb{K}$. Eine Abbildung $p: X \rightarrow Y$ heißt *stetiges n -homogenes Polynom* (oder *stetiges homogenes Polynom vom Grad n*), wenn ein $T \in \mathcal{L}^n(X, Y)$ existiert mit $p(x) = T(x, \dots, x)$ für alle $x \in X$. Der Vektorraum aller stetigen n -homogenen Polynome wird mit $\mathcal{P}^n(X, Y)$ bezeichnet und mit der Norm $\|p\| := \sup\{\|p(x)\| : x \in B_X\}$, $p \in \mathcal{P}^n(X, Y)$, versehen; diese Norm ist vollständig, wenn Y vollständig ist; im Fall von $X \neq \{0\}$ gilt $\|p\| = \sup_{x \neq 0} \frac{\|p(x)\|}{\|x\|^n}$. $\mathcal{P}^0(X, Y)$ ist der Vektorraum $P^0(X, Y)$ aller konstanten Abbildungen von X nach Y . Es ist $\mathcal{P}^1(X, Y) = \mathcal{L}(X, Y)$. Für $n \in \mathbb{N}$ ist $\mathcal{P}^n(X)$ als $\mathcal{P}^n(X, X)$ definiert.

2.4.14 (SCHAEFER [255](1966), Seite 19). Sei X_ν , $\nu \in I$, eine Familie von topologischen Vektorräumen, alle über den gleichen Körper $\mathbb{K}$. Dann ist das direkte Produkt $\prod_{\nu \in I} X_\nu$, versehen mit der Produkttopologie, wieder ein topologischer Vektorraum über $\mathbb{K}$; er wird wieder mit $\prod_{\nu \in I} X_\nu$ bezeichnet und ist genau dann Hausdorff'sch, wenn alle X_ν , $\nu \in I$, Hausdorff'sch sind.

2.4.15 Innere topologische direkte Summe (SCHAEFER [255](1966)). Sei X ein topologischer Vektorraum über $\mathbb{K}$, $n \in \mathbb{N}^\times$ und seien $U_1, \dots, U_n$ Untervektorräume über $\mathbb{K}$ von X mit $X = \sum_{i=1}^n U_i$. Für jedes $i \in \{1, \dots, n\}$ bezeichne X_i den Vektorraum U_i , versehen mit der von X induzierten Relativtopologie. Für jedes $i \in \{1, \dots, n\}$ sei p_i die Abbildung $X \rightarrow X$, $x_1 + \dots + x_n \mapsto x_i$, wobei $x_j \in U_j$ für alle $j \in \{1, \dots, n\}$ gilt. Dann ist jedes p_i , $i \in \{1, \dots, n\}$, eine Projektion aus $L(X)$. Für jedes $i \in \{1, \dots, n\}$ bezeichne q_i die Astriktion von p_i auf U_i , also die surjektive und offene Abbildung $X \rightarrow X_i$, $x \mapsto q_i(x) := p_i(x)$.

Für alle $i \in \{1, \dots, n\}$ gilt, dass die Projektion p_i genau dann stetig ist, wenn die Abbildung q_i stetig ist.

Beweis. Sei $i \in \{1, \dots, n\}$. Sei p_i stetig und V_i eine Nullumgebung von X_i . Dann gibt es eine Nullumgebung V von X mit $V_i = V \cap U_i$. Da p_i stetig ist, gibt es eine Nullumgebung U von X mit $p_i(U) \subseteq V$. Wegen $p_i(X) \subseteq U_i$ also $p_i(U) \subseteq V_i$ und q_i ist stetig. Sei nun umgekehrt q_i stetig und V eine Nullumgebung von X . Dann ist $V_i := V \cap U_i$ eine Nullumgebung von X_i . Da q_i stetig ist, existiert eine Nullumgebung U von X mit $q_i(U) \subseteq V_i$, also $p_i(U) \subseteq V_i$, $p_i(U) \subseteq V$ und p_i ist stetig. $\square$

Die bijektive Abbildung $S: \prod_{i=1}^n X_i \rightarrow \sum_{i=1}^n U_i$, $(x_1, \dots, x_n) \mapsto \sum_{i=1}^n x_i$ ist stetig. Falls S ein Isomorphismus, das heißt, ein Homöomorphismus ist, dann wird X die *topologische direkte Summe* (genauer *innere topologische direkte Summe*, im Englischen *internal topological direct sum*) der $U_1, \dots, U_n$ genannt. Diese Homöomorphie liegt genau dann vor, wenn alle Projektionen p_i , $i \in \{1, \dots, n\}$, stetig sind; dazu beachte man nur, dass die Umkehrabbildung von S gleich $S^{-1}: X \rightarrow \prod_{i=1}^n X_i$, $x \mapsto (q_1(x), \dots, q_n(x))$, ist.

Ist X die topologische direkte Summe zweier Untervektorräume U und V von X , dann sagt man, dass U das *topologische Komplement* von V in X sei. Eine dafür notwendige, aber im Allgemeinen nicht hinreichende Bedingung ist, dass U und V beide abgeschlossen sind. Ist X die innere direkte algebraische Summe einer Familie von abgeschlossenen Unterräumen U_ν , $\nu \in I$, so sagt man, dass der topologische Vektorraum X über $\mathbb{K}$ die *direkte Summe* (genauer *innere direkte Summe*) der U_ν , $\nu \in I$, sei. Ist speziell ein topologischer Vektorraum X über $\mathbb{K}$ die innere direkte Summe zweier Untervektorräume U und V von X , so sagt man, dass U das *Komplement* von V in X sei.

Die beiden Begriffe „topologisches Komplement“ und „Komplement“ fallen zum Beispiel für alle Fréchet-Räume — soll heißen, die metrisierbaren, vollständigen lokal konvexen Vektorräume über $\mathbb{K}$ — zusammen (EDWARDS [88](1965), Seite 458, Übung 6.2).

Selbst endlichdimensionale Unterräume eines topologischen Vektorraumes über $\mathbb{K}$ müssen keinen topologischen Komplementärraum besitzen. Zum Beispiel besitzt im Funktionenraum $L_p(\mu)$, $0 < p < 1$, kein endlichdimensionaler Untervektorraum einen topologischen Komplementärraum, siehe KÖTHE [187](1966), Seite 162. Ist aber X ein lokal konvexer Vektorraum über $\mathbb{K}$, der Hausdorff'sch ist, so besitzt jeder endlichdimensionale Untervektorraum von X einen topologischen Komplementärraum (KÖTHE [187](1966), Seite 242).

Ein Beispiel eines normierten Raumes mit einem komplementierten Unterraum, der nicht topologisch komplementiert ist, findet sich bei MEGGINSON [211](1998), Seite 299, Beispiel 3.2.13.

Für die Situation, dass X die topologische direkte Summe zweier Untervektorräume U und V von X ist, gibt es mehrere zueinander äquivalente Sprechweisen, beispielsweise: (a) U ist *topologisch komplementierbar* (im Englischen *topological complementable*) in X ; (b) U ist *topologisch komplementiert* (im Englischen *topological complemented*); (c) U und V sind *topologisch komplementär*; (d) UPMEIER [279](1985), Seite 42, nennt solch ein U im Englischen auch *split*, wo dieses Wort je nach Verwendung ein Nomen, Adjektiv oder Verb sein kann. Er schreibt dann $X = \begin{pmatrix} U \\ V \end{pmatrix}$.

Ist X ein normierter Vektorraum mit einem topologisch komplementierten Untervektorraum U von X und $p \in L(X)$ die (eindeutig bestimmte) Projektion auf U , so sagt man im Fall von $\|p\| = 1$, dass U *Norm-1-komplementiert* sei, oder einfach auch nur, dass U *1-komplementiert* sei.

2.4.16 Äußere topologische direkte Summe (ROBERTSON und ROBERTSON [247](1967); BOURBAKI [35](1987), II.4.5; EDWARDS [88](1965), Seite 430; KÖTHE [187](1966), Seite 217).

Sei X_ν , $\nu \in I$, eine Familie von topologischen Vektorräumen, wobei alle X_ν Hausdorff'sch und Vektorräume über den gleichen Körper $\mathbb{K}$ sind. Sei $\mathfrak{u}_\nu$ für jedes $\nu \in I$ die Menge aller offenen Nullumgebungen in X_ν . Dann bildet die Menge

$$\mathfrak{B}_\delta := \left\{ \bigoplus_{\nu \in I} U_\nu \subseteq \bigoplus_{\nu \in I} X_\nu : (U_\nu)_{\nu \in I} \in \prod_{\nu \in I} \mathfrak{u}_\nu \right\}$$

eine Nullumgebungsbasis einer Topologie δ auf der direkten Summe $\bigoplus_{\nu \in I} X_\nu$.

Bezeichne π die von der Produkttopologie auf $\prod_{\nu \in I} X_\nu$ auf der direkten Summe $\bigoplus_{\nu \in I} X_\nu$ induzierte Topologie. Dann ist die Menge $\mathfrak{S}_\pi := \{(\prod_{\nu \in I} Y_\nu) \cap \bigoplus_{\nu \in I} X_\nu \subseteq \bigoplus_{\nu \in I} X_\nu : \exists \mu \in I \forall \nu \in I \setminus \{\mu\} : Y_\mu \in \mathfrak{U}_\mu, Y_\nu = X_\nu\}$, also

$$\mathfrak{S}_\pi = \left\{ \bigoplus_{\nu \in I} Y_\nu \subseteq \bigoplus_{\nu \in I} X_\nu : \exists \mu \in I \forall \nu \in I \setminus \{\mu\} : Y_\mu \in \mathfrak{U}_\mu, Y_\nu = X_\nu \right\},$$

eine Nullumgebungssubbasis für die Topologie π . Folglich ist die Menge

$$\mathfrak{B}_\pi := \left\{ \bigoplus_{\nu \in I} Y_\nu \subseteq \bigoplus_{\nu \in I} X_\nu : \exists F \subseteq I \text{ endlich} : \forall \mu \in F : Y_\mu \in \mathfrak{U}_\mu, \right. \\ \left. \forall \nu \in I \setminus F : Y_\nu = X_\nu \right\}$$

eine Nullumgebungsbasis von π .

Offensichtlich gilt $\mathfrak{S}_\pi \subseteq \mathfrak{B}_\pi \subseteq \mathfrak{B}_\delta$, so dass also nach 1.2.17(c) die Relativtopologie π gröber als die Topologie δ ist.

In Anlehnung an ROBERTSON und ROBERTSON wird die Topologie π hier als die *Produkttopologie auf $\bigoplus_{\nu \in I} X_\nu$* bezeichnet. $\bigoplus_{\nu \in I} X_\nu$, versehen mit der Topologie δ , nennt KÖTHE die *topologische direkte Summe* der X_ν .

Sind die X_ν alle lokal konvex und Hausdorff'sch, so ist die Topologie δ ebenfalls lokal konvex und vollständig für vollständige X_ν .

Im Rest des vorliegenden Abschnittes 2.4.16 seien ab hier alle X_ν lokal konvex. Sowohl die Produkttopologie π als auch die Topologie δ auf $\bigoplus_{\nu \in I} X_\nu$ besitzen noch zwei ästhetische Mängel. Erstens zeigen diese Topologien in der Theorie der Dualität nicht unbedingt immer die Eigenschaften, die man gerne hätte. So gilt zum Beispiel für die Topologie δ , falls alle X_ν Hausdorff'sch sind, dass der duale Raum von $\bigoplus_{\nu \in I} X_\nu$ nur der Teilraum ist von $\prod_{\nu \in I} X_\nu^*$, der alle $(x_\nu^*)_{\nu \in I}$, $x_\nu^* \in X_\nu^*$, mit höchstens abzählbar vielen $x_\nu^* \neq 0$ enthält. Und zweitens wird zwar durch jeweils beide Topologien auf jedem X_ν die ursprüngliche Topologie induziert, aber im Allgemeinen ist die Topologie δ auf $\bigoplus_{\nu \in I} X_\nu$ nicht die feinste, die das leistet.

Bezeichnet man für jedes $\nu \in I$ mit i_ν die kanonische Injektion von X_ν in $\bigoplus_{\nu \in I} X_\nu$, so nennt man die finale lokal konvexe Topologie τ auf $\bigoplus_{\nu \in I} X_\nu$ bezüglich der Funktionenfamilie i_ν , $\nu \in I$, die *Topologie der lokal konvexen direkten Summe* oder *lokal konvexe direkte Summen-Topologie* auf $\bigoplus_{\nu \in I} X_\nu$ und $\bigoplus_{\nu \in I} X_\nu$ heißt dann die *lokal konvexe direkte Summe* der lokal konvexen Räume X_ν , $\nu \in I$; der Deutlichkeit halber schreibt man sie auch als $(\bigoplus_{\nu \in I} X_\nu)_f$, wobei der Index „f“ für „final“ steht. BOURBAKI nennt diese Topologie τ die *direkte Summe der Topologien der X_ν , $\nu \in I$* , oder einfach auch nur die *direkte Summen-Topologie*. Es sei hier nur angemerkt, dass diese Topologie τ gleich dem sogenannten *induktiven Limes* der Topologien der X_ν , $\nu \in I$, unter der Funktionenfamilie i_ν , $\nu \in I$, ist. Sei $\mathfrak{A}_\nu$ für jedes $\nu \in I$ eine Basis aus absolutkonvexen Umgebungen in X_ν . Dann bildet auf $\bigoplus_{\nu \in I} X_\nu$ die Menge

$$\left\{ \text{co} \left(\bigcup_{\nu \in I} i_\nu(A_\nu) \right) \subseteq \bigoplus_{\nu \in I} X_\nu : (A_\nu)_{\nu \in I} \in \prod_{\nu \in I} \mathfrak{A}_\nu \right\}$$

eine Umgebungsbasis der lokal konvexen direkten Summen-Topologie τ .

Diese Topologie hat die folgenden Eigenschaften. Auf jedem X_ν induziert die lokal konvexe direkte Summen-Topologie τ die ursprüngliche Topologie. Die lokal konvexe direkte Summen-Topologie τ auf $\bigoplus_{\nu \in I} X_\nu$ ist feiner als die Topologie δ auf $\bigoplus_{\nu \in I} X_\nu$. Für jede endliche Teilmenge J von der Indexmenge I fallen die lokal konvexe direkte Summen-Topologie τ , die Topologie δ und die Produkttopologie π auf $\bigoplus_{\nu \in J} X_\nu$, betrachtet als

Unterraum von $\bigoplus_{\nu \in I} X_\nu$, zusammen. Für abzählbar viele Summanden X_ν , alle X_ν Hausdorff'sch, stimmen die lokal konvexe direkte Summen-Topologie τ und die Topologie δ auf $\bigoplus_{n=1}^\infty X_n$ überein. Ist die Indexmenge I unendlich und enthält jedes X_ν , $\nu \in I$, eine von X_ν verschiedene Nullumgebung U_ν , so ist die Topologie δ echt feiner als die Produkttopologie π . Denn: $\bigoplus_{\nu \in I} U_\nu$ ist zwar dann ein Element von $\mathfrak{B}_\delta$, aber es gibt kein Element $\bigoplus_{\nu \in I} Y_\nu$ aus $\mathfrak{B}_\pi$ mit $\bigoplus_{\nu \in I} Y_\nu \subseteq \bigoplus_{\nu \in I} U_\nu$, so dass also aus 1.2.17(f) die Behauptung mit $x = 0$ folgt.

Die lokal konvexe direkte Summen-Topologie τ auf $\bigoplus_{\nu \in I} X_\nu$ ist genau dann Hausdorff'sch, wenn jedes X_ν Hausdorff'sch ist und in diesem Fall ist einerseits jedes X_ν in $\bigoplus_{\nu \in I} X_\nu$ abgeschlossen und andererseits die lokal konvexe direkte Summen-Topologie τ auf $\bigoplus_{\nu \in I} X_\nu$ genau dann vollständig, wenn jeder der Räume X_ν vollständig ist.

Die lokal konvexe direkte Summen-Topologie τ auf $\bigoplus_{\nu \in I} X_\nu$ zeigt besonders angenehme Eigenschaften hinsichtlich der Bildung von Dualräumen. Einerseits ist nämlich der topologische Dualraum der lokal konvexen direkten Summe $\bigoplus_{\nu \in I} X_\nu$ das Produkt der topologischen Dualräume und andererseits ist der topologische Dualraum des topologischen Produktes $\prod_{\nu \in I} X_\nu$ die lokal konvexe direkte Summe der topologischen Dualräume:

$$\left(\bigoplus_{\nu \in I} X_\nu\right)^* = \prod_{\nu \in I} X_\nu^* \quad \text{und} \quad \left(\prod_{\nu \in I} X_\nu\right)^* = \bigoplus_{\nu \in I} X_\nu^*.$$

Siehe dazu auch 2.4.37 und 2.4.38.

2.4.17 Funktionenräume. Sei X ein topologischer Raum. Im Folgenden werden einige Funktionenräume genannt, die man mit komponentenweise definierter Addition und gegebenenfalls Multiplikation von auf X definierten Funktionen erhalten kann.

(a) Sei Y ein topologischer Vektorraum über $\mathbb{K}$. Dann bezeichnet $C(X, Y)$ den Vektorraum über $\mathbb{K}$ aller stetigen Funktionen von X nach Y . Ist Y eine Algebra über $\mathbb{K}$, dann ist auch $C(X, Y)$ eine Algebra über $\mathbb{K}$.

(b) Sei Y ein normierter Vektorraum über $\mathbb{K}$. Dann bezeichnet $C^b(X, Y)$ den Vektorraum über $\mathbb{K}$ aller stetigen Y -wertigen Funktionen f auf X , die beschränkt sind, das heißt, für die gilt $\|f\|_\infty := \sup\{\|f(x)\| : x \in X\} < \infty$. Ist Y eine normierte Algebra über $\mathbb{K}$, dann ist auch $C^b(X, Y)$ eine normierte Algebra über $\mathbb{K}$.

(c) Mit $C_0(X, \mathbb{K})$ wird die Unter algebra über $\mathbb{K}$ von $C^b(X, \mathbb{K})$ der im Unendlichen verschwindenden Funktionen bezeichnet. Dabei heißt eine Funktion $f \in C^b(X, \mathbb{K})$ im Unendlichen verschwindend, wenn für jedes $\varepsilon > 0$ eine kompakte Teilmenge K von X existiert mit $|f(x)| < \varepsilon$ für alle x außerhalb von K .

(d) Versehen mit der Supremumsnorm $\|\cdot\|_\infty$, sind $C^b(X, \mathbb{K})$ und $C_0(X, \mathbb{K})$ kommutative Banachalgebren über $\mathbb{K}$. Ist X kompakt, dann gilt $C(X, \mathbb{K}) = C^b(X, \mathbb{K}) = C_0(X, \mathbb{K})$.

(e) (DALES [56](2000), Seite 157). $C_0(X, \mathbb{K})$ ist genau dann unital, wenn X kompakt ist.

(f) Die Definitionen von im vorliegenden Text zwar verwendeten, aber nicht definierten Begriffen aus der Maßtheorie finden sich bei FREMLIN [99](2000-2008); so etwa σ -Algebra von Teilmengen einer vorgegebenen Menge (ebd., 111A), Maßraum (X, Σ, μ) , messbare Mengen, Maß μ auf X (ebd., 112A), Borel- σ -Algebra (ebd., 4A3A), Lebesguemaß (ebd., 115E); des Weiteren für beliebige Maßräume (X, Σ, μ) und $p \in \{0\} \cup [1, \infty]$ mit hier unten geschriebenem Index „ p “ die Räume $\mathcal{L}_p(X, \Sigma, \mu)$, kurz $\mathcal{L}_p(\mu)$, und $L_p(X, \Sigma, \mu)$, kurz $L_p(\mu)$, und deren Komplexifizierungen $\mathcal{L}_{p(\mathbb{C})}(X, \Sigma, \mu)$, kurz $\mathcal{L}_{p(\mathbb{C})}(\mu)$, und $L_{p(\mathbb{C})}(X, \Sigma, \mu)$, kurz $L_{p(\mathbb{C})}(\mu)$, (ebd., Kapitel 24; Komplexifizierungen: ebd., 241J, 242P, 243K; siehe auch 354Y1). Die ebd. in der dortigen Symbolik konsequent angewandte Unterscheidung von Funktionen und deren Äquivalenzklassen wird hier nicht verwendet. Bezüglich des Falles $0 < p < 1$ sei an die Bemerkung in 2.4.15 erinnert. Falls nicht ausdrücklich etwas anderes vereinbart ist, meint für $X = \mathbb{R}^n$ der Ausdruck $L_p(X)$ stets $L_p(X, \Sigma, \lambda)$ mit

Σ die σ -Algebra der Lebesgue-messbaren Teilmengen des $\mathbb{R}^n$ und λ das Lebesguemaß auf $\mathbb{R}^n$; insbesondere ist dann für jedes $A \in \Sigma$ mit dem Ausdruck $L_p(A)$ der Raum $L_p(A, \Sigma_A, \lambda_A)$ mit $\Sigma_A := \{E \in \Sigma : E \subseteq A\}$ und $\lambda_A := \lambda \upharpoonright \Sigma_A$ (ebd., 131B) gemeint. $L_p(\mathbb{C})(A)$, $\mathcal{L}_p(A)$ und $\mathcal{L}_{p(\mathbb{C})}(A)$ entsprechend.

2.4.18 Minimale, aber nicht schiefminimale Idempotente. Wie das folgende Beispiel zeigt, muss ein minimal idempotentes Element eines Ringes nicht notwendig auch schiefminimal idempotent sein.

Sei dazu X ein zusammenhängender, nicht leerer topologischer Raum und sei $\mathbb{Z}$ mit der diskreten Topologie versehen. Dann betrachte den Ring $C(X, \mathbb{Z})$, versehen mit der in 2.1.45 definierten partiellen Ordnung $\leq$. Bezeichne e seine Eins $X \rightarrow \mathbb{Z}$, $x \mapsto 1$. Da $C(X, \mathbb{Z})$ nur aus den konstanten Abbildungen von X nach $\mathbb{Z}$ besteht, gilt $C(X, \mathbb{Z}) \simeq \mathbb{Z}$. Wegen $eC(X, \mathbb{Z})e = C(X, \mathbb{Z})$ ist e nicht schiefminimal idempotent. e ist aber minimal idempotent, denn: Es gilt zwar für jedes $f \in C(X, \mathbb{Z})$ offensichtlich $f = ef = fe$. Da in $\mathbb{Z}$ aber $x^2 = x$ äquivalent zu $x \in \{0, 1\}$ ist, sind nur die Nullfunktion $X \rightarrow \mathbb{Z}$, $x \mapsto 0$ und e idempotent.

Ein weiteres Beispiel liefert der Funktionenraum $C([0, 1], \mathbb{K})$. Offensichtlich ist in ihm das einzige, von der Nullfunktion verschiedene, idempotente Element die konstante Abbildung $x \mapsto 1$ und dieses ist somit minimal idempotent, aber zweifellos nicht schiefminimal idempotent.

2.4.19 Produkte von normierten Räumen (ABRAMOVICH und ALIPRANTIS [2](2002), Seite 281; CONWAY [52](1990), Seite 72; MATHIEU [208](1998), Seite 68; MEGGINSON [211](1998), Seiten 61 und 491; PEDERSEN [233](1979), Seite 6; PEDERSEN [234](1989), Seite 49).

Sei X_ν , $\nu \in I$, eine Familie von normierten Vektorräumen, alle über den gleichen Körper $\mathbb{K}$. Für jedes $p \in \mathbb{R}$, $1 \leq p < \infty$, definiert man die Abbildung $\prod_{\nu \in I} X_\nu \rightarrow \mathbb{R} \cup \{\infty\}$,

$$\|x\|_p := \left(\sup \left\{ \sum_{\nu \in F} \|x_\nu\|^p : F \subseteq I \text{ endlich} \right\} \right)^{1/p}$$

für alle $x := (x_\nu)_{\nu \in I} \in \prod_{\nu \in I} X_\nu$, wobei natürlich $\infty^{1/p}$ als ∞ gesetzt wird. Die entsprechende Abbildung $\prod_{\nu \in I} X_\nu \rightarrow \mathbb{R} \cup \{\infty\}$ für $p = \infty$ definiert man als

$$\|x\|_\infty := \sup_{\nu \in I} \|x_\nu\|$$

für alle $x := (x_\nu)_{\nu \in I} \in \prod_{\nu \in I} X_\nu$.

Auf dem Vektorraum $\bigoplus_{\nu \in I} X_\nu$ ist für jedes $1 \leq p \leq \infty$ die Abbildung $\|\cdot\|_p$ eine Norm; sie wird *p-Norm* genannt.

Für jedes $1 \leq p \leq \infty$ definiert man den Vektorraum

$$\ell_p\left(\bigoplus_{\nu \in I} X_\nu\right) := \left(\bigoplus_{\nu \in I} X_\nu\right)_p := \left\{x \in \prod_{\nu \in I} X_\nu : \|x\|_p < \infty\right\}.$$

Für jedes $1 \leq p \leq \infty$ ist die Abbildung $\|\cdot\|_p$ auf $\ell_p(\bigoplus_{\nu \in I} X_\nu)$ eine Norm und der dadurch entstehende normierte Vektorraum $\ell_p(\bigoplus_{\nu \in I} X_\nu)$ ist genau dann vollständig, wenn alle X_ν , $\nu \in I$, vollständig sind.

Des Weiteren setzt man:

$$c_0\left(\bigoplus_{\nu \in I} X_\nu\right) := \left(\bigoplus_{\nu \in I} X_\nu\right)_0 := \left\{(x_\nu)_{\nu \in I} \in \prod_{\nu \in I} X_\nu : \left(\begin{array}{c} I \rightarrow \mathbb{R}, \\ \nu \mapsto \|x_\nu\| \end{array}\right) \in C_0(I, \mathbb{R})\right\},$$

wobei die Indexmenge I mit der *diskreten Topologie* versehen sei. Als ein Unterraum von $\ell_\infty(\bigoplus_{\nu \in I} X_\nu)$ trägt der Raum $c_0(\bigoplus_{\nu \in I} X_\nu)$ die ∞ -Norm. Ist $I = \mathbb{N}$, so hat man

$$c_0\left(\bigoplus_{n \in \mathbb{N}} X_n\right) = \left\{(x_n)_{n \in \mathbb{N}} \in \prod_{n \in \mathbb{N}} X_n : \lim_{n \rightarrow \infty} \|x_n\| = 0\right\};$$

auch dieser Vektorraum ist genau dann vollständig, wenn alle X_n , $n \in \mathbb{N}$, vollständig sind.

Sind alle X_ν , $\nu \in I$, vollständig, dann ist für jedes $1 \leq p < \infty$ der Banachraum $\ell_p(\bigoplus_{\nu \in I} X_\nu)$ die Vervollständigung von $(\bigoplus_{\nu \in I} X_\nu, \|\cdot\|_p)$, während $c_0(\bigoplus_{\nu \in I} X_\nu)$ die Vervollständigung von $(\bigoplus_{\nu \in I} X_\nu, \|\cdot\|_\infty)$ ist.

Sei X ein Banachraum über $\mathbb{K}$. Dann gilt für jede beliebige Indexmenge I , wenn man $X_\nu := X$ für alle $\nu \in I$ setzt:

$$\ell_\infty\left(\bigoplus_{\nu \in I} X_\nu\right) = C^b(I, X) \quad \text{und} \quad c_0\left(\bigoplus_{\nu \in I} X_\nu\right) = C_0(I, X).$$

Speziell für $X = \mathbb{K}$ setzt man dann

$$\ell_p(I) := \ell_p\left(\bigoplus_{\nu \in I} X_\nu\right) \quad \text{für alle } 1 \leq p \leq \infty \quad \text{und} \quad c_0(I) := c_0\left(\bigoplus_{\nu \in I} X_\nu\right);$$

und weiter

$$\ell_p := \ell_p(\mathbb{N}) \quad \text{für alle } 1 \leq p \leq \infty \quad \text{und} \quad c_0 := c_0(\mathbb{N}),$$

um dabei übliche Folgenräume ℓ_p , $1 \leq p \leq \infty$, und c_0 zu erhalten. Bezeichnet v das Element von ℓ_∞ mit $v_n = 1$ für alle $n \in \mathbb{N}$, so wird wie üblich der Folgenraum $\mathbb{K} \cdot v + c_0$ aller konvergenten Folgen mit c bezeichnet.

Seien $X_1, \dots, X_n$ normierte Vektorräume, alle über den gleichen Körper $\mathbb{K}$. Ist $1 \leq p \leq \infty$, so schreibt man anstelle von $\ell_p(\prod_{k=1}^n X_k)$ auch $X_1 \oplus_p \dots \oplus_p X_n$. Ist $n \in \mathbb{N}^\times$ beliebig, aber fest, so sind auf $\prod_{k=1}^n X_k$ die Normen $\|\cdot\|_p$, $1 \leq p \leq \infty$, alle zueinander äquivalent und induzieren jeweils die Produkttopologie. Für $n \in \mathbb{N}^\times$, $1 \leq p \leq \infty$ und $1 \leq j \leq n$ ist die kanonische Einbettung von X_j in $\prod_{k=1}^n X_k$ isometrisch isomorph zu X_j und in $\ell_p(\prod_{k=1}^n X_k)$ abgeschlossen.

Sei X_n , $n \in \mathbb{N}$, eine Folge von normierten Vektorräumen, alle über den gleichen Körper $\mathbb{K}$. Sei $1 \leq p < \infty$ und $1 < q \leq \infty$ mit $p + q = pq$. Dann gilt

$$\ell_q\left(\bigoplus_{n \in \mathbb{N}} X_n^*\right) \cong \ell_p\left(\bigoplus_{n \in \mathbb{N}} X_n\right)^*.$$

Für endliche Summen gilt dies auch für die Kombination $p = \infty$, $q = 1$.

2.4.20. Sei X ein normierter Vektorraum über $\mathbb{K}$. Sei $p \in \mathbb{R}$ mit $1 \leq p \leq \infty$. Seien U und V zwei abgeschlossene Untervektorräume von X mit $X = U \oplus_p V$. Dann ist die Abbildung $T: X/U \rightarrow V$, $x + U \mapsto v$, wobei $x = u + v$, $u \in U$, $v \in V$ gilt, eine lineare surjektive Isometrie. In Zeichen gilt also $X/U \cong V$.

Beweis. Die Linearität und Surjektivität sind klar. Sei $x \in X$. Schreibe $x = u + v$ für gewisse $u \in U$, $v \in V$. Dann gilt: $\|x + U\| = \inf\{\|x - w\| : w \in U\} = \inf\{\|u + v - w\| : w \in U\} = \inf\{\|(u - w) + v\| : w \in U\} = \inf\{\|w + v\| : w \in U\}$. Der letzte Term ist für $1 \leq p < \infty$ gleich $\inf\{(\|w\|^p + \|v\|^p)^{1/p} : w \in U\}$ und für $p = \infty$ gleich $\inf\{\max\{\|w\|, \|v\|\} : w \in U\}$. Dies ist in beiden Fällen gleich $\|v\|$, also gleich $\|T(x + U)\|$. $\square$

2.4.21 Extremalpunkte. Sei X ein normierter Vektorraum über $\mathbb{K}$ mit $X = U \oplus_1 V$ für zwei jeweils von $\{0\}$ verschiedenen Untervektorräumen U und V von X . Dann gilt $\text{ex } B_X = \text{ex } B_U \cup \text{ex } B_V$.

Beweis. „ $\subseteq$ “: Sei $x \in \text{ex } B_X$. Also $x = u + v$ für zwei gewisse $u \in B_U$ und $v \in B_V$. x ist ein Punkt von S_X . Angenommen, es wäre $u \neq 0$ und $v \neq 0$. Wegen $1 = \|x\| = \|u\| + \|v\|$,

also $0 < \|u\| < 1$ und $0 < \|v\| < 1$. Dann könnte man aber x schreiben als $x = u + v = \|u\| \frac{u}{\|u\|} + \|v\| \frac{v}{\|v\|}$, also $x = \lambda \frac{u}{\|u\|} + (1 - \lambda) \frac{v}{\|v\|}$ mit $\lambda = \|u\|$, im Widerspruch dazu, dass x ein Extrempunkt von B_X ist. Es ist also entweder $u = 0$ oder $v = 0$. Liege der Fall $v = 0$ vor. Also $x \in B_U$. Als Extrempunkt von B_X ist x nun a fortiori ein Extrempunkt von B_U . Der andere Fall entsprechend.

„ $\supseteq$ “: Sei $x \in \text{ex } B_U$. Seien $a, b \in B_X$ und sei $0 < \lambda < 1$ mit $x = \lambda a + (1 - \lambda)b$. Es ist $a = a_U + a_V$ und $b = b_U + b_V$ für gewisse $a_U, b_U \in B_U$ und $a_V, b_V \in B_V$. Somit $x = \lambda a_U + (1 - \lambda)b_U + \lambda a_V + (1 - \lambda)b_V$. Da $x \in U$, gilt $\lambda a_V + (1 - \lambda)b_V = 0$, also $x = \lambda a_U + (1 - \lambda)b_U$ und nach Voraussetzung folgt $x = a_U = b_U$. Somit $\|a_U\| = \|b_U\| = 1$, $\|a_V\| = \|a\| - \|a_U\| \leq 1 - 1 = 0$, $a_V = 0$, analog $b_V = 0$, also $x = a = b$, das heißt, $x \in \text{ex } B_X$. Für $x \in \text{ex } B_V$ entsprechend. $\square$

2.4.22 Orthogonalität. Ist X ein topologischer Vektorraum über $\mathbb{K}$, so bezieht sich, wenn nicht ausdrücklich etwas anderes vereinbart ist, der in 2.1.4 definierte Begriff *orthogonal* wie allgemein hin üblich auf das Dualsystem (X, X^*) .

In der Theorie der Banachräume ist auch die folgende Definition üblich: Ist X ein normierter Vektorraum über $\mathbb{K}$, so heißen $x, y \in X$ *orthogonal*, wenn für alle $\lambda \in \mathbb{K}$ die Ungleichung $\|x\| \leq \|x + \lambda y\|$ gilt; siehe DIESTEL [65](1975), Seite 24.

ALFSEN und SHULTZ [8](2001) definieren zwei Elemente x und y des positiven Kegels gewisser normierter Vektorräume über $\mathbb{R}$ (im Englischen sogenannte *base norm spaces*, ebd., Seite 9) als *orthogonal*, wenn $\|x - y\| = \|x\| + \|y\|$ gilt. NEAL und RUSSO [223](2004) definieren zwei Elemente eines normierten Vektorraumes über $\mathbb{C}$ als *orthogonal*, falls $\|x + y\| = \|x - y\| = \|x\| + \|y\|$ gilt.

2.4.23 Ein Trennungskriterium. Man betrachte das Dualsystem (X, X^*) für einen topologischen Vektorraum X über $\mathbb{K}$. Gemäß Bemerkungen 2.2.34 ist (X, X^*) in X^* trennend. Ist X ein lokal konvexer Hausdorffraum, so ist nach einem Korollar (WERNER [291](2002), Seite 379, Korollar VIII.2.13) des Satzes von Hahn-Banach das Dualsystem (X, X^*) auch in X trennend. Siehe auch 2.4.29.

2.4.24 Ultraprodukte (HEINRICH [130](1980); JOHNSON und LINDENSTRAUSS [158](2001), Seiten 55 und 455). Sei I eine Menge und $\mathcal{U}$ ein Ultrafilter auf I . Sei mit diesen X_ν , $\nu \in I$, eine Familie von Banachräumen, alle über den gleichen Körper $\mathbb{K}$. Setze

$$N_{\mathcal{U}} := \left\{ (x_\nu)_{\nu \in I} \in \ell_\infty \left(\bigoplus_{\nu \in I} X_\nu \right) : \lim_{\mathcal{U}} \|x_\nu\| = 0 \right\}.$$

Ein $(x_\nu)_{\nu \in I} \in \ell_\infty(\bigoplus_{\nu \in I} X_\nu)$ ist also genau dann in $N_{\mathcal{U}}$, wenn für jede Nullumgebung $U \subseteq \mathbb{R}$ die Menge $\{\nu \in I : \|x_\nu\| \in U\}$ ein Element des Filters $\mathcal{U}$ ist. $N_{\mathcal{U}}$ ist ein abgeschlossener Unterraum von $\ell_\infty(\bigoplus_{\nu \in I} X_\nu)$. Der Quotient $\ell_\infty(\bigoplus_{\nu \in I} X_\nu)/N_{\mathcal{U}}$, versehen mit der kanonischen Quotientennorm, heißt das *Ultraprodukt* (genauer *Banachraum-Ultraprodukt*) (bezüglich $N_{\mathcal{U}}$) und wird mit $(\prod_{\nu \in I} X_\nu)_{\mathcal{U}}$ bezeichnet. Ist $(x_\nu)_{\nu \in I}$ ein Element von $\ell_\infty(\bigoplus_{\nu \in I} X_\nu)$, so wird die dazugehörige Äquivalenzklasse im Ultraprodukt $(\prod_{\nu \in I} X_\nu)_{\mathcal{U}}$ mit $(x_\nu)_{\mathcal{U}}$ bezeichnet; es gilt $\|(x_\nu)_{\mathcal{U}}\| = \lim_{\mathcal{U}} \|x_\nu\|$. Sind alle X_ν , $\nu \in I$, gleich einem Raum X , so wird dieses Ultraprodukt mit $X_{\mathcal{U}}^I$ oder einfach auch nur mit $X_{\mathcal{U}}$ (bei DINEEN $X^{\mathcal{U}}$) bezeichnet und *Ultrapotenz* (im Englischen *ultrapower*, im Deutschen aber üblicherweise auch Ultraprodukt) genannt.

2.4.25 Definition. Sei $(X, Y; B)$ ein Dualsystem. Mit den Bezeichnungen von 2.2.22 wird die Initialtopologie auf X für die Funktionenfamilie $\text{ran}(R)$ die von Y auf X erzeugte *schwache Topologie* (im Französischen *la topologie faible*) genannt und wird mit $\sigma(X, Y)$ bezeichnet. Diese Topologie ist lokal konvex.

Wenn nicht etwas anderes vereinbart ist, ist für $Y \subseteq X^\#$ die Bilinearform B immer die Restriktion der kanonischen Bilinearform von X auf $X \times Y$, und für B wird dann auch $\langle \cdot, \cdot \rangle$ geschrieben.

Ist X ein topologischer Vektorraum über $\mathbb{K}$, so heißt die Topologie $w := \sigma(X, X^*)$ die *schwache Topologie* von X und die Topologie $w^* := \sigma(X^*, X)$ die *schwach*-Topologie* von X^* .

2.4.26 (KELLEY und NAMIOKA et al. [185](1963), Seite 138). Sei $(X, Y; B)$ ein Dualsystem. Betrachte den Produktraum $P := \prod_{\nu \in Y} X_\nu$ mit $X_\nu = \mathbb{K}$ für alle $\nu \in Y$, versehen mit der Produkttopologie. P ist der Vektorraum $\mathbb{K}^Y$ aller $\mathbb{K}$ -wertigen Funktionen auf Y , versehen mit der Topologie der punktweisen Konvergenz. P ist vollständig, lokal konvex und Hausdorff'sch mit $Y^\#$ als einen abgeschlossenen Unterraum. Sei S analog zu 2.2.22 die Abbildung $X \rightarrow P$, $x \mapsto B(x, \cdot)$. Dann ist ein $O \subseteq X$ genau dann $\sigma(X, Y)$ -offen, wenn es ein offenes $U \subseteq P$ mit $O = S^{-1}(U)$ gibt.

Beweis. Sei R die Abbildung aus 2.2.22. Sei $O \subseteq X$ ein Element der $\sigma(X, Y)$ -Subbasis von X . Dann existiert nach Definition 2.4.25 ein $y \in Y$ und ein offenes $K \subseteq \mathbb{K}$ mit $O = R(y)^{-1}(K)$. Also $O = \{x \in X : B(x, y) \in K\}$. Somit $O = \{x \in X : S(x) \in U\}$ mit $U = \{f \in P : f(y) \in K\}$. Also $O = S^{-1}(U)$ und da U ein typisches Element einer Subbasis von P ist, gilt die Behauptung. $\square$

Man beachte, dass für eine $\sigma(X, Y)$ -abgeschlossene Menge $A \subseteq X$ zwar $S(A)$ abgeschlossen in $S(X)$ ist, aber nicht unbedingt $S(A)$ in P abgeschlossen zu sein braucht.

2.4.27 Schwache Umgebungen. (a) Sei $(X, Y; B)$ ein Dualsystem. Gemäß Definition 2.4.25 ist für alle $\varepsilon > 0$, allen $x \in X$ und allen $y \in Y$ das Urbild unter $R(y) = B(\cdot, y)$ von der offenen ε -Kugel um $B(x, y)$ eine $\sigma(X, Y)$ -Umgebung von x ; sie wird mit $U_{y, \varepsilon}(x)$ bezeichnet.

(b) Sei $(X, Y; B)$ ein Dualsystem. Für jedes $y \in Y$ ist $p_y : x \mapsto |B(x, y)|$ eine Halbnorm auf X .

(c) Sei X ein Vektorraum über $\mathbb{K}$. Sei P eine Menge von Halbnormen auf X . Für jede endliche Teilmenge $F \subseteq P$ und für jedes $\varepsilon > 0$ setze man $U_{F, \varepsilon} := \{x \in X : p(x) \leq \varepsilon \quad \forall p \in F\}$ und $V_{F, \varepsilon} := \{x \in X : p(x) < \varepsilon \quad \forall p \in F\}$. Dann bildet sowohl die Menge aller solcher $U_{F, \varepsilon}$ als auch die Menge aller solcher $V_{F, \varepsilon}$ jeweils eine Nullumgebungssubbasis genau einer Vektorraumtopologie, die *die von P erzeugte Topologie* heißt; diese Topologie wird mit $\sigma(X, P)$ bezeichnet, ist lokal konvex und gleich der Initialtopologie der Funktionenfamilie $\{x \mapsto p(x - y) : p \in P, y \in X\}$. Dass die $U_{F, \varepsilon}$ und die $V_{F, \varepsilon}$ Nullumgebungssubbasen ein und derselben Topologie sind, wird in TAYLOR und LAY [275](1980), Seite 107, bewiesen.

(d) Sei X ein Vektorraum über $\mathbb{K}$, Q eine Menge von Halbnormen auf X und dann $P \subseteq Q$. Dann gilt $\sigma(X, P) \subseteq \sigma(X, Q)$.

(e) Jede lokal konvexe Topologie auf einem Vektorraum X über $\mathbb{K}$ kann durch eine Menge von stetigen Halbnormen auf X erzeugt werden.

(f) Als Folgerung hat man, siehe zum Beispiel HORVÁTH [142](1966), Seite 98: Seien X und Y zwei lokal konvexe Vektorräume über $\mathbb{K}$. Sei P bzw. Q eine Menge von stetigen Halbnormen auf X bzw. Y , die die lokal konvexe Topologie von X bzw. Y erzeugt. Dann ist eine lineare Abbildung $T : X \rightarrow Y$ genau dann stetig, wenn für alle $q \in Q$ ein nicht leeres $F \subseteq P$ und ein $\lambda > 0$ existiert, so dass für alle $x \in X$ die Ungleichung $q(Tx) \leq \lambda \cdot \max\{p(x) : p \in F\}$ erfüllt ist.

2.4.28 Schwach kompakte Abbildungen. Seien X und Y zwei lokal konvexe Räume über $\mathbb{K}$. Ein $T \in L(X, Y)$ heißt *schwach kompakt*, wenn es eine Nullumgebung von X gibt, deren Bild unter T in Y relativ schwach kompakt ist. Die Menge der schwach kompakten

Elemente von $L(X, Y)$ wird mit $K_w(X, Y)$ bezeichnet. Man setzt $K_w(X) := K_w(X, X)$. Es gilt $K(X, Y) \subseteq K_w(X, Y) \subseteq \mathcal{L}(X, Y)$.

2.4.29 Schwach und Hausdorff. (a) Sei $(X, Y; B)$ ein Dualsystem. Die schwache Topologie $\sigma(X, Y)$ ist genau dann Hausdorff'sch, wenn das Dualsystem trennend in X ist. Mit 2.4.23 folgt daher insbesondere: Für jeden topologischen Vektorraum X über $\mathbb{K}$ ist die Topologie $\sigma(X^*, X)$ auf X^* Hausdorff'sch.

(b) (HEUSER [133](1975), Seiten 333 und 353). Sei P eine Menge von Halbnormen auf einen Vektorraum X über $\mathbb{K}$. P wird *total* genannt, wenn aus $p(x) = 0$ für alle $p \in P$ stets $x = 0$ folgt. Die durch P erzeugte lokal konvexe Topologie ist genau dann Hausdorff'sch, wenn P total ist.

(c) Sei $(X, Y; B)$ ein Dualsystem. Jede $\sigma(X, Y)$ -stetige Linearform auf X kann als $B(\cdot, y)$ für ein $y \in Y$ geschrieben werden. Das y ist eindeutig, wenn das Dualsystem in Y trennend ist, und in diesem Fall kann man also Y mit der Menge aller $\sigma(X, Y)$ -stetigen Linearformen auf X identifizieren. Insbesondere gilt für den Fall, dass Y ein Untervektorraum von X^* ist, dass ein Funktional $\ell \in X^\#$ genau dann $\sigma(X, Y)$ -stetig ist, wenn $\ell \in Y$ gilt. Siehe auch Satz 2.4.35.

(d) (BOURBAKI [35](1987), II.6.2). Sei $(X, Y; B)$ ein Dualsystem und U ein Untervektorraum von X . Dann ist U genau dann $\sigma(X, Y)$ -dicht in X , wenn das Dualsystem (U, Y) in Y trennend ist. (Siehe auch 2.5.9.)

(e) (KADISON und RINGROSE [167], Seite 29). Sei X ein Vektorraum über $\mathbb{K}$ und S eine Teilmenge von $X^\#$. Die gemäß den Bemerkungen 2.4.27(b) und (c) durch die Halbnormen $x \mapsto |\langle x, s \rangle|$, $s \in S$, erzeugte lokal konvexe Topologie wird mit $\sigma(X, S)$ bezeichnet. Es gilt: $\sigma(X, S) = \sigma(X, \text{span}(S))$. Trennt S die Punkte von X , so gilt zusätzlich $(X, \sigma(X, S))^* = \text{span}(S)$.

(f) (BOURBAKI [35](1987), II.6.2). Als eine Folgerung aus (d) hat man: Sei X ein Vektorraum über $\mathbb{K}$ und U ein Unterraum über $\mathbb{K}$ von $X^\#$. Sei B die Restriktion der kanonischen Bilinearform von X auf $X \times U$. Dann gilt: Das Dualsystem $(X, U; B)$ ist genau dann trennend, wenn U $\sigma(X^\#, X)$ -dicht in $X^\#$ ist. Siehe auch 2.5.10.

(g) Als eine Folgerung von (c) und 2.4.23 hat man: Ist X ein lokal konvexer Hausdorffraum, so gilt $X = (X^*, \sigma(X^*, X))^*$. Insbesondere trägt also X die schwach*-Topologie $\sigma((X^*, \sigma(X^*, X))^*, X^*)$; diese ist aber offensichtlich identisch mit der schwachen Topologie von X .

2.4.30 Zueinander duale Abbildungen. Seien (X, Y) und (V, W) zwei Dualsysteme. Die beiden dazugehörigen Bilinearformen seien hier beide mit $\langle \cdot, \cdot \rangle$ bezeichnet. Seien $T: X \rightarrow V$ und $R: W \rightarrow Y$ zwei lineare Abbildungen.

(a) (KELLEY und NAMIOKA [185](1963), Seite 199; BOURBAKI [35](1987), II.6.4). Die Abbildungen T und R heißen *dual* (zueinander), wenn für alle $x \in X, w \in W$ die Gleichung

$$\langle T(x), w \rangle = \langle x, R(w) \rangle$$

erfüllt ist.

Zu der Abbildung T existiert genau dann eine Abbildung S , die dual zu T ist, wenn T $\sigma(X, Y) - \sigma(V, W)$ -stetig ist, und dies wiederum ist genau dann der Fall, wenn für jedes $w \in W$ das lineare Funktional $x \mapsto \langle T(x), w \rangle$ $\sigma(X, Y)$ -stetig ist. Sind die Abbildungen T und R zueinander dual, so ist also $R: \sigma(W, V) - \sigma(Y, X)$ -stetig; des Weiteren ist R genau dann die einzige Abbildung, die zu T dual ist, wenn das Dualsystem (X, Y) in Y trennend ist.

(b) (SCHAEFER [255](1967), Seite 128). Liege nun der Fall vor, dass das Dualsystem (X, Y) in Y und das Dualsystem (V, W) in W trennend ist. Gemäß 2.2.31 sei Y dann mit einem Untervektorraum von $X^\#$ und W mit einem Untervektorraum von $V^\#$ identifiziert.

Dann ist die Abbildung T genau dann $\sigma(X, Y) - \sigma(V, W)$ -stetig, wenn

$$T^\#(W) \subseteq Y$$

gilt; offensichtlich ist in diesem Fall $T^\# \upharpoonright W$ eine Abbildung, die zu T dual ist.

Beweis. Sei $T^\#(W) \subseteq Y$ und $w \in W$. Betrachte auf X die Funktion $g: x \mapsto \langle T(x), w \rangle$, also $g(x) = \langle x, (T^\# \upharpoonright W)(w) \rangle = \langle x, y \rangle$ für alle $x \in X$ und für $y := T^\#(w)$. Nach Voraussetzung ist $y \in Y$, das heißt, die Funktion g ist $\sigma(X, Y)$ -stetig. Nach (a) ist somit die Abbildung T $\sigma(X, Y) - \sigma(V, W)$ -stetig.

Sei nun umgekehrt die Abbildung T $\sigma(X, Y) - \sigma(V, W)$ -stetig. Sei $w \in W$. Betrachte auf X wieder die Funktion $g: x \mapsto \langle T(x), w \rangle$, also $g(x) = \langle x, T^\#(w) \rangle$ für alle $x \in X$. Da $T^\#(w) \in X^\#$, ist die Funktion g $\sigma(X, X^\#)$ -stetig, insbesondere also $\sigma(X, Y)$ -stetig und nach 2.4.29(c) ist somit $T^\#(w) \in Y$. Da dann offensichtlich für alle $x \in X$, $w \in W$ die Gleichung $\langle T(x), w \rangle = \langle x, (T^\# \upharpoonright W)(w) \rangle$ gilt, ist nach (a) $T^\# \upharpoonright W$ die Abbildung, die zu T dual ist. $\square$

2.4.31 Bemerkung. Seien X und Y zwei normierte Vektorräume über $\mathbb{K}$ und sei M ein Untervektorraum von X^* . Sei T ein isometrischer Isomorphismus von X auf Y . Setze $N := (T^*)^{-1}(M)$. Dann ist wegen $(T^*)^{-1} = (T^{-1})^*$ (MEGGINSON [211](1998), Satz 3.1.15) $N = (T^{-1})^*(M)$ ein Untervektorraum von Y^* und man kann das Dualsystem (Y, N) betrachten. Für alle $x \in X$, $m \in M$ gilt: $\langle x, m \rangle = \langle T^{-1}Tx, m \rangle = \langle Tx, (T^{-1})^*m \rangle$.

2.4.32 Bemerkung und Definition. Seien X und Y zwei normierte Vektorräume über $\mathbb{K}$. Ein Dualsystem $(X, Y; B)$ wird *stetig* genannt, wenn die Bilinearform B stetig ist.

Ist $(X, Y; B)$ ein stetiges Dualsystem, so gilt für die Abbildungen S und R von 2.2.22: $\text{ran}(S) \subseteq Y^*$ und $\text{ran}(R) \subseteq X^*$. Daher definiert man für ein stetiges Dualsystem $(X, Y; B)$ die Abbildungen S und R als:

$$\begin{aligned} S: X &\rightarrow Y^*, & x &\mapsto B(x, \cdot); \\ R: Y &\rightarrow X^*, & y &\mapsto B(\cdot, y). \end{aligned}$$

Die Abbildungen S und R sind linear und stetig. Wenn die Abbildung S eine surjektive Isometrie ist, dann heißt X eine *Dual von Y bezüglich der Bilinearform B* . Analog heißt Y eine *Dual von X bezüglich der Bilinearform B* , wenn die Abbildung R eine surjektive Isometrie ist.

2.4.33 Satz (KÖTHE [187](1966), Seite 238). *Sei (X, τ) ein lokal konvexer Hausdorffraum. Ist $U \subseteq X$ ein τ -abgeschlossener Unterraum, so ist U orthogonalabgeschlossen bezüglich des Dualsystems (X, X^*) . Jeder bezüglich (X, X^*) orthogonalabgeschlossene Unterraum $U \subseteq X$ ist sogar $\sigma(X, X^*)$ -abgeschlossen in X . Speziell für normierte Vektorräume gilt also, dass genau die abgeschlossenen Unterräume die orthogonalabgeschlossenen Unterräume sind.*

2.4.34 Lemma. *Sei X ein normierter Vektorraum über $\mathbb{K}$, $U \subseteq X^*$ ein Untervektorraum und $\ell \in X^\#$. Sei $\ell \upharpoonright B_X$ $\sigma(X, U)$ -stetig. Dann ist $\ell \in X^*$. Sei $\varepsilon > 0$, dann gibt es $\ell_1, \dots, \ell_n \in U$ mit der Eigenschaft, dass für alle $x \in B_X$ mit $|\ell_1(x)| + \dots + |\ell_n(x)| < 1$ die Ungleichung $|\ell(x)| < \varepsilon$ gilt. Daraus folgt: $|\ell(x)| \leq \varepsilon\|x\| + \|\ell\| \cdot (|\ell_1(x)| + \dots + |\ell_n(x)|)$ für alle $x \in X$.*

ℓ lässt sich darstellen als $\ell = \varphi_1 + \varphi_2$ mit $\varphi_1 \in X^$, $\|\varphi_1\| \leq \varepsilon$ und $\varphi_2 \in U$.*

Beweis. (KADISON und RINGROSE [167](1991), Übung 1.9.14.) Da $\ell \upharpoonright B_X$ $\sigma(X, U)$ -stetig ist, ist $\ell \upharpoonright B_X$ stetig. Insbesondere ist $\ell \upharpoonright B_X$ bei 0 stetig. Es gibt also ein $t > 0$, so

dass für alle $x \in t \cdot B_X$ die Ungleichung $|\ell(x)| < 1$ gilt. Somit gilt für alle $x \in B_X$ (also $tx \in t \cdot B_X$) die Ungleichung $|\ell(x)| < \frac{1}{t}$, das heißt, $\|\ell\| = \sup \{|\ell(x)| : x \in B_X\} \leq \frac{1}{t}$ und somit $\ell \in X^*$. Da $\ell \upharpoonright B_X$ bei 0 $\sigma(X, U)$ -stetig ist, gibt es für gegebenes $\varepsilon > 0$ eine $\sigma(X, U)$ -Nullumgebung $U_{F, \delta} = \{x \in B_X : |\varphi(x)| < \delta, \varphi \in F\}$ mit einem $F = \{f_1, \dots, f_n\} \subseteq U$, $n \in \mathbb{N}^\times$, und $\delta > 0$, so dass aus $x \in U_{F, \delta}$ die Ungleichung $|\ell(x)| < \varepsilon$ folgt, oder mit anderen Worten, aus $x \in U_{\frac{1}{n\delta}F, \frac{1}{n}}$ die Ungleichung $|\ell(x)| < \varepsilon$ folgt. Die ℓ_ν , $\nu = 1, \dots, n$, in der Behauptung erhält man also per $\ell_\nu := \frac{1}{n\delta}f_\nu$, $\nu = 1, \dots, n$.

Für die nächste Ungleichung genügt es wegen der Homogenität sie nur für $\|x\| = 1$ zu zeigen. Gilt dabei noch $|\ell_1(x)| + \dots + |\ell_n(x)| < 1$, so folgt sie direkt aus der davor gezeigten Ungleichung. Gilt neben $\|x\| = 1$ aber $|\ell_1(x)| + \dots + |\ell_n(x)| \geq 1$, so ist sie sowieso klar.

Die durch $p_1(x) := \varepsilon\|x\|$ und $p_2(x) := \|\ell\| \cdot (|\ell_1(x)| + \dots + |\ell_n(x)|)$ definierten Halbnormen p_1 und p_2 erfüllen die Voraussetzungen des Lemmas 2.2.20, wonach es also $\varphi_1, \varphi_2 \in X^\#$ gibt mit $\ell = \varphi_1 + \varphi_2$, $|\varphi_1(x)| \leq \varepsilon\|x\|$ und $|\varphi_2(x)| \leq \|\ell\| \cdot (|\ell_1(x)| + \dots + |\ell_n(x)|)$ für alle $x \in X$. Somit $\varphi_1 \in X^*$, $\|\varphi_1\| \leq \varepsilon$ und, da φ_2 auf dem Durchschnitt aller Kerne der $\ell_1, \dots, \ell_n$ verschwindet, ist φ_2 eine Linearkombination der $\ell_1, \dots, \ell_n$ (siehe zum Beispiel WERNER [291](2002), Seite 381, Lemma VIII.3.3), das heißt, $\varphi_2 \in U$. $\square$

Anschließend an die Bemerkung 2.4.29(c) hat man:

2.4.35 Satz (STRĂTILĂ und ZSIDÓ [270](1979), Seite 14). *Sei X ein Banachraum über $\mathbb{K}$, $U \subseteq X^*$ ein Unterraum und $\ell \in X^\#$. Dann gilt:*

- (a) $\ell \upharpoonright B_X$ ist $\sigma(X, U)$ -stetig $\Leftrightarrow \ell \in \text{cl}(U)$.
- (b) Die Topologien $\sigma(X, U)$ und $\sigma(X, \text{cl}(U))$ stimmen auf B_X überein.
- (c) Wenn U abgeschlossen ist und $\ell \upharpoonright B_X$ $\sigma(X, U)$ -stetig ist, dann ist ℓ $\sigma(X, U)$ -stetig.

Beweis. (a): „ $\Rightarrow$ “: Sei $\ell \upharpoonright B_X$ $\sigma(X, U)$ -stetig. Nach Lemma 2.4.34 ist $\ell \in X^*$ und für jedes $\varepsilon > 0$ existiert ein $\varphi_1 \in X^*$ mit $\|\varphi_1\| \leq \varepsilon$ und ein $\varphi_2 \in U$ mit $\ell = \varphi_1 + \varphi_2$, das heißt, $\|\ell - \varphi_2\| \leq \varepsilon$. Somit ist ℓ ein Berührungspunkt von U und damit ein Element von $\text{cl}(U)$. „ $\Leftarrow$ “: Sei $\ell \in \text{cl}(U)$. Es existiert also eine Folge $(\ell_n)_{n \in \mathbb{N}}$ mit $\ell_n \in U$ für alle $n \in \mathbb{N}$ und $\sup_{x \in B_X} |\ell(x) - \ell_n(x)| \rightarrow 0$ für $n \rightarrow \infty$. Sei $x_0 \in B_X$ und $\varepsilon > 0$. Sei $n_0 \in \mathbb{N}$ derart, dass $\|\ell - \ell_n\| < \frac{\varepsilon}{3}$ für alle $n \geq n_0$. Sei V eine $\sigma(X, U)$ -Umgebung von x_0 , so dass für alle $x \in V \cap B_X$ die Ungleichung $|\ell_n(x_0) - \ell_n(x)| < \frac{\varepsilon}{3}$ gilt. Dann gilt für alle $x \in V \cap B_X$ und $n \geq n_0$: $|\ell(x_0) - \ell(x)| = |\ell(x_0) - \ell_n(x) - \ell(x) + \ell_n(x)| \leq |\ell(x_0) - \ell_n(x)| + |\ell(x) - \ell_n(x)| = |\ell(x_0) - \ell_n(x_0) - \ell_n(x) + \ell_n(x_0)| + |\ell(x) - \ell_n(x)| < |\ell(x_0) - \ell_n(x_0)| + |\ell_n(x_0) - \ell_n(x)| + \frac{\varepsilon}{3} < \frac{\varepsilon}{3} + \frac{\varepsilon}{3} + \frac{\varepsilon}{3} = \varepsilon$, das heißt, $\ell \upharpoonright B_X$ ist $\sigma(X, U)$ -stetig.

(b): Es ist zu zeigen, dass $\sigma(X, U)|_{B_X}$ feiner als $\sigma(X, \text{cl}(U))|_{B_X}$ ist. Betrachte gemäß Satz 1.2.17(j) eine Subbasis von $\sigma(X, \text{cl}(U))|_{B_X}$ bei 0, etwa diejenige, die aus allen Mengen der Form $\{x \in B_X : |\ell(x)| < 1\}$, $\ell \in \text{cl}(U)$, besteht. Sei V solch ein Subbasiselement für ein $\ell \in \text{cl}(U)$. Nach (a) ist $\ell \upharpoonright B_X$ $\sigma(X, U)$ -stetig, insbesondere bei 0. Es existiert also eine $\sigma(X, U)$ -Nullumgebung W mit $|\ell(0) - \ell(x)| < 1$ für alle $x \in W \cap B_X$, also $|\ell(x)| < 1$ für alle $x \in W \cap B_X$. Das heißt, $W \cap B_X \subseteq V$, und V umfasst ein Element einer $\sigma(X, U)|_{B_X}$ -Basis bei 0. $\square$

2.4.36 Satz (KELLEY und NAMIOKA et al. [185](1963), Seite 120). *Sei X ein topologischer Vektorraum über $\mathbb{K}$ und U ein Unterraum von X . Dann gilt:*

- (a) Die Abbildung $T: (X/U)^* \rightarrow U^\perp$, $y^* \mapsto x^* := (x \mapsto y^*(x + U))$ ist ein Vektorraum-Isomorphismus. Es gilt $x^*x = (Ty^*)(x) = y^*(x + U)$ für alle $x \in X$. Wenn X normiert und U abgeschlossen ist, dann ist T eine Isometrie.

- (b) Die Abbildung $T: X^*/U^\perp \rightarrow U^*$, $x^* + U^\perp \mapsto x^* \upharpoonright U$ ist ein Vektorraum-Monomorphismus. Falls X lokal konvex ist, hat man als Fortsetzungsversion des Satzes von Hahn-Banach die Aussage, dass T ein Vektorraum-Isomorphismus ist. Wenn X normiert ist, dann ist T eine Isometrie.

Man beachte, dass die Abbildung T in (b) von Satz 2.4.36 für nicht lokal konvexe Vektorräume über $\mathbb{K}$ im Allgemeinen nicht surjektiv ist, während die analoge Abbildung des Dualsystems $(X, X^\#)$ für jeden Vektorraum X über $\mathbb{K}$ eine Vektorraum-Isomorphie ist; siehe 2.2.36.

2.4.37 Schwache Produkttopologien (CHOQUET [46](1969), Seite 56; BOURBAKI [35](1987), II.6.6). Sei (X_ν, Y_ν) , $\nu \in I$, eine Familie von Dualsystemen, alle über den gleichen Körper $\mathbb{K}$. Setze $X := \prod_{\nu \in I} X_\nu$, $Y := \bigoplus_{\nu \in I} Y_\nu$ und $\langle x, y \rangle := \sum_{\nu \in I} \langle x_\nu, y_\nu \rangle$ für alle $x = (x_\nu)_{\nu \in I} \in X$ und $y = (y_\nu)_{\nu \in I} \in Y$. Dann ist $(X, Y; \langle \cdot, \cdot \rangle)$ ein Dualsystem und $\sigma(X, Y)$ ist die Produkttopologie der Topologien $\sigma(X_\nu, Y_\nu)$, $\nu \in I$, auf den X_ν , $\nu \in I$. Im Allgemeinen ist $\sigma(Y, X)$ nicht die von $\prod_{\nu \in I} Y_\nu$ induzierte Relativtopologie.

Das Dualsystem (X, Y) ist genau dann trennend in X (bzw. Y), wenn für alle $\nu \in I$ die Dualsysteme (X_ν, Y_ν) trennend in X_ν (bzw. Y_ν) sind. Wenn das Dualsystem (X, Y) in X (bzw. in Y) trennend ist, dann ist ein Untervektorraum U von X (bzw. Y), der orthogonal zu einem Y_μ (bzw. X_μ) (kanonisch identifiziert mit einem Untervektorraum von Y (bzw. X)) ist, ein Untervektorraum von $\prod_{\nu \in I \setminus \{\mu\}} X_\nu$ (bzw. $\bigoplus_{\nu \in I \setminus \{\mu\}} Y_\nu$).

2.4.38 (CHOQUET [46](1969), Seite 57). Aus 2.4.37 folgt: Sei X_ν , $\nu \in I$, eine Familie von lokal konvexen Vektorräumen über $\mathbb{K}$ und X_ν^* , $\nu \in I$, ihre Dualräume. Dann ist der Dualraum von $X := \prod_{\nu \in I} X_\nu$, versehen mit der Produkttopologie, gleich $\bigoplus_{\nu \in I} (X_\nu^*)$. Siehe auch 2.4.16.

2.4.39 (BOURBAKI [35](1987), II.6.6). Ebenfalls aus 2.4.37 folgt: Sei (X, Y) ein trennendes Dualsystem. Wenn $(X, \sigma(X, Y))$ die topologische direkte Summe zweier Unterräume U und V ist, dann ist $(Y, \sigma(Y, X))$ die topologische direkte Summe der beiden Unterräume $U^\perp$ und $V^\perp$.

2.4.40 $(\tau_X - \overline{\tau_Y})$ -halbstetig (GILES, GREGORY und SIMS [108](1978); KAMENSKII, OBUKHOVSKII und ZECCA [168](2001), Seite 6).

(a) Sei (X, τ_X) ein topologischer Raum, (Y, τ_Y) ein topologischer Vektorraum über $\mathbb{K}$, $D \subseteq X$, $x_0 \in D$ und $f: D \rightarrow \mathcal{P}(Y)$ eine Abbildung. Die Abbildung f heißt $(\tau_X - \overline{\tau_Y})$ -*oberhalbstetig* im Punkt x_0 , wenn für jede τ_Y -Nullumgebung V eine τ_X -Umgebung U von x_0 existiert, so dass für alle $x \in U \cap D$ die Inklusion $f(x) \subseteq f(x_0) + V$ erfüllt ist.

Die Abbildung f heißt $(\tau_X - \overline{\tau_Y})$ -*unterhalbstetig* im Punkt x_0 , wenn für jede τ_Y -Nullumgebung V eine τ_X -Umgebung U von x_0 existiert, so dass für alle $x \in U \cap D$ die Inklusion $f(x_0) \subseteq f(x) + V$ erfüllt ist.

Man bemerke: Falls Y mit einer Metrik d versehen ist, die die Topologie τ_Y induziert, dass dann, da die $U(0; r)$, $r > 0$, eine Umgebungsbasis von Y bilden, die Abbildung f genau dann $(\tau_X - \overline{\tau_Y})$ -oberhalbstetig im Punkt x_0 ist, wenn für jedes $r > 0$ eine τ_X -Umgebung U von x_0 existiert, so dass für alle $x \in U \cap D$ die Inklusion $f(x) \subseteq f(x_0) + U(0; r)$ erfüllt ist.

(b) Sei wie bei (a) wieder (X, τ_X) ein topologischer Raum, (Y, τ_Y) ein topologischer Vektorraum über $\mathbb{K}$. Sei $x_0 \in X$ und $f: X \rightarrow \mathcal{P}(Y)$ eine Abbildung. Es sei an 1.2.8 erinnert. Ist nun die Abbildung f oberhalbstetig im Punkt x_0 , so ist sie auch $(\tau_X - \overline{\tau_Y})$ -oberhalbstetig im Punkt x_0 .

Beweis. Sei f oberhalbstetig im Punkt x_0 . Sei V eine τ_Y -Nullumgebung. Dann ist $G := f(x_0) + V$ eine offene Menge, für die offensichtlich $f(x_0) \subseteq G$ gilt. Nach Voraussetzung existiert eine τ_X -Umgebung U von x_0 , so dass für alle $x \in U$ die Inklusion $f(x) \subseteq G$,

also $f(x) \subseteq f(x_0) + V$, erfüllt ist. Das heißt, f ist $(\tau_X - \overline{\tau_Y})$ -oberhalbstetig im Punkt x_0 . $\square$

Falls $f(x_0)$ kompakt ist, gilt auch die umgekehrte Implikation.

Beweis. Sei f im Punkt x_0 $(\tau_X - \overline{\tau_Y})$ -oberhalbstetig und $f(x_0)$ kompakt. Sei $G \subseteq Y$ offen mit $f(x_0) \subseteq G$. Nach 2.4.4 existiert eine Nullumgebung V mit $f(x_0) + V \subseteq G$ und somit nach Voraussetzung dazu auch eine Umgebung U von x_0 , so dass für alle $x \in U$ die Inklusion $f(x) \subseteq f(x_0) + V$, insbesondere also $f(x) \subseteq G$ gilt. $\square$

(c) Sei (X, τ_X) ein topologischer Raum, (Y, d) ein metrischer Raum, $D \subseteq X$, $x_0 \in D$ und $f: D \rightarrow \mathcal{P}(Y)$ eine Abbildung. Die Abbildung f heißt $(\tau_X - \overline{d})$ -oberhalbstetig im Punkt x_0 , wenn für jedes $\varepsilon > 0$ eine τ_X -Umgebung U von x_0 existiert, so dass für alle $x \in U \cap D$ die Inklusion $f(x) \subseteq \{y \in Y : \text{dist}(y, f(x_0)) < \varepsilon\}$ erfüllt ist. Für die $(\tau_X - \overline{d})$ -Unterhalbstetigkeit bei x_0 wird entsprechend die Inklusion $f(x_0) \subseteq \{y \in Y : \text{dist}(y, f(x)) < \varepsilon\}$ gefordert.

(d) Sei wie bei (c) wieder (X, τ_X) ein topologischer Raum und (Y, d) ein metrischer Raum. Sei $x_0 \in X$ und $f: X \rightarrow \mathcal{P}(Y)$ eine Abbildung. Ähnlich wie bei (b) sieht man: Ist die Abbildung f oberhalbstetig im Punkt x_0 , so ist sie auch $(\tau_X - \overline{d})$ -oberhalbstetig im Punkt x_0 .

Beweis. Sei f oberhalbstetig im Punkt x_0 und $\varepsilon > 0$. Die Menge $G := \{y \in Y : \text{dist}(y, f(x_0)) < \varepsilon\}$ ist offen und umfasst offensichtlich $f(x_0)$. Sofort wie bei (b) sieht man hier die $(\tau_X - \overline{d})$ -Oberhalbstetigkeit von f im Punkte x_0 . $\square$

Falls $f(x_0)$ kompakt ist, gilt auch die umgekehrte Implikation.

Beweis. Sei f im Punkt x_0 $(\tau_X - \overline{d})$ -oberhalbstetig und $f(x_0)$ kompakt. Sei $G \subseteq Y$ offen mit $f(x_0) \subseteq G$. Als eine abgeschlossene Teilmenge der kompakten Menge $f(x_0)$ ist der Rand $\partial(f(x_0))$ von $f(x_0)$ kompakt. Betrachte die Abstandsfunktion $g := \text{dist}(\cdot, \partial G) \upharpoonright \partial(f(x_0))$. Da jeder Punkt des Randes $\partial(f(x_0))$ ein innerer Punkt von G ist, gilt für alle $y \in \partial(f(x_0))$ die Ungleichung $g(y) > 0$. Als eine stetige Funktion auf einer kompakten Menge nimmt g ihr Minimum an, etwa ε_0 . Wegen $g > 0$ ist $\varepsilon_0 > 0$. Somit liegt die offene Menge $\{y \in Y : \text{dist}(y, f(x_0)) < \varepsilon_0\}$ in G . Nach Voraussetzung folgt, dass eine τ_X -Umgebung U von x_0 existiert, so dass für alle $x \in U$ die Inklusion $f(x) \subseteq \{y \in Y : \text{dist}(y, f(x_0)) < \varepsilon_0\}$ erfüllt ist, insbesondere also $f(x) \subseteq G$ gilt. Das heißt, f ist oberhalbstetig im Punkt x_0 . $\square$

(e) Sei wie bei (d) wieder (X, τ_X) ein topologischer Raum, (Y, d) ein metrischer Raum, $x_0 \in X$, $f: X \rightarrow \mathcal{P}(Y)$ eine Abbildung. Dann gilt: Ist die Abbildung f $(\tau_X - \overline{d})$ -unterhalbstetig im Punkt x_0 , so ist sie auch unterhalbstetig im Punkt x_0 .

Beweis. Sei f $(\tau_X - \overline{d})$ -unterhalbstetig im Punkt x_0 . Sei $G \subseteq Y$ offen mit $f(x_0) \cap G \neq \emptyset$. Sei $y \in f(x_0) \cap G$ und dann $\varepsilon > 0$ derart, dass $U(y; \varepsilon) \subseteq G$. Nach Voraussetzung gibt es eine τ_X -Umgebung U von x_0 , so dass für alle $x \in U$ die Inklusion $f(x_0) \subseteq \{z \in Y : \text{dist}(z, f(x)) < \varepsilon\}$ gilt. Wegen $y \in f(x_0)$, gilt also für jedes $x \in U$ die Ungleichung $\text{dist}(y, f(x)) < \varepsilon$, das heißt, $f(x) \cap U(y; \varepsilon) \neq \emptyset$ und somit auch $f(x) \cap G \neq \emptyset$. $\square$

Falls $f(x_0)$ kompakt ist, gilt auch die umgekehrte Implikation.

Beweis. Siehe KAMENSKII, OBUKHOVSKII und ZECCA ebd. $\square$

2.5 Der Bipolarensatz

2.5.1 Definition (HOLMES [137](1975); KÖTHE [187](1966); REED und SIMON [245](1980); WERNER [291](2002); KELLEY und NAMIOKA et al. [185](1963), Seite 141; ROBERTSON und ROBERTSON [247](1967)).

Sei $(X, Y; B)$ ein Dualsystem, $M \subseteq X$ und $N \subseteq Y$. Die *Polare* von M ist

$$M^\circ := \{y \in Y : \operatorname{Re} B(x, y) \leq 1 \quad \forall x \in M\},$$

die *Polare* von N ist

$$N^\circ := \{x \in X : \operatorname{Re} B(x, y) \leq 1 \quad \forall y \in N\}.$$

Die *Absolutpolare* von M ist

$$M^\bullet := \{y \in Y : |B(x, y)| \leq 1 \quad \forall x \in M\},$$

die *Absolutpolare* von N ist

$$N^\bullet := \{x \in X : |B(x, y)| \leq 1 \quad \forall y \in N\}.$$

2.5.2 Bemerkung. Es sei an die Definition 1.1.3 erinnert, wonach in der Situation von Definition 2.5.1 die Polaren die rechten und linken Träger bezüglich der Relation $\{(x, y) \in X \times Y : \operatorname{Re} B(x, y) \leq 1\}$, und die Absolutpolaren die rechten und linken Träger bezüglich der Relation $\{(x, y) \in X \times Y : |B(x, y)| \leq 1\}$ sind. Des Weiteren sei auf die somit gemäß 1.1.19 geltenden Aussagen für die Polaren- und Absolutpolarenbildung hingewiesen!

2.5.3 (KÖTHE [187](1966), Seite 246). Sei $(X, Y; B)$ ein Dualsystem und $M \subseteq X$. Es gilt $M^\bullet \subseteq M^\circ$. Gilt $S_{\mathbb{K}} \cdot M \subseteq M$, was zum Beispiel der Fall ist, wenn M kreisförmig ist oder gar absolutkonvex ist, so gilt $M^\bullet = M^\circ$.

Beweis. Gelte $S_{\mathbb{K}} \cdot M \subseteq M$ und sei $x \in M$ und $y \in M^\circ$. Also $\operatorname{Re} B(x, y) \leq 1$, und da für jedes $\lambda \in S_{\mathbb{K}}$ auch λx in M ist, gilt $\operatorname{Re} B(\lambda x, y) \leq 1$ für alle $\lambda \in S_{\mathbb{K}}$. Wegen $B(\lambda x, y) = \lambda B(x, y)$ gilt demnach $|\lambda B(x, y)| \leq 1$ für alle $\lambda \in S_{\mathbb{K}}$. Insbesondere für $\lambda = 1$ folgt $|B(x, y)| \leq 1$, das heißt, $y \in M^\bullet$. $\square$

2.5.4. Sei X ein normierter Vektorraum über $\mathbb{K}$. Dann gilt $B_X^\bullet = B_X^\circ = B_{X^*}$.

2.5.5 (KÖTHE [187](1966), Seite 247; BOURBAKI [35](1987), II.6.3; TAYLOR und LAY [275](1980), Seite 161).

Sei $(X, Y; B)$ ein Dualsystem. Seien M und N Teilmengen von X . Offensichtlich gilt $M^\bullet = (B_{\mathbb{K}} \cdot M)^\bullet$ und somit nach 2.5.3 folglich $M^\bullet = (B_{\mathbb{K}} \cdot M)^\circ$, das heißt, die Absolutpolare ist gleich der Polaren der kreisförmigen Hülle. M° ist stets konvex und $M^\bullet$ ist stets absolutkonvex. Daraus folgt $M^{\bullet\bullet} = M^{\circ\circ} = (B_{\mathbb{K}} \cdot M)^{\circ\circ}$. M° und $M^\bullet$ sind $\sigma(Y, X)$ -abgeschlossen.

Für alle $\lambda \in \mathbb{K}^\times$ gilt $(\lambda \cdot M)^\circ = \frac{1}{\lambda} \cdot M^\circ$ und $(\lambda \cdot M)^\bullet = \frac{1}{\lambda} \cdot M^\bullet$. Wenn M die Menge N absorbiert, dann absorbiert N° die Menge M° . Für jede Familie M_ν , $\nu \in I$, I eine Indexmenge, von Teilmengen von X gilt die Inklusion

$$\left(\bigcap_{\nu \in I} M_\nu \right)^\circ \supseteq \operatorname{cl} \left(\sigma(Y, X); \operatorname{co} \left(\bigcup_{\nu \in I} M_\nu^\circ \right) \right);$$

ein Kriterium für Gleichheit findet man in Korollar 2.5.7.

2.5.6 Bipolarensatz. Sei $(X, Y; B)$ ein Dualsystem und $M \subseteq X$. Dann gilt: $M^{\circ\circ} = cl(\sigma(X, Y); co(\{0\} \cup M))$ und $M^{\bullet\bullet} = cl(\sigma(X, Y); co(\{0\} \cup B_{\mathbb{K}} \cdot M))$, mit anderen Worten ist $M^{\bullet\bullet}$ die kleinste absolutkonvexe $\sigma(X, Y)$ -abgeschlossene Menge, die M umfasst.

Beweis. Hier sei nur auf die Aussage „mit anderen Worten“ eingegangen: Sie folgt aus der Tatsache, dass von einer kreisförmigen Menge sowohl der Abschluss innerhalb eines topologischen Vektorraumes als auch die konvexe Hülle wieder kreisförmig ist, siehe TAYLOR und LAY [275](1980), Seite 102. $\square$

2.5.7 Korollar (KELLEY und NAMIOKA et al. [185](1963), Seite 142). In der Situation der letzten Mengeninklusion von 2.5.5 gilt Gleichheit, also $\left(\bigcap_{\nu \in I} M_{\nu}\right)^{\circ} = cl\left(\sigma(Y, X); co\left(\bigcup_{\nu \in I} M_{\nu}^{\circ}\right)\right)$, wenn alle M_{ν} , $\nu \in I$, $I \neq \emptyset$, $\sigma(X, Y)$ -abgeschlossen und konvex sind und die $0 \in X$ enthalten.

Ist M_{ν} , $\nu \in I$, I eine nicht leere Indexmenge, eine Familie von absolutkonvexen $\sigma(X, Y)$ -abgeschlossenen Teilmengen von X , so gilt $\left(\bigcap_{\nu \in I} M_{\nu}\right)^{\bullet} = \left(\bigcap_{\nu \in I} M_{\nu}^{\bullet\bullet}\right)^{\bullet} = \left(\left(\bigcup_{\nu \in I} M_{\nu}^{\bullet}\right)^{\bullet}\right)^{\bullet} = \left(\bigcup_{\nu \in I} M_{\nu}^{\bullet}\right)^{\bullet\bullet}$, also $\left(\bigcap_{\nu \in I} M_{\nu}\right)^{\bullet} = \left(B_{\mathbb{K}} \cdot \left(\bigcup_{\nu \in I} M_{\nu}^{\bullet}\right)\right)^{\circ\circ} = \left(\bigcup_{\nu \in I} M_{\nu}^{\bullet}\right)^{\circ\circ} = cl\left(\sigma(Y, X); co\left(\bigcup_{\nu \in I} M_{\nu}^{\bullet}\right)\right)$.

2.5.8 Normierender Unterraum. Sei X ein normierter Vektorraum über $\mathbb{K}$ und U ein Untervektorraum von X^* . Dann ist U genau dann ein für X normierender Raum, wenn B_U schwach*-dicht in B_{X^*} ist.

Beweis. Betrachte das Dualsystem (X^*, X) . Sei U normierend für X . Wegen der Gleichungskette $(B_{X^*})_{\circ}^{\circ} = B_X^{\circ} = B_{X^*}$ ist B_{X^*} schwach*-abgeschlossen. Daher ist $cl(w^*; B_U) \subseteq B_{X^*}$. Angenommen, es gäbe ein $x^* \in B_{X^*}$ mit $x^* \notin cl(w^*; B_U)$. Nach dem Satz von Hahn-Banach gibt es dann ein schwach*-stetiges $\ell \in X^{**}$, ein $c \in \mathbb{R}$ und ein $\varepsilon > 0$ mit $\operatorname{Re} \ell(u^*) \leq c$ für alle $u^* \in cl(w^*; B_U)$ und $\operatorname{Re} \ell(x^*) \geq c + \varepsilon$. Da $(X^*, w^*)^* = i_X(X)$ gilt, existiert ein $x \in X$ mit $\ell(y^*) = y^*(x)$ für alle $y^* \in X^*$. Insbesondere gilt also für alle $u^* \in B_U$ die Ungleichung $\operatorname{Re} u^*(x) \leq c$, und da B_U zirkulär ist, sogar $|u^*(x)| \leq c$. Da U als normierend vorausgesetzt wurde, ist somit $\|x\| \leq c$. Somit gilt $|x^*(x)| \leq c \cdot \|x^*\| \leq c$ im Widerspruch zu $\operatorname{Re} x^*(x) \geq c + \varepsilon$.

Sei nun umgekehrt B_U schwach*-dicht in B_{X^*} . Sei $x \in X$. Wegen $\|x\| = \max\{|x^*(x)| : x^* \in B_{X^*}\}$ existiert ein $\xi^* \in B_{X^*}$ mit $\|x\| = \xi^*(x)$. Sei $\varepsilon > 0$. Da B_U schwach*-dicht in B_{X^*} ist, existiert ein $u^* \in B_U$ mit $|u^*(x) - \xi^*(x)| < \varepsilon$, also $||u^*(x)| - |\xi^*(x)|| < \varepsilon$ und somit $||\|x\| - |u^*(x)|| < \varepsilon$. Da $\sup\{|x^*(x)| : x^* \in B_U\} \leq \|x\|$, folgt $\|x\| = \sup\{|x^*(x)| : x^* \in B_U\}$. $\square$

Es sei angemerkt, dass, wenn U normierend ist, das Dualsystem (U, X) in X trennend ist, was nach 2.4.29(d) äquivalent ist zu der schwach*-Dichtheit von U in X^* (die ja auch unmittelbar aus der eben gezeigten schwach*-Dichtheit von B_U in B_{X^*} folgt).

2.5.9 (HOLMES [137](1975), Seite 68). Als eine einfache Folgerung aus dem Bipolarensatz hat man die bereits als Spezialfall in 2.4.29(d) enthaltende Aussage:

Sei X ein topologischer Vektorraum über $\mathbb{K}$ und A eine Teilmenge von X . Dann ist $\operatorname{span}(A)$ genau dann $\sigma(X, X^*)$ -dicht in X , wenn $A^{\perp} = \{0\}$ bezüglich des Dualsystems (X, X^*) gilt, mit anderen Worten also genau dann, wenn das Dualsystem $(\operatorname{span}(A), X^*)$ in X^* trennend ist.

Beweis. Setze $M := \text{span}(A)$ und $N := \text{cl}(\sigma(X, X^*); M)$. Sei M $\sigma(X, X^*)$ -dicht in X , also $N = X$ und damit auch $N^\perp = X^\perp$. Da nach 2.4.23 (X, X^*) in X^* trennend ist, ist nach 2.2.29(e) $N^\perp = \{0\}$. Nach dem Bipolarensatz gilt $N = M^{\perp\perp}$, also $M^\perp = M^{\perp\perp\perp} = N^\perp = \{0\}$ und nach Gleichung (2.26) aus 2.2.27 folgt $A^\perp = \{0\}$.

Ist nun umgekehrt $A^\perp = \{0\}$, so ist wegen $M^\perp \subseteq A^\perp$ auch $M^\perp = \{0\}$. Also $M^{\perp\perp} = \{0\}^\perp = X$ und nach dem Bipolarensatz ist M $\sigma(X, X^*)$ -dicht in X . $\square$

2.5.10. Sei X ein topologischer Vektorraum über $\mathbb{K}$. Der Beweis der Folgerung 2.5.9 zeigt: Ist A eine Teilmenge von X^* , so folgt aus $A^\perp = \{0\}$ bezüglich des Dualsystems (X^*, X) die $\sigma(X^*, X)$ -Dichtheit von $\text{span}(A)$ in X^* . Hinreichend dafür, dass auch hier die Umkehrung gilt, ist, dass das Dualsystem (X, X^*) in X trennend ist, siehe dazu 2.4.23.

Eine an die Bemerkung 2.4.29(f) erinnernde Aussage für das Dualsystem (X, X^*) bekommt man mit 2.2.33, wonach dann eine Teilmenge A von X^* genau dann die Punkte von X trennt, wenn $A^\perp = \{0\}$ bezüglich des Dualsystems (X^*, X) gilt.

2.5.11 Satz von Goldstine. Für jeden normierten Vektorraum X über $\mathbb{K}$ gilt:

- (a) $i_X(B_X)$ ist schwach*-dicht in $B_{X^{**}}$.
- (b) $i_X(X)$ ist schwach*-dicht in X^{**} .

Beweis. Zu (a): (TAYLOR und LAY [275](1980), Seite 177, Theorem 10.7). Betrachte das Dualsystem (X^{**}, X^*) . Wegen $i_X(B_X)^\circ = B_{X^*}$ gilt mit 2.5.4 $i_X(B_X)^{\circ\circ} = B_{X^{**}}$ und mit dem Bipolarensatz folgt die Behauptung. Für einen Beweis von (a) direkt mittels des Satzes von Hahn-Banach siehe DUNFORD und SCHWARTZ [78](1958), Seite 424, Theorem 5. (Die dort verwendete Idee wurde hier bereits im Widerspruchsbeweis von 2.5.8 verwendet.)

Zu (b): Folgt unmittelbar aus (a) (ebd., Seite 425): Dazu bemerke man nur, dass der schwach*-Abschluss von $i_X(X)$ in X^{**} ein Untervektorraum von X^{**} ist, der nach (a) $B_{X^{**}}$ enthält. Man kann auch mit Netzen argumentieren (MEGGINSON [211](1998), Seite 233, Korollar 2.6.28): Gelte (a) und sei $x^{**} \in X^{**\times}$. Setze $y^{**} := x^{**}/\|x^{**}\|$. Dann existiert ein Netz (x_ν) in B_X mit $i_X(x_\nu) \xrightarrow{w^*} y^{**}$, also $i_X(\|x^{**}\| \cdot x_\nu) \xrightarrow{w^*} x^{**}$. Ein einfacher direkter Beweis von (b) mittels des Bipolarensatzes geht wie folgt: Betrachte die beiden Dualsysteme (X, X^*) und (X^*, X^{**}) . Dann gilt $\hat{X}^{\circ\circ} = \hat{X}_\perp^\perp = X^{\perp\perp} = \{0\}^\perp = X^{**}$ und mit dem Bipolarensatz folgt die Aussage (b). (Oder alternativ: Wegen $\hat{X}_\perp = \{0\}$ und dann mit der anfänglichen Folgerung von 2.5.10 mit $A := i_X(X) \subseteq (X^*)^*$.) Für eine andere Argumentation zu (b) siehe EDWARDS [88](1965), Seite 523. $\square$

2.5.12 Dualität und Annihilator. Besonders im Zusammenhang mit M -Summanden in 2.5.17 ist die folgende Aussage von Interesse. Sei X ein Banachraum über $\mathbb{K}$. Seien M und N abgeschlossene Unterräume von X . Dann gilt

$$X = M \dot{+} N \quad \Leftrightarrow \quad X^* = M^\perp \dot{+} N^\perp.$$

Beweis. „ $\Rightarrow$ “: Nach 2.4.15 ist die Projektion $P \in L(X)$ mit $M = \text{ran}(P)$ und $N = \ker(P)$ stetig; dann ist die zu P duale Abbildung P^* eine stetige Projektion von X^* mit $\text{ran}(P^*) = N^\perp$ und $\ker(P^*) = M^\perp$.

„ $\Leftarrow$ “: Ist $N^\perp = \{0\}$, dann klar mit $M = \{0\}$ und $N = X$. Seien im Folgenden also M und N nichttriviale Unterräume von X .

(a) $M \cap N = \{0\}$:

Sei $x \in M \cap N$ und $x^* \in X^*$. Nach Voraussetzung ist $x^* = m^\perp + n^\perp$ mit eindeutigen $m^\perp \in M^\perp$ und $n^\perp \in N^\perp$. Somit $x^*(x) = m^\perp(x) + n^\perp(x) = 0 + 0 = 0$. Da das Dualsystem (X, X^*) in X trennend ist (2.4.23), muss mit Interpretation 2.2.28(c) $x = 0$ sein.

(b) $M \dot{+} N$ ist dicht in X :

Nach den beiden Korollaren 2.5.9 und 2.6.11 (der Vorgriff sei hier gestattet, da es auch hier schon zeigbar ist; siehe die dortige Bemerkung im Beweis) ist die Dichtheit von $M \dot{+} N$ in X äquivalent zu der Aussage $(M \dot{+} N)^\perp = \{0\}$. Sei also $x^* \in X^*$ mit $x^* \upharpoonright M \dot{+} N = 0$. Dann ist $x^* \upharpoonright M = 0$ und $x^* \upharpoonright N = 0$. Also $x^* \in M^\perp \cap N^\perp$ und nach Voraussetzung somit $x^* = 0$.

(c) $M \dot{+} N$ ist abgeschlossen in X :

(i) Die Abbildungen $X^* \rightarrow X^*$, $m^\perp + n^\perp \mapsto m^\perp$ beziehungsweise $n^\perp$ sind stetig:

Es ist $\|m^\perp + n^\perp\| \leq \|m^\perp\| + \|n^\perp\|$; also ist die Abbildung $id: M^\perp \oplus_1 N^\perp \rightarrow (M^\perp \dot{+} N^\perp, \|\cdot\|_{X^*})$ stetig. Nach einem Korollar (WERNER [291](2002), Seite 153, Korollar IV.3.4) des Satzes von der offenen Abbildung ist dann auch id^{-1} stetig, id also eine Isomorphie und somit insbesondere nach unten beschränkt, das heißt, es existiert ein $c > 0$ mit $c \cdot (\|m^\perp\| + \|n^\perp\|) \leq \|m^\perp + n^\perp\|$. Folglich gilt $\|m^\perp\| \leq \frac{1}{c} \|m^\perp + n^\perp\|$ und $\|n^\perp\| \leq \frac{1}{c} \|m^\perp + n^\perp\|$.

(ii) $P: m + n \mapsto m$ ist stetig auf $M \dot{+} N$:

$$\begin{aligned} \|m\| &= \sup_{\|x^*\|=1} |x^*(m)| = \sup_{\|x^*\|=1} |m^\perp(m) + n^\perp(m)| = \sup_{\|x^*\|=1} |n^\perp(m)| \\ &= \sup_{\|x^*\|=1} |n^\perp(m+n)| \leq \sup_{\|x^*\|=1} \|n^\perp\| \|m+n\| \stackrel{(i)}{\leq} \sup_{\|x^*\|=1} \frac{1}{c} \|m^\perp + n^\perp\| \|m+n\| = \\ &= \frac{1}{c} \|m+n\|. \end{aligned}$$

(iii) $(M \dot{+} N, \|\cdot\|_X) \simeq M \oplus_1 N$:

Sei $m \in M$, $n \in N$. Dann gilt $\|n\| \leq \|m+n\| + \|m\| \stackrel{(ii)}{\leq} \|m+n\| + \frac{1}{c} \|m+n\| = (1 + \frac{1}{c}) \|m+n\|$. Nochmals (ii) gibt $\|m\| + \|n\| \leq (1 + \frac{2}{c}) \|m+n\|$. Also $(1 + \frac{2}{c})^{-1} (\|m\| + \|n\|) \leq \|m+n\| \leq \|m\| + \|n\|$.

(iv) Nach WERNER [291](2002), Seite 35, Satz I.3.3, ist $M \oplus_1 N$ vollständig und mit (iii) ist $M \dot{+} N$ in X vollständig, das heißt, mit WERNER [291](2002), Seite 4, Lemma I.1.3(b), ist $M \dot{+} N$ in X abgeschlossen. $\square$

Anmerkung zum Beweis: Hat man nicht nur $X^* = M^\perp \dot{+} N^\perp$ gegeben, sondern sogar $X^* = M^\perp \oplus_1 N^\perp$, so gilt natürlich (c)(i) ohne erst den Satz von der offenen Abbildung bemühen zu müssen.

2.5.13 Satz von Alaoglu-Bourbaki. Sei X ein topologischer Vektorraum über $\mathbb{K}$ und U eine Nullumgebung in X . Dann sind $U^\bullet$ und $U^\circ = \sigma(X^*, X)$ -kompakt.

Beweis. Siehe zum Beispiel HEUSER [133](1975), Seite 373. Dort wird die Aussage für $U^\bullet$ bewiesen. Daraus folgt die Aussage für U° wie folgt: Wegen Bemerkung 2.4.3(d) enthält U eine kreisförmige Nullumgebung V ; $V^\circ = V^\bullet$ ist dann $\sigma(X^*, X)$ -kompakt und umfasst die $\sigma(X^*, X)$ -abgeschlossene Menge U° und somit ist U° ebenfalls $\sigma(X^*, X)$ -kompakt. $\square$

Wegen 2.5.4 ist insbesondere für jeden normierten Vektorraum X über $\mathbb{K}$ B_{X^*} $\sigma(X^*, X)$ -kompakt.

2.5.14 Schwache Quotiententopologien (BOURBAKI [35](1987), II.6.5 und KELLEY und NAMIOKA et al. [185](1963), Seite 147). Sei (X, Y) ein Dualsystem. Sei U ein Untervektorraum von X und V ein Untervektorraum von Y mit $\langle U, V \rangle = \{0\}$. Dann sind in kanonischer Weise $(U, Y/V)$ und $(X/U, V)$ zwei Dualsysteme (siehe auch 2.2.37) und es gilt:

(a) $\sigma(X, Y)|_U = \sigma(U, Y/V)$.

(b) Die von $\sigma(X, Y)$ auf X/U induzierte Quotiententopologie ist $\sigma(X/U, U^\perp)$.

- (c) Die von $\sigma(Y, X)$ auf Y/V induzierte Quotiententopologie ist $\sigma(Y/V, V^\perp)$. Offensichtlich gilt $\sigma(Y/V, U) = \sigma(Y/V, U + Y^\perp)$. Es gilt genau dann $\sigma(Y/V, U) = \sigma(Y/V, V^\perp)$, wenn $U + Y^\perp = V^\perp$ gilt. Falls das Dualsystem (X, Y) in X trennend ist, so ist die Topologie $\sigma(Y/U^\perp, U)$ genau dann gleich der von $\sigma(Y, X)$ auf $Y/U^\perp$ induzierten Quotiententopologie, wenn U $\sigma(X, Y)$ -abgeschlossen ist.
- (d) Speziell für $U = V^\perp$ hat man also: Die Quotiententopologie auf $X/V^\perp$ ist gleich $\sigma(X/V^\perp, \text{cl}(\sigma(Y, X); V))$.

Gemäß (c) ist $\sigma(X/V^\perp, V)$ echt schwächer als die von $\sigma(X, Y)$ auf $X/V^\perp$ induzierte Quotiententopologie, falls das Dualsystem (X, Y) in Y trennend ist und V nicht $\sigma(Y, X)$ -abgeschlossen ist.

Ist X ein lokal konvexer Vektorraum, so hat man mit (a) und Satz 2.4.36(b) wieder den Satz von Hahn-Banach vorliegen: $\sigma(X, X^*)|_U = \sigma(U, U^*)$.

2.5.15. Sei X ein Banachraum über $\mathbb{K}$ und U ein abgeschlossener Unterraum des Dualraumes X^* . Ist U reflexiv, so ist U schwach*-abgeschlossen.

Beweis. Sei U reflexiv. Nach WERNER [291](2002), Theorem VIII.3.18 ist dies äquivalent zu der $\sigma(U, U^*)$ -Kompaktheit von B_U . Nach der in 2.5.14 erwähnten Formulierung des Satzes von Hahn-Banach ist dies wiederum äquivalent zu der $\sigma(X^*, X^{**})$ -Kompaktheit von B_U . Da $\sigma(X^*, X)$ gröber als $\sigma(X^*, X^{**})$ ist, folgt die schwach*-Kompaktheit von B_U . Nach WERNER [291](2002), Korollar VIII.3.16 zum Satz von Krein-Smuljan, ist B_U schwach*-abgeschlossen. $\square$

2.5.16 Kontraktionen. Seien X und Y ein normierter Vektorraum über $\mathbb{K}$.

(a) Ein $T \in \mathcal{L}(X, Y)$ mit $\|T\| \leq 1$ heißt *kontraktiv* oder eine *Kontraktion*. Ist $p \in \mathcal{L}(X)$ eine Projektion, so gilt bekanntlich wegen $\|p\| = \|p^2\| \leq \|p\|^2$ entweder $\|p\| = 0$ oder $\|p\| \geq 1$. Eine Projektion $p \in \mathcal{L}(X)$ heißt *bikontraktiv*, falls sowohl p als auch $p^\perp$ kontraktiv sind.

(b) Für die in 2.4.9 definierte kanonische topologische Inklusionsabbildung i_X von X nach X^{**} , $x \mapsto (x^* \mapsto x^*x)$, hat man: Die Abbildung $i_{X^*} \circ i_X^*$ ist die kanonische Projektion von X^{***} auf X^* , wird mit π_{X^*} bezeichnet und auch als die *Dixmier-Projektion* angesprochen. Ihr Kern ist gleich $(i_X(X))^\perp$, bezogen auf das Dualsystem (X^{**}, X^{***}) ; siehe dazu auch HARMAND, WERNER und WERNER [126](1992), Seite 117, Lemma III.2.4. Es gilt

$$X^{***} = i_{X^*}(X^*) \dot{+} i_X(X)^\perp, \quad (2.31)$$

siehe DIXMIER [74](1948), Seite 1066, Theorem 15. Vergleiche auch mit Gleichung (2.33) in 2.5.17. Des Weiteren siehe auch 2.11.13.

(c) BROWN [39](1976) zeigt: Ist X der Folgenraum ℓ_1 , so gilt $\|(\pi_{X^*})^\perp\| = 2$. Der Folgenraum $X = \ell_1$ ist also ein Beispiel, in dem π_{X^*} nicht bikontraktiv ist. Für Weiteres zu bikontraktiven Abbildungen siehe 3.1.11.

(d) (HOLMES [137](1975), Seite 126; MEGGINSON [211](1998), Seite 232). Sei U ein Unterraum von X^* . Dann gilt $\|i_{X,U}\| \leq 1$.

Beweis. $\|i_{X,U}(x)\| = \|\ell_x\|$ für $\ell_x \in U^*$ mit $\ell_x(u) = u(x)$ für alle $u \in U$. Also $\|i_{X,U}(x)\| = \sup\{|u(x)| : u \in B_U\} \leq \sup\{|x^*x| : x^* \in B_{X^*}\} = \|x\|$. $\square$

$i_{X,U}$ ist $\sigma(X, X^*) - \sigma(U^*, U)$ -stetig, ja sogar $\sigma(X, U) - \sigma(U^*, U)$ -stetig.

Beweis. Sei (x_ν) ein Netz in X und $x \in X$. Dann gilt: $x_\nu \xrightarrow{w} x \Leftrightarrow \forall x^* \in X^* : \langle x_\nu, x^* \rangle_{(X, X^*)} \rightarrow \langle x, x^* \rangle_{(X, X^*)} \Rightarrow \forall u \in U : \langle x_\nu, u \rangle \rightarrow \langle x, u \rangle$ (sowohl in (X, X^*) als auch in (X, U)) $\Leftrightarrow i_{X,U}(x_\nu) \xrightarrow{w^*} i_{X,U}(x)$. $\square$

Sei S die Abbildung von 2.2.22 bezogen auf das Dualsystem (X, U) . Da dann die Abbildung $i_{X,U}$ gleich der Astriktion $U^* \upharpoonright S$ ist, ist wegen Bemerkung 2.2.28(a)(b) die Abbildung $i_{X,U}$ genau dann injektiv, wenn U die Punkte von X trennt und nach Korollar 2.5.9 ist dies wiederum genau dann der Fall, wenn U $\sigma(X^*, X)$ -dicht in X^* ist, oder dazu äquivalent, $U^\perp = \{0\}$ im Dualsystem (X^*, X) gilt. Da X ein Banachraum ist, gilt nach 2.4.23 $(X, \sigma(X, X^*))^* = X^*$. Setze $V := \text{ran}(i_{X,U})$. Nach 2.5.14(a) gilt $\sigma(U^*, U)|_V = \sigma(V, U)$. Wegen $V \subseteq U^\#$ ist nach 2.2.34 das Dualsystem (V, U) trennend in V . Daher gilt gemäß 2.4.29(e) $(V, \sigma(V, U))^* = U$. Somit ist die duale Abbildung von $V \upharpoonright i_{X,U}: (X, \sigma(X, X^*)) \rightarrow (V, \sigma(V, U))$ eine Abbildung von U nach X^* . Ist nun $V \upharpoonright i_{X,U}$ injektiv, so ist die besagte duale Abbildung $(V \upharpoonright i_{X,U})^*$ eine Bijektion von U auf X^* und in diesem Fall gilt für die eine im letzten Beweis vorkommende Implikation „ $\Rightarrow$ “ auch die Umkehrung „ $\Leftarrow$ “. Das heißt, falls $i_{X,U}$ injektiv ist, ist nicht nur

$$\text{ran}(i_{X,U}) \upharpoonright i_{X,U}: (X, \sigma(X, U)) \rightarrow (\text{ran}(i_{X,U}), \sigma(U^*, U)|_{\text{ran}(i_{X,U})}),$$

sondern auch

$$\text{ran}(i_{X,U}) \upharpoonright i_{X,U}: (X, \sigma(X, X^*)) \rightarrow (\text{ran}(i_{X,U}), \sigma(U^*, U)|_{\text{ran}(i_{X,U})})$$

ein Homöomorphismus.

2.5.17 L -Summanden und M -Ideale (ALFSEN und EFFROS [7] (1972), BEHRENDTS [21] (1979), HARMAND, WERNER und WERNER [126] (1992)).

Sei X ein Banachraum über $\mathbb{K}$ und $P \in L(X)$ eine Projektion. Wie bereits in 2.3.28 eingeführt, bezeichne dann $P^\perp$ die zu P komplementäre Projektion $\text{Id}_X - P$. Die Projektion P heißt eine L -Projektion, wenn für alle $x \in X$ die Gleichung $\|x\| = \|Px\| + \|P^\perp x\|$ gilt. P heißt eine M -Projektion, wenn für alle $x \in X$ die Gleichung $\|x\| = \max\{\|Px\|, \|P^\perp x\|\}$ gilt.

Genauso wie $\mathbb{K}$ im Allgemeinen als Platzhalter für $\mathbb{R}$ oder $\mathbb{C}$ steht, steht forthin $\mathbb{P}$ für den Buchstaben L oder M . Des Weiteren bezeichnet $\mathbb{P}^q$ den Buchstaben aus der Menge $\{L, M\} \setminus \{\mathbb{P}\}$.

Jede $\mathbb{P}$ -Projektion ist bikontraktiv.

Beweis. Sei P eine L -Projektion, dann ist P kontraktiv: $\|x\| = \|Px\| + \|x - Px\| \geq \|Px\| - \|x\| + \|Px\| \Leftrightarrow \|Px\| \leq \|x\|$. Da für jede Projektion $P \in L(X)$ die Gleichung $P = P^{\perp\perp}$ gilt, ist mit jeder $\mathbb{P}$ -Projektion P auch $P^\perp$ eine $\mathbb{P}$ -Projektion. $\square$

Ein Unterraum $J \subseteq X$ heißt ein $\mathbb{P}$ -Summand, wenn er das Bild einer $\mathbb{P}$ -Projektion $P \in L(X)$ ist. U ist in diesem Fall notwendigerweise abgeschlossen. Ein abgeschlossener Unterraum $J \subseteq X$ ist also genau dann ein $\mathbb{P}$ -Summand, wenn es einen abgeschlossenen Unterraum $K \subseteq X$ gibt, so dass X isometrisch isomorph zu $J \oplus_p K$ ist mit $p = 1$ im Fall von $\mathbb{P} = L$ und mit $p = \infty$ im Fall von $\mathbb{P} = M$; nach BEHRENDTS [21] (1979), Lemma 1.2(i), ist dabei der Raum K eindeutig durch J bestimmt; er wird der zu J komplementäre $\mathbb{P}$ -Summand genannt und mit $J^\perp$ bezeichnet. Nach ebd., Lemma 1.2(ii), gilt die eben angeführte genau-dann-Aussage auch, wenn man die Räume J und K nicht als abgeschlossen voraussetzt. Nach ebd., Lemma 1.4, stehen die $\mathbb{P}$ -Projektionen von X in eindeutiger Korrespondenz zu den $\mathbb{P}$ -Summanden von X .

Es gibt Banachräume, die außer den beiden trivialen M -Summanden $\{0\}$ und dem Raum selbst keine weiteren M -Summanden aufweisen. (Zum Beispiel Hilberträume über $\mathbb{R}$; ALFSEN und EFFROS [7] (1972), Seite 129.) Um in manchen Fällen trotzdem in X eine gewisse nicht triviale M -Struktur zu erkennen, schwächt man die Definition eines M -Summanden etwas ab, indem man einen Unterraum J von X sich als ähnlich wie ein M -Summand verhaltend erwartet, wenn der Annihilator $J^\perp$ sich im Dualraum X^* ähnlich verhält wie der Annihilator eines M -Summanden. Dabei fällt besonders

die folgende duale Eigenschaft von $\mathbb{P}$ -Projektionen ins Auge (BEHREND [\[21\]](#)(1979), Seite 14, Satz 1.5): Eine Projektion $P \in \mathcal{L}(X)$ ist genau dann eine $\mathbb{P}$ -Projektion, wenn die zu P duale Abbildung $P^* \in \mathcal{L}(X^*)$ eine $\mathbb{P}^q$ -Projektion ist. Insbesondere gilt: Der Annihilator eines $\mathbb{P}$ -Summanden ist ein $\mathbb{P}^q$ -Summand. Als abgeschwächte Version eines M -Summanden definiert man:

Ein abgeschlossener Unterraum $J \subseteq X$ heißt ein M -Ideal in X , wenn der Annihilator von J ein L -Summand im Dualraum X^* ist. Nach [2.4.20](#) (dort mit $p = 1$) und Satz [2.4.36](#)(b) ist ein Untervektorraum J von X genau dann ein M -Ideal in X , wenn

$$X^* \cong J^\perp \oplus_1 J^* \quad (2.32)$$

gilt; siehe dazu auch BEHREND [\[21\]](#)(1979), Seite 35, und HARMAND, WERNER und WERNER [\[126\]](#)(1992), Seite 11, Satz I.1.12.

Folglich hat man: Ist J ein M -Ideal in X und ist J reflexiv, so ist, da allgemein ein Banachraum genau dann reflexiv ist, wenn sein Dualraum reflexiv ist (WERNER [\[291\]](#)(2002), Seite 105, Satz III.3.4(b)), nach [2.5.15](#) J^* schwach*-abgeschlossen (in $J^\perp \oplus_1 J^*$).

M -Ideale lassen sich geometrisch charakterisieren (ALFSEN und EFFROS [\[7\]](#)(1972), Seite 122, Theorem 5.9 und BEHREND [\[21\]](#)(1979), Seite 47, Theorem 2.17): Ein abgeschlossener Unterraum J von X ist genau dann ein M -Ideal, wenn er die sogenannte 3-Kugel-Eigenschaft erfüllt, soll heißen, immer wenn B_1, B_2, B_3 drei offene Kugeln mit $B_1 \cap B_2 \cap B_3 \neq \emptyset$ und $B_i \cap J \neq \emptyset, i = 1, 2, 3$, sind, gilt $B_1 \cap B_2 \cap B_3 \cap J \neq \emptyset$.

Ist J ein M -Ideal in X , so sind die folgenden vier Aussagen äquivalent (BEHREND [\[21\]](#)(1979), Seite 34, Satz 2.2):

- (i) J ist ein M -Summand.
- (ii) $J^{\perp\perp}$ ist schwach*-abgeschlossen.
- (iii) Die L -Projektion auf $J^\perp$ ist schwach*-stetig.
- (iv) Für jedes $x \in X$ existiert ein $y \in J$, so dass für alle $\ell \in J^{\perp\perp}$ die Gleichung $\ell(x) = \ell(y)$ gilt.

Beweis. Interessant ist hier die Implikation (ii) $\Rightarrow$ (i) und nur diese wird hier bewiesen. Gelte also (ii). Betrachte die beiden Dualsysteme (X, X^*) und (X^*, X^{**}) . Nach dem Bipolarensatz [2.5.6](#) ist $J^{\perp\perp} = J^{\perp\perp\perp\perp}$, insbesondere ist $J^{\perp\perp\perp} \subseteq X$ schwach*-abgeschlossen, also auch abgeschlossen. Nach [2.5.12](#) ist $X = J \dot{+} J^{\perp\perp\perp}$. Man beachte die allgemeine Tatsache $i_X(U) \subseteq U^{\perp\perp}$ für alle $U \subseteq X$. (Denn: $x \in U \Rightarrow \forall y \in U^\perp : y(x) = 0 \Leftrightarrow \forall y \in U^\perp : \hat{x}(y) = 0 \Leftrightarrow \hat{x} \in U^{\perp\perp}$.) Hier also $i_X(J) \subseteq J^{\perp\perp}$ und $i_X(J^{\perp\perp\perp}) \subseteq J^{\perp\perp\perp\perp\perp} = J^{\perp\perp\perp}$. Somit $i_X(X) = i_X(J \dot{+} J^{\perp\perp\perp}) \subseteq X^{**} = J^{\perp\perp} \dot{+} J^{\perp\perp\perp}$. Da $J^{\perp\perp}$ und $J^{\perp\perp\perp}$ komplementäre M -Summanden sind, gilt auch $X^{**} \cong J^{\perp\perp} \oplus_\infty J^{\perp\perp\perp}$, also $i_X(X) \cong i_X(J) \oplus_\infty i_X(J^{\perp\perp\perp})$, das heißt, $X \cong J \oplus_\infty J^\perp$ mit $J^\perp = J^{\perp\perp\perp}$ und J ist ein M -Summand, womit die Implikation (ii) $\Rightarrow$ (i) gezeigt ist. $\square$

Ist J ein M -Ideal in X und ist J reflexiv, so wurde vorhin bemerkt, dass $J^{\perp\perp}$ schwach*-abgeschlossen ist; nach der vorstehenden Äquivalenz von (i) und (ii) ist J dann also ein M -Summand.

Sei K ein M -Summand von X^* . Nach BEHREND [\[21\]](#)(1979), Seite 114, Theorem 5.6, ist K schwach*-abgeschlossen. Daraus folgt (ebd., Theorem 5.7(i)), dass bezüglich des Dualsystems (X, X^*) $K^\perp$ ein L -Summand in X ist mit $K = K^{\perp\perp}$ und $K^{\perp\perp} = K^{\perp\perp}$. Da somit ein Unterraum von X , dessen Annihilator ein M -Summand ist, automatisch ein L -Summand ist, wird üblicherweise auf eine analoge Definition eines L -Ideals verzichtet. Um aber einige Aussagen etwas übersichtlicher formulieren zu können, wird der Begriff des L -Ideals trotzdem verwendet und kann als synonym zu dem Begriff des L -Summanden aufgefasst werden.

Ist X ein Banachraum über $\mathbb{C}$, so ist ein Unterraum J von X genau dann ein $\mathbb{P}$ -Ideal in X , wenn er ein $\mathbb{P}$ -Ideal in $X_{\mathbb{R}}$ ist (BEHRENDT [21](1979), Seite 22, Theorem 1.12, Satz 2.3).

Ein Banachraum X über $\mathbb{K}$ heißt ein $\mathbb{P}$ -*eingebetteter Raum*, wenn $i_X(X)$ ein $\mathbb{P}$ -Ideal im Bidualraum X^{**} ist. Wegen Gleichung (2.32) ist X genau dann ein M -eingebetteter Raum, wenn bezüglich des Dualsystems (X^{**}, X^{***})

$$X^{***} = i_X(X)^{\perp} \oplus_1 i_{X^*}(X^*) \quad (2.33)$$

gilt. (Es sei an Gleichung (2.31) in 2.5.16(b) erinnert.)

Ist X ein Banachraum über $\mathbb{C}$, so ist X genau dann $\mathbb{P}$ -eingebettet, wenn $X_{\mathbb{R}}$ $\mathbb{P}$ -eingebettet ist (HARMAND, WERNER und WERNER [126](1992), Seite 101, Bemerkung (c)).

Beispiele von M -eingebetteten Räumen sind $c_0(I)$, I eine beliebige Menge, und $K(H)$, H ein Hilbertraum über $\mathbb{K}$. Beispiele von L -eingebetteten Räumen sind die Funktionenräume $L_1(\mu)$, Prädualen (siehe Unterkapitel 2.10) von gewissen Tripeln, siehe 4.7.57, und die Dualräume von M -eingebetteten Räumen.

2.6 Beschränkte Mengen

2.6.1 Definition (HEUSER [133](1975), Seite 335). Sei X ein Vektorraum über $\mathbb{K}$ und Y eine nicht leere Menge. Für jedes $y \in Y$ sei ℓ_y ein Element aus $X^{\#}$. Eine Teilmenge $S \subseteq Y$ heißt X -*beschränkt*, wenn $\sup\{|\ell_y(x)| : y \in S\} < \infty$ ist für alle $x \in X$.

2.6.2. Ist in der Situation von Definition 2.6.1 $S \subseteq Y$ eine X -beschränkte Menge, dann ist $p_S : x \mapsto \sup\{|\ell_y(x)| : y \in S\}$ auf X eine Halbnorm. Ist $\mathfrak{S}$ ein System von X -beschränkten Mengen, dann erzeugt $\{p_S : S \in \mathfrak{S}\}$ eine Vektorraumtopologie auf X und diese Topologie wird *die von $\mathfrak{S}$ erzeugte Topologie* genannt.

Sei $(X, Y; B)$ ein Dualsystem. Setze $\ell_y := B(\cdot, y)$ für alle $y \in Y$. Dann erzeugt jedes System $\mathfrak{S}$ X -beschränkter Teilmengen von Y vermöge den Halbnormen $p_S : x \mapsto \sup\{|\ell_y(x)| : y \in S\}$, $S \in \mathfrak{S}$, eine Vektorraumtopologie auf X , die man $\mathfrak{S}$ -Topologie nennt; sie ist die Topologie der gleichmäßigen Konvergenz auf den Mengen von $\mathfrak{S}$.

Die schwache Topologie $\sigma(X, Y)$ wird durch das System aller einpunktigen Teilmengen von Y erzeugt. Weitere Beispiele anderer Mengensysteme $\mathfrak{S}$, die die schwache Topologie $\sigma(X, Y)$ erzeugen, findet man bei HORVÁTH [142](1966), Seite 197, Beispiel 1.

Bezüglich der $\mathfrak{S}$ -Topologie siehe auch TAYLOR und LAY [275](1980), Seite 167.

2.6.3 Definition (BOURBAKI [35](1987), III.1.2). Sei X ein topologischer Vektorraum über $\mathbb{K}$. Eine Teilmenge $A \subseteq X$ heißt *beschränkt*, wenn sie von jeder Nullumgebung in X absorbiert wird.

2.6.4. Gemäß der Anmerkung bei BOURBAKI [35](1987), III.1.2, ist eine Teilmenge A eines topologischen Vektorraumes X über $\mathbb{K}$ genau dann beschränkt, wenn für jede Nullumgebung U von X ein $\lambda > 0$ existiert, so dass $A \subseteq \lambda \cdot U$ gilt. Äquivalent dazu genügt es, die besagten Nullumgebungen dabei als kreisförmig anzunehmen (MEGGINSON [211](1998), Seite 170, Korollar 2.2.10).

2.6.5 Satz (HEUSER [133](1975), Satz 93.3, Seite 357). *Sei die Topologie τ des lokal konvexen Raumes X über $\mathbb{K}$ erzeugt durch ein Halbnormensystem P . Dann gilt: $S \subseteq X$ τ -beschränkt $\Leftrightarrow \forall p \in P : p$ auf $S \subseteq X$ beschränkt.*

Daraus folgt für jedes Dualsystem $(X, Y; B)$: $S \subseteq X$ ist $\sigma(X, Y)$ -beschränkt $\Leftrightarrow \forall y \in Y : B(\cdot, y)$ bleibt auf S beschränkt $\Leftrightarrow \forall y \in Y : \sup\{|B(x, y)| : x \in S\} < \infty \Leftrightarrow S \subseteq X$ ist Y -beschränkt.

2.6.6 Bemerkung. Ist X ein normierter Vektorraum über $\mathbb{K}$, so ist jede beschränkte Teilmenge von X $\sigma(X, X^*)$ -beschränkt.

Beweis. Sei $A \subseteq X$ und $y \in X^{*\times}$. Dann gilt: $\sup\{|y(x)| : x \in A\} = \|y\| \cdot \sup\{|\frac{y}{\|y\|}(x)| : x \in A\} \leq \|y\| \sup_{x \in A} \sup_{x^* \in B_{X^*}} |x^*(x)| = \|y\| \sup\{\|x\| : x \in A\}$. $\square$

Mit dem Satz von Banach-Steinhaus zeigt man, dass auch die Umkehrung dieser Aussage gilt, siehe TAYLOR und LAY [275](1980), Seite 170, Theorem 9.2, oder HORVÁTH [142](1966), Seite 202, Beispiel 5. Die somit vorliegende Äquivalenz gilt sogar, wenn X allgemeiner nur als ein lokal konvexer Hausdorffraum vorausgesetzt wird, siehe CHOQUET [46](1969), Seite 74, Theorem 23.11. Siehe auch hier den Satz von Mackey, 2.6.12.

2.6.7 Definition (HEUSER [133](1975), Seite 365). Sei (X, τ) ein lokal konvexer Vektorraum über $\mathbb{K}$. Sei (X, Y) ein in Y trennendes Dualsystem. Dann heißt die Topologie τ *zulässig* (im Englischen *compatible*) für das Dualsystem (X, Y) , wenn $Y = (X, \tau)^*$ gilt.

2.6.8 (BOURBAKI [35](1987), IV.1.1). Eine Topologie τ eines lokal konvexen Vektorraumes X über $\mathbb{K}$ ist also genau dann zulässig für ein Dualsystem (X, Y) , wenn die Astriktion $(X, \tau)^* \upharpoonright R$ der Abbildung $R: Y \rightarrow X^\#$ von Definition 2.2.22 auf $(X, \tau)^*$ bijektiv ist.

2.6.9 Bemerkung. Per Definition ist also die Topologie eines lokal konvexen Raumes X zulässig für das Dualsystem (X, X^*) . Die Topologie $\sigma(X, Y)$ ist offenbar die grösste der für das Dualsystem (X, Y) zulässigen Topologien. Es existiert eine feinste für (X, Y) zulässige Topologie; sie wird die *Mackey-Topologie* genannt und mit $\tau(X, Y)$ bezeichnet. Nach dem Satz von Mackey-Arens ist eine lokal konvexe Topologie genau dann auf X zulässig für das in Y trennende Dualsystem (X, Y) , wenn sie feiner als $\sigma(X, Y)$ und gröber als $\tau(X, Y)$ ist. Die feinste $\mathfrak{G}$ -Topologie auf X wird die *starke Topologie* genannt, mit $\beta(X, Y)$ bezeichnet und ist feiner als $\tau(X, Y)$ und im Allgemeinen nicht zulässig (also einen größeren Dualraum ergibt). Ist X ein normierter Raum, dann ist zwar die Norm-Topologie von X^* gleich $\beta(X^*, X)$, aber im Allgemeinen ist die Norm-Topologie von X nur gleich $\tau(X, X^*)$; siehe KÖTHE [187] (1966), Seite 265, für ein Beispiel eines normierten Raumes, wo die Norm-Topologie verschieden von der starken Topologie ist. Ist aber X ein Banachraum, so ist die Norm-Topologie von X auch gleich $\beta(X, X^*)$.

Als ein Korollar des Satzes 2.4.33 hat man:

2.6.10 Satz (HEUSER [133](1975), Seite 365). Sei (X, Y) ein in Y trennendes Dualsystem. Sei K eine konvexe Teilmenge von X . Dann ist K in jeder oder in keiner der zulässigen Topologien abgeschlossen.

2.6.11 Korollar. Sei X ein Vektorraum. Seien τ_1 und τ_2 zwei lokal konvexe Topologien auf X , so dass $(X, \tau_1)^* = (X, \tau_2)^*$ gilt. Sei K eine konvexe Teilmenge von X . Dann gilt $cl(\tau_1; K) = cl(\tau_2; K)$. (Siehe auch Bemerkung 1.2.5.)

Beweis. Da $(X, \tau_1)^* = (X, \tau_2)^*$ gilt, kann man schlichtweg von $X^* := (X, \tau_1)^*$ reden. Nach der Bemerkung 2.2.34 ist (X, X^*) ein in X^* trennendes Dualsystem. Nach Bemerkung 2.6.9 sind sowohl τ_1 als auch τ_2 zulässig für das Dualsystem (X, X^*) . $cl(\tau_1; K)$ ist eine in (X, τ_1) abgeschlossene konvexe Menge. Nach dem Satz 2.6.10 ist $cl(\tau_1; K)$ dann auch in (X, τ_1) eine abgeschlossene konvexe Menge. Somit ist $cl(\tau_1; K) \supseteq cl(\tau_2; K)$. Analog zeigt man „ $\subseteq$ “. Einen direkten Beweis dieses Korollars mit dem Satz von Hahn-Banach ohne Verwendung des Satzes 2.6.10 findet sich bei WERNER [291](2002), Seite 382, Satz VIII.3.5. $\square$

2.6.12 Satz von Mackey. Sei (X, Y) ein in Y trennendes Dualsystem und A eine Teilmenge von X . Dann ist A in jeder oder in keiner der für das Dualsystem (X, Y) zulässigen Topologien beschränkt.

Beweis. BOURBAKI [35](1987), IV.1.1, Satz 1(ii); KÖTHE [187] (1966), Seite 255, Satz 7; HORVÁTH [142](1966), Seite 209, Theorem 3. $\square$

2.6.13 Scheiben (BOURGIN [36](1983); GILES [109](1982)). Sei $(X, Y; B)$ ein Dualsystem. Es sei an Satz 2.6.5 und dessen anschließende Bemerkung 2.6.6 erinnert.

(a) Ist A eine $\sigma(X, Y)$ -beschränkte Teilmenge von X und $y \in Y$, so bezeichnet $M(A, y)$ die erweiterte reelle Zahl $\sup \{\operatorname{Re} B(x, y) : x \in A\}$.

(b) Ist A eine $\sigma(Y, X)$ -beschränkte Teilmenge von Y und $x \in X$, so bezeichnet $M(A, x)$ die erweiterte reelle Zahl $\sup \{\operatorname{Re} B(x, y) : y \in A\}$.

(c) Ist A eine $\sigma(X, Y)$ -beschränkte Teilmenge von X , $y \in Y^\times$ und $\varepsilon > 0$, dann heißt die Menge

$$S(A, y, \varepsilon) := \{x \in A : \operatorname{Re} B(x, y) > M(A, y) - \varepsilon\}$$

die durch y und ε bestimmte Scheibe (im Englischen *slice*) von A (oder einfach auch nur eine Scheibe von A). Falls klar ist, von welcher Menge A eine Scheibe betrachtet werden soll, wird anstatt $S(A, y, \varepsilon)$ auch einfach nur $S(y, \varepsilon)$ geschrieben. Man bemerke, dass $S(A, y, \varepsilon)$ eine relativ $\sigma(X, Y)$ -offene Teilmenge von A ist. Speziell für den Fall, dass X ein topologischer Vektorraum über $\mathbb{K}$ ist, werden die dann in X^* bezüglich des Dualsystems (X^*, X) betrachtbaren Scheiben schwach*-Scheiben genannt. Etwas vorgreifend sei diesbezüglich auf Lemma 2.10.3 hingewiesen.

2.6.14 Beulen (GILES [109](1982), Seite 220; DAY [61](1973), Seite 105). Sei X ein normierter Vektorraum über $\mathbb{K}$. Sei A eine beschränkte, nicht leere Teilmenge von X . Sei $\varepsilon > 0$. Ein $x \in A$ heißt ε -Beulpunkt (im Englischen ε -denting point) von A , falls $x + \varepsilon \cdot \operatorname{int}(B_X)$ eine Scheibe von A enthält, die x enthält. Ein $x \in A$ heißt Beulpunkt von A , falls x ein ε -Beulpunkt für jedes $\varepsilon > 0$ ist. Unmittelbar bemerkt man, dass die auf A eingeschränkte identische Abbildung von X an jedem Beulpunkt von A (w - n -)stetig ist. Bezüglich der Umkehrung dieser Aussage siehe 2.6.15. Ist X ein Dualraum, so definiert man entsprechend ε -schwach*-Beulpunkt und schwach*-Beulpunkt mittels schwach*-Scheiben.

Ein $x \in A$ ist genau dann ein ε -Beulpunkt von A , wenn gilt:

$$x \notin \overline{\operatorname{co}}(A \setminus (x + \varepsilon \cdot \operatorname{int}(B_X))). \quad (2.34)$$

Beweis. „ $\Leftarrow$ “: Sei $x \in A$. Setze $V := \overline{\operatorname{co}}(A \setminus (x + \varepsilon \cdot \operatorname{int}(B_X)))$. Liege der Fall $x \notin V$ vor. Nach dem Satz von Hahn-Banach existiert ein $r > 0$ und ein $x^* \in X^{\times\ast}$, so dass für alle $v \in V$ gilt: $\operatorname{Re} x^*(x) > r \geq \operatorname{Re} x^*(v)$, also $\operatorname{Re} x^*(x) > r \geq M(V, x^*)$. Setze $\alpha := M(A, x^*) - r$. Dass ein $y \in A$ in der Scheibe $S(A, x^*, \alpha)$ enthalten ist, heißt somit: $\operatorname{Re} x^*(y) > M(A, x^*) - \alpha = r \geq M(V, x^*) = \sup_{v \in V} \operatorname{Re} x^*(v)$ und y ist kein Element von V . Dann aber muss y ein Element von $(x + \varepsilon \cdot \operatorname{int}(B_X)) \cap A$ sein, da sonst Widerspruch. Das heißt, die Scheibe $S(A, x^*, \alpha)$ ist eine Teilmenge von $(x + \varepsilon \cdot \operatorname{int}(B_X)) \cap A$ und x ist ein ε -Beulpunkt von A .

„ $\Rightarrow$ “: Sei x ein ε -Beulpunkt von A . Sei dementsprechend $S(A, x^*, \alpha)$ eine Scheibe von A mit $x \in S(A, x^*, \alpha) \subseteq x + \varepsilon \cdot \operatorname{int}(B_X)$. Setze wieder $V := \overline{\operatorname{co}}(A \setminus (x + \varepsilon \cdot \operatorname{int}(B_X)))$. Also $V \subseteq \overline{\operatorname{co}}(A \setminus S(A, x^*, \alpha))$. Ein $\xi \in A$ ist genau dann ein Element von $A \setminus S(A, x^*, \alpha)$, wenn gilt: $\operatorname{Re} x^*(\xi) \leq M(A, x^*) - \alpha$. Also ist die Zahlenmenge $\operatorname{Re} x^*(\overline{\operatorname{co}}(A \setminus S(A, x^*, \alpha)))$ eine Teilmenge von $(-\infty, M(A, x^*) - \alpha]$. Da x in der Scheibe $S(A, x^*, \alpha)$ enthalten ist,

ist $\operatorname{Re} x^*(x) > M(A, x^*) - \alpha$ und somit ist x kein Element von $\overline{\operatorname{co}}(A \setminus S(A, x^*, \alpha))$, also insbesondere kein Element von V . $\square$

Jeder Beulpunkt von A ist ein Extrempunkt von A .

Beweis. Sei x ein Beulpunkt von A . Seien $a, b \in A$ und sei $\lambda \in (0, 1)$ mit $\lambda a + (1 - \lambda)b = x$. Angenommen, a und b wären von x verschieden. Dann gäbe es ein $\varepsilon > 0$, so dass $x + \varepsilon \cdot \operatorname{int}(B_X)$ weder a noch b enthalten würde. a und b wären also in der Menge $\overline{\operatorname{co}}(A \setminus (x + \varepsilon \cdot \operatorname{int}(B_X)))$ enthalten; da diese Menge aber konvex ist, enthielte sie auch x , Widerspruch. $\square$

Sei weiterhin wie gehabt A eine beschränkte, nicht leere Teilmenge von X . A heißt *beulbar* (im Englischen *dentable*), wenn A einen ε -Beulpunkt von A für jedes $\varepsilon > 0$ enthält.

A ist genau dann beulbar, wenn A für jedes $\varepsilon > 0$ eine nicht leere Scheibe enthält, deren Durchmesser kleiner als ε ist.

Beweis. Sei A beulbar. Für jedes $\varepsilon > 0$ existiert also ein $x_\varepsilon \in A$, so dass $A \cap (x_\varepsilon + \varepsilon \cdot \operatorname{int}(B_X))$ eine Scheibe von A enthält, die x_ε enthält. Somit enthält A für jedes $\varepsilon > 0$ eine Scheibe mit Durchmesser kleiner oder gleich 2ε .

Enthalte A nun für beliebig kleines $\varepsilon > 0$ stets eine nicht leere Scheibe, deren Durchmesser kleiner als ε ist. Sei $\varepsilon > 0$ beliebig, $S = S(A, x^*, \alpha)$ eine Scheibe von A mit einem Durchmesser kleiner als $\varepsilon/3$ und $x \in S$. Dann gilt $S \subseteq (x + \frac{2}{3}\varepsilon \cdot B_X) \cap A \subseteq (x + \varepsilon \cdot \operatorname{int}(B_X)) \cap A$, also $\overline{\operatorname{co}}(A \setminus (x + \varepsilon \cdot \operatorname{int}(B_X))) \subseteq \overline{\operatorname{co}}(A \setminus S)$. Es gilt $\operatorname{Re} x^* \overline{\operatorname{co}}(A \setminus S) \subseteq (-\infty, M(A, x^*) - \alpha]$. Da $x \in S$, gilt $\operatorname{Re} x^*(x) > M(A, x^*) - \alpha$. Somit $x \notin \overline{\operatorname{co}}(A \setminus S)$, insbesondere $x \notin \overline{\operatorname{co}}(A \setminus (x + \varepsilon \cdot \operatorname{int}(B_X)))$, das heißt, x ist ein ε -Beulpunkt von A . $\square$

Sei A weiterhin wie gehabt. Dann gilt:

Ein $x \in A$ ist genau dann ein Beulpunkt von A , wenn für jedes $\varepsilon > 0$ eine Scheibe von A existiert, deren Durchmesser kleiner oder gleich ε ist, und die x enthält.

Beweis. Sei x ein Beulpunkt von A . Sei $\varepsilon > 0$. Dann ist x insbesondere ein $\frac{\varepsilon}{2}$ -Beulpunkt, das heißt, x ist in einer Scheibe von A enthalten, die einen Durchmesser kleiner oder gleich ε aufweist.

Sei umgekehrt für jedes $\varepsilon > 0$ eine Scheibe von A vorhanden, deren Durchmesser kleiner oder gleich ε ist, und die x enthält. Sei $\varepsilon > 0$ beliebig. Sei S eine Scheibe mit $x \in S$ und Durchmesser kleiner oder gleich $\frac{\varepsilon}{3}$. Das heißt, $S \subseteq x + \frac{2}{3}\varepsilon \cdot B_X \subseteq x + \varepsilon \cdot \operatorname{int}(B_X)$ und x ist ein ε -Beulpunkt von A . $\square$

A ist genau dann beulbar, wenn $\operatorname{cl} \operatorname{co} A$ beulbar ist (BOURGIN [36](1983), Seite 29).

Sei K eine konvexe Teilmenge von X und $x \in K$. x heißt ein *stark extremer Punkt* (im Englischen *strongly extremal point*) von K , wenn für alle $\varepsilon > 0$ ein $\delta > 0$ existiert, so dass x nicht die Mitte irgendeines Segmentes (siehe 2.2.43) der Länge ε ist, welches in $K + \delta \cdot B_X$ liegt. x ist genau dann ein stark extremer Punkt eines nicht leeren, beschränkten, abgeschlossenen, konvexen K , wenn x ein Beulpunkt von K ist (DAY, ebd.). Man beachte, dass BOURGIN [36](1983) zwischen *strongly* (ebd., Seite 67) und *strong* (ebd., Seite 340) exposed points unterscheidet, wobei die letzteren ebenfalls genau die Beulpunkte eines wie eben verwendeten K sind, siehe auch hier 2.11.11.

In der englischen Literatur finden sich auch sogenannte *strong extreme points*, die von den hier definierten stark extremen Punkten unterschieden werden müssen, siehe LIN, LIN und TROYANSKI [199](1988).

2.6.15 Stetigkeitspunkte (LIN, LIN und TROYANSKI [199](1988)). Sei X ein Banachraum über $\mathbb{K}$, A eine nicht leere, beschränkte, abgeschlossene, konvexe Teilmenge von X und $x \in A$. Dann ist x genau dann ein Beulpunkt von A , wenn sowohl die Abbildung $A \upharpoonright \text{Id}_X \upharpoonright A$ ($w - n$)-stetig bei x ist als auch x ein Extrempunkt von A ist.

2.6.16 Atome in Maßräumen (FREMLIN [99](2001), 211I-K). Sei (X, Σ, μ) ein Maßraum. Eine Menge $A \in \Sigma$ heißt ein *Atom für μ* , falls $\mu(A) > 0$ und für alle $E \in \Sigma$ mit $E \subseteq A$ gilt $\mu(E) = 0$ oder $\mu(A \setminus E) = 0$. μ oder (X, Σ, μ) heißt *atomlos*, falls es für μ keine Atome gibt. μ oder (X, Σ, μ) heißt *rein atomar* (im Englischen *purely atomic*), falls jedes $E \in \Sigma$ mit $\mu(E) \neq 0$ ein Atom für μ umfasst.

2.6.17 Radon-Nikodým-Eigenschaft (BOURGIN [36](1983), Seite 31 und 60). Sei X ein Banachraum über $\mathbb{K}$. Sei K eine abgeschlossene, beschränkte, konvexe Teilmenge von X . Dann weist K genau dann die sogenannte *Radon-Nikodým-Eigenschaft* (im Englischen *Radon-Nikodým property*, kurz *RNP*) auf, wenn jede nicht leere Teilmenge von K beulbar ist.

Sei nun K eine abgeschlossene, konvexe Teilmenge von X (die also nicht notwendig beschränkt zu sein braucht). Dann sagt man, K habe die *Radon-Nikodým-Eigenschaft*, wenn jede abgeschlossene, beschränkte, konvexe Teilmenge von K die Radon-Nikodým-Eigenschaft aufweist. Insbesondere weist also X genau dann die Radon-Nikodým-Eigenschaft auf, wenn B_X die Radon-Nikodým-Eigenschaft aufweist.

Beispiele von Räumen, die die Radon-Nikodým-Eigenschaft aufweisen, sind: Reflexive Banachräume, Banachräume mit Fréchet-glattem Prädual (siehe Unterkapitel 2.10 und 4.1.11), $\ell_1(I)$ für jede Menge I , $C_p(H)$ für $1 \leq p < \infty$ (Schattenklassen, siehe 3.6.1). Jede konvexe, schwach kompakte Teilmenge von X besitzt die Radon-Nikodým-Eigenschaft.

Beispiele von Räumen, die die Radon-Nikodým-Eigenschaft *nicht* aufweisen, sind: $L_1[0, 1]$, $L_1(X, \Sigma, \mu)$ für nicht rein atomares μ (also zum Beispiel, wenn $\mu \neq 0$ und X atomlos ist), c_0 , c , ℓ_∞ , $L_\infty[0, 1]$ und $C(\Omega)$ für jeden unendlichen, kompakten Hausdorffraum Ω .

Weitere Beispiele und äquivalente Formulierungen, wann genau ein Banachraum die Radon-Nikodým-Eigenschaft aufweist, finden sich bei DIESTEL und UHL [66](1977), Seiten 217-219.

Ist Y ein Banachraum über $\mathbb{K}$, dann heißt ein $T \in \mathcal{L}(X, Y)$ *stark Radon-Nikodým* (im Englischen *strong Radon-Nikodým*), falls $\text{cl}(T(B_X))$ die Radon-Nikodým-Eigenschaft hat.

2.6.18 Definition. Seien X und Y zwei topologische Vektorräume über $\mathbb{K}$. Eine Familie $\mathfrak{F}$ von Abbildungen $f: X \rightarrow Y$ heißt *gleichgradig stetig im Punkt $x_0 \in X$* , wenn es zu jeder Nullumgebung (oder auch nur zu jeder Basisnullumgebung) V in Y eine Nullumgebung U in X gibt, so dass gilt:

$$f(x) - f(x_0) \in V \quad \text{für alle } x \in X \text{ mit } x - x_0 \in U \quad \text{und für alle } f \in \mathfrak{F}.$$

Die Familie wird *gleichstetig* genannt, wenn sie in jeden Punkt von X gleichgradig stetig ist.

2.6.19 Bemerkungen (HEUSER [133](1975), Seite 369).

(a) Seien X und Y zwei topologische Vektorräume über $\mathbb{K}$. Ist $S \subseteq L(X, Y)$, so ist S offensichtlich genau dann gleichstetig, wenn S in einem beliebigen Punkt $x_0 \in X$ gleichgradig stetig ist. Da dies also insbesondere für $x_0 = 0$ gilt, ist $S \subseteq L(X, Y)$ genau dann gleichstetig, wenn für jede Nullumgebung V in Y die Menge $\bigcap \{\ell^{-1}(V) : \ell \in S\}$ eine Nullumgebung in X ist.

(b) Sei X ein lokal konvexer Vektorraum über $\mathbb{K}$ und $S \subseteq X^*$. Dann ist S genau dann gleichstetig, wenn es eine stetige Halbnorm p auf X gibt, so dass für alle $x \in X$ und

$\ell \in S$ die Ungleichung $|\ell(x)| \leq p(x)$ gilt; in diesem Fall ist S auch $\sigma(X^*, X)$ -beschränkt (Zum Beweis siehe ebd.).

(c) Sei X ein topologischer Vektorraum über $\mathbb{K}$ und $S \subseteq X^*$. Dann ist S genau dann gleichstetig, wenn es eine Nullumgebung U in X mit $S \subseteq U^\bullet$ gibt (ebd.), oder dazu äquivalent, wenn $S^\bullet$ eine Nullumgebung in X ist.

Beweis. Ist S gleichstetig, so ist $U := S^\bullet = \bigcap \{\ell^{-1}(B_{\mathbb{K}}) : \ell \in S\}$ eine Nullumgebung in X mit $S \subseteq U^\bullet$. Sei nun umgekehrt U eine Nullumgebung in X mit $S \subseteq U^\bullet$. Sei V eine Nullumgebung in $\mathbb{K}$, ohne Einschränkung etwa $V = \varepsilon \cdot B_{\mathbb{K}}$ für ein $\varepsilon > 0$. Dann gilt $\bigcap \{\ell^{-1}(V) : \ell \in S\} = \varepsilon \cdot S^\bullet \supseteq \varepsilon \cdot U^{\bullet\bullet} \supseteq \varepsilon \cdot U$. $\square$

2.6.20 Satz (HEUSER [133](1975), Seite 369). *Sei X ein lokal konvexer Raum über $\mathbb{K}$. Dann wird die Topologie von dem System $\mathfrak{S}$ aller gleichstetigen Teilmengen von X^* erzeugt.*

Aus dem Satz von Alaoglu-Bourbaki folgt mit Bemerkung 2.6.19(c) sofort der

2.6.21 Satz. *Sei X ein topologischer Vektorraum über $\mathbb{K}$ und $S \subseteq X^*$ gleichstetig. Dann ist S relativ $\sigma(X^*, X)$ -kompakt.*

Für Banachräume über $\mathbb{K}$ gilt auch die Umkehrung. Genauer hat man:

2.6.22 (TAYLOR und LAY [275](1980), Seite 174). *Sei X ein Banachraum über $\mathbb{K}$ und $S \subseteq X^*$. Dann sind die folgenden Aussagen äquivalent:*

- (a) S ist $\sigma(X^*, X)$ -beschränkt.
- (b) S ist $\sigma(X^*, X^{**})$ -beschränkt.
- (c) S ist Norm-beschränkt.
- (d) S ist gleichstetig.
- (e) S ist relativ $\sigma(X^*, X)$ -kompakt.

2.7 Topologische Algebren

2.7.1 Definition. Sei A eine Algebra über $\mathbb{K}$, versehen mit einer Topologie, so dass A ein topologischer Vektorraum über $\mathbb{K}$ ist. Die Multiplikation der Algebra A heißt *getrennt stetig*, wenn für alle $a \in A$ die durch a bestimmten Links- und Rechts-Multiplikationen (siehe Definition 2.1.18) alle stetig sind. Ist die Multiplikation der Algebra A getrennt stetig, so heißt A eine *topologische Algebra* über $\mathbb{K}$.

2.7.2 Bemerkung (PALMER [231](1994), Seite 8). Einige Autoren verlangen in der Definition einer topologischen Algebra noch, dass sie Hausdorff'sch ist. PALMER [231](1994) und DALES [56](2000) zum Beispiel nennen eine so wie hier in Definition 2.7.1 definierte topologische Algebra eine *semitopologische Algebra*. NAIMARK [221](1972), Seite 167, definiert eine topologische Algebra A analog wie hier, nur dass er A nicht wie hier als einen topologischen Vektorraum, sondern als einen lokal konvexen Raum voraussetzt.

Man beachte, dass in einem topologischen Vektorraum die Trennungsaxiome T_0 , T_1 , T_2 , $(T_1 + T_3)$ und $(T_1 + T_{3a})$ äquivalent sind (MEGGINSON [211](1998), Seite 174).

Nach Bemerkung 2.3.2(b) ist jede normierte Algebra über $\mathbb{K}$ eine topologische Algebra über $\mathbb{K}$.

$GL(X)$ (bzw. $G\mathcal{L}(X)$) bezeichnet die Gruppe, die aus den invertierbaren Elementen von $L(X)$ (bzw. $\mathcal{L}(X)$) besteht.

2.7.3 (UPMEIER [279](1985), Seite 33). Sei A eine nichtassoziative Banachalgebra über $\mathbb{K}$. Die Menge

$$\text{Aut}(A) := \{T \in G\mathcal{L}(A) : \forall x, y \in A : T(xy) = (Tx)(Ty)\}$$

aller stetigen Automorphismen von A ist eine abgeschlossene Untergruppe von $G\mathcal{L}(A)$. Für einen Banachraum X über $\mathbb{K}$ wird die Banach-Lie-Algebra $(\mathcal{L}(X))^\ominus$ mit $\mathfrak{gl}(X)$ bezeichnet. Die Menge

$$\text{aut}(A) := \{\delta \in \mathfrak{gl}(A) : \delta \text{ Derivation von } A\}$$

aller stetigen Derivationen von A ist eine abgeschlossene Unteralgebra von $\mathfrak{gl}(A)$, also eine Banach-Lie-Algebra. Das folgende Diagramm kommutiert:

$$\begin{array}{ccc} \mathfrak{gl}(A) & \xrightarrow{\exp} & G\mathcal{L}(A) \\ \cup & & \cup \\ \text{aut}(A) & \xrightarrow{\exp} & \text{Aut}(A) \end{array}$$

2.7.4 Satz (PALMER [231](1994), Seite 10, Satz 1.1.6). Sei A eine topologische Algebra über $\mathbb{K}$, die Hausdorff'sch ist. Dann ist für alle Teilmengen S von A die Menge S' abgeschlossen. Somit ist das Zentrum einer Algebra und jede maximal kommutative Unteralgebra abgeschlossen.

Beweis. Nach Voraussetzung ist insbesondere für alle $s \in A$ die Abbildung $L_s - R_s$ (siehe Definition 2.1.18) stetig. Die Abgeschlossenheit von S' folgt dann mit Satz 2.3.26(k). Das Zentrum ist A' und Satz 2.3.26(f) zeigt $C = C'$. $\square$

2.7.5 Bemerkung. Sei A eine topologische Algebra über $\mathbb{K}$, die Hausdorff'sch ist und S eine Teilmenge von A . Dann gilt $S' = (\text{cl } S)'$.

Beweis. Sei $x \in S'$ und $y \in \text{cl}(S)$. Sei $(y_\nu) \subseteq S$ ein Netz, das gegen y konvergiert. Dann gilt $xy = x \lim y_\nu = \lim (xy_\nu) = \lim (y_\nu x) = (\lim y_\nu) x = yx$, also $S' \subseteq (\text{cl } S)'$. $\square$

2.7.6 Satz. Sei A eine topologische Algebra über $\mathbb{K}$, die Hausdorff'sch ist und S eine Teilmenge von A . Dann sind die beiden folgenden Aussagen äquivalent:

$$(a) \ S \text{ abgeschlossen} \Leftrightarrow S = S''.$$

$$(b) \ \text{cl}(S) = S''.$$

Beweis. (a) $\Rightarrow$ (b): Gelte (a), also $\text{cl}(S) = \text{cl}(S)''$. Nach Bemerkung 2.7.5 gilt $\text{cl}(S)' = S'$, also $\text{cl}(S)'' = S''$. Somit $\text{cl}(S) = S''$, das heißt, (b).

$$(b) \Rightarrow (a): S \text{ abgeschlossen} \Leftrightarrow S = \text{cl}(S) = S''. \quad \square$$

2.7.7 Bemerkung (Fortsetzung von 2.1.33). Sei A eine Hausdorff'sche topologische Algebra über $\mathbb{K}$ und p ein idempotentes Element von A . Dann sind die Mengen Ap , pA und pAp in A abgeschlossen.

Beweis. Die Abbildung $T: a \mapsto a - ap$ ist stetig. Also ist $Ap = T^{-1}(\{0\})$ abgeschlossen. $\square$

2.8 Darstellungstheorie

2.8.1 Definition (PALMER [231](1994), Seiten 225, 440 und 456). Sei A eine Algebra über $\mathbb{K}$ und X ein Vektorraum über $\mathbb{K}$. Eine *Darstellung* (bzw. *Antidarstellung*) von A in $L(X)$ oder auf X ist ein Homomorphismus (bzw. Antihomomorphismus) $T: A \rightarrow L(X)$. Für $a \in A$ schreibt man T_a für den mit a assoziierten *Darsteller* $T(a)$. Genauer bezeichnet man dann eine Darstellung als (T, X) . Wenn von einer Darstellung T einer Algebra A gesprochen wird, ohne mit anzugeben, auf welchem Vektorraum die Darsteller T_a , $a \in A$, agieren, so wird dieser Vektorraum mit X_T bezeichnet. Ein Unterraum $U \subseteq X$ ist *T-invariant* (oder einfach *invariant*, wenn klar ist, welche Darstellung gemeint ist), wenn jeder Darsteller T_a , $a \in A$, ihn in sich abbildet.

Die bilineare Abbildung $B: A \times X \rightarrow X$, $(a, x) \mapsto T_a(x)$ wird *die zu der Darstellung von A auf X assoziierte bilineare Abbildung* oder die bilineare Abbildung der Darstellung von A auf X genannt und auch als $\langle \cdot, \cdot \rangle_T$ geschrieben. Somit kann man von dem *Bilinearsystem* $(A, X; \langle \cdot, \cdot \rangle_T)$ einer Darstellung T von A auf X sprechen. Wenn klar ist, welche Darstellung gemeint ist, wird wie üblich das Bilinearsystem als (A, X) und die bilineare Abbildung als $\langle \cdot, \cdot \rangle$ geschrieben.

Die Darstellung (bzw. Antidarstellung) T von A auf X heißt

- (a) *treu* (im Englischen *faithful*), wenn das Bilinearsystem (A, X) in A trennend ist, das heißt, wenn T injektiv ist.
- (b) *trivial*, wenn $T_a = 0$ für alle $a \in A$ gilt.
- (c) *irreduzibel* (genauer *algebraisch irreduzibel* oder auch *streng irreduzibel*), wenn $\{0\}$ und X die einzigen T -invarianten Unterräume von X sind und T nicht trivial ist.
- (d) *zyklisch*, wenn es ein $x \in X$ gibt mit $X = \{T_a x : a \in A\}$; solch ein x heißt dann ein *zyklischer Vektor*. Eine zyklische Darstellung schreibt man dann auch als (T, X, x) .
- (e) *äquivalent* zu einer Darstellung (bzw. Antidarstellung) S von A auf einem Vektorraum Y über $\mathbb{K}$, wenn es eine lineare Bijektion $U: X \rightarrow Y$ gibt mit $S_a U = U T_a$ für alle $a \in A$; solch ein U heißt dann eine *Äquivalenz*.
- (f) *nicht in X ausgeartet* oder *in X trennend* oder *wesentlich* (im Englischen *essential*), wenn das Bilinearsystem (A, X) in X trennend ist. Die hierfür von manchen Autoren verwendete Bezeichnung *proper* wird hier nicht verwendet, damit keine Missverständnisse mit den Begriffen aus Definition 3.5.49 auftreten können. (Siehe auch Bemerkung 2.2.29.)
- (g) *nicht ausgeartet*, wenn T sowohl *treu* als auch *wesentlich* ist, also wenn die zu T assoziierte bilineare Abbildung nicht ausgeartet ist.

Trägt X eine Norm, dann heißt T

- (h) *normiert*, wenn $T_a \in \mathcal{L}(X)$ für alle $a \in A$ gilt.
- (i) *topologisch zyklisch*, wenn es einen zyklischen Vektor $x \in X$ gibt, so dass $T_A x = \{T_a x : a \in A\}$ dicht in X ist; solch ein x heißt dann ein *topologisch zyklischer Vektor*.
- (j) *topologisch irreduzibel*, wenn $\{0\}$ und X die einzigen abgeschlossenen T -invarianten Unterräume von X sind und T nicht trivial ist.
- (k) *topologisch äquivalent* zu einer Darstellung S von A auf einem normierten Vektorraum Y über $\mathbb{K}$, wenn es eine homöomorphe lineare Bijektion $U: X \rightarrow Y$ gibt mit $S_a U = U T_a$ für alle $a \in A$; solch ein U heißt dann eine *topologische Äquivalenz*.

2.8.2 Bemerkung. Manche Autoren nennen eine Darstellung $T: A \rightarrow L(X)$ *nicht ausgeartet*, wenn sie wesentlich ist, ohne dabei zu verlangen, dass sie treu zu sein hat.

2.8.3 Definition. Sei A eine Algebra über $\mathbb{K}$ und T eine Darstellung oder Antidarstellung von A auf einem Vektorraum X über $\mathbb{K}$. A heißt *wesentlich*, wenn T wesentlich ist.

2.8.4 Definition. Sei A eine Algebra über $\mathbb{K}$. Mit den Bezeichnungen von Definition 2.1.18 heißt die Abbildung $S: A \rightarrow L(A)$, $a \mapsto S_a := L_a$ die *links-reguläre Darstellung* von A . Entsprechend heißt die Abbildung $R: A \rightarrow L(A)$, $a \mapsto R_a$ die *rechts-reguläre Darstellung* von A (die genau genommen natürlich eine Antidarstellung ist).

2.8.5. Die Abbildung S (bzw. R) in Definition 2.8.4 ist eine Darstellung (bzw. Antidarstellung) von A auf seinen zugrunde liegenden Vektorraum über $\mathbb{K}$. Ein Unterraum U der Algebra A ist genau dann S -invariant, wenn er ein Linksideal ist.

2.8.6 Definition. Sei T eine Darstellung von A auf X , Y ein T -invarianter Unterraum von X und B eine Unteralgebra von A . Sprachlich sind zwei verschiedene Restriktionen zu unterscheiden.

Die Abbildung $a \mapsto T_a \upharpoonright Y$ ist eine Darstellung von A auf Y ; sie wird die *Restriktion von der Darstellung T zu Y* genannt, mit T^Y bezeichnet und auch eine *Subdarstellung* von T genannt.

Die Restriktion $T \upharpoonright B$ der Abbildung T auf B wird die *Restriktion des Homomorphismus T* genannt.

Wenn keine Mißverständnisse möglich sind, wird die Subdarstellung T^Y auch mit $T \upharpoonright Y$ bezeichnet.

2.8.7 (RICKART [246](1960), Seite 55). Sei A eine Algebra über $\mathbb{K}$. Das Jacobson-Radikal $\mathfrak{J}(A)$ von A ist gleich dem Durchschnitt der Kerne von allen algebraisch irreduziblen Darstellungen auf Vektorräumen über $\mathbb{K}$ von A .

2.9 Skalarbereichsänderung von Moduln und Algebren

Es sei an die Vereinbarung erinnert, dass das Wort *Modul* für *Linksmodul* steht.

2.9.1 (BOURBAKI [34](1974), II.1.13; HUNGERFORD [143](1980)). Seien R und S zwei Ringe. Sei $\Psi: R \rightarrow S$ ein Ringhomomorphismus. Sei $(X, +, \cdot_S)$ ein S -Modul. Dann existiert auf X eine kanonische Struktur eines R -Moduls:

Als Addition wähle man die Addition auf $(X, +, \cdot_S)$. Die Multiplikation $\cdot_R$ mit einem Skalar $\lambda \in R$ definiere man folgendermaßen:

$$\lambda \cdot_R x := \Psi(\lambda) \cdot_S x \quad \text{für alle } \lambda \in R, x \in X.$$

Der für $\Psi: R \rightarrow S$ und S -Modul X soeben definierte R -Modul heißt *der mit Ψ und der S -Modul-Struktur assoziierte R -Modul* und wird mit X_Ψ bezeichnet; man sagt auch, dass die R -Modul-Struktur auf X per einem *Pullback* entlang Ψ gegeben sei. Wenn keine Missverständnisse auftreten können, wird für X_Ψ auch einfach nur X geschrieben. Nebenbei sei angemerkt, dass BOURBAKI, ebd., für X_Ψ die Bezeichnungen $\Psi_*(X)$ und $X_{[R]}$ verwendet.

2.9.2 Bemerkung (KLINGENBERG und KLEIN [186](1972), Seite 78; BOURBAKI [34](1974)). Sei $\Psi: R \rightarrow S$ ein unitaler Ringhomomorphismus zwischen zwei Körpern R und S . Dann ist Ψ injektiv. Daher ist $\Psi(R)$ ein Unterkörper von S , der isomorph zu R ist, und man identifiziert vermöge Ψ die Körper R und $\Psi(R)$ und fasst R als Unterkörper von

S auf. Es wird also die Operation $\cdot_S$ eingeschränkt auf die Multiplikation mit Skalaren aus dem Unterkörper R von S . In diesem Fall heißt X_Ψ *der durch Ψ -Einschränkung des Skalarbereichs entstandene R -Modul* und diese Bezeichnung wird auch verwendet, wenn R und S keine Körper sind.

2.9.3 Bemerkung. Seien R und S zwei Ringe. Sei $\Psi: R \rightarrow S$ ein Ringhomomorphismus. Sei X ein unitaler R -Modul und sei Y ein unitaler S -Modul. Dann sind äquivalent:

- (a) $T: X \rightarrow Y$ Ψ -semilinear.
- (b) $T: X \rightarrow Y_\Psi$ linear.

Damit sind die Begriffe für lineare Abbildungen auf Ψ -semilineare Abbildungen übertragbar: Eine Ψ -semilineare Abbildung $T: (X, +, \cdot_R) \rightarrow (Y, +, \cdot_S)$ wird als lineare Abbildung von $(X, +, \cdot_R)$ nach Y_Ψ aufgefasst.

2.9.4 Satz (KLINGENBERG und KLEIN [186](1972), Seiten 80-85; BOURBAKI [34](1974), II.1.13). *Seien R und S zwei Ringe mit Eins, sei $\Psi: R \rightarrow S$ ein unitaler Ringhomomorphismus und seien X und Y zwei unitale S -Moduln. Dann gelten die folgende Aussagen:*

- (a) *Ist U ein Untermodul von X , dann ist U_Ψ ein Untermodul des R -Moduls X_Ψ . Falls Ψ surjektiv ist, gilt auch die Umkehrung: Jeder Untermodul des R -Moduls X_Ψ ist ein Untermodul des S -Moduls X .*
- (b) *Sind R und S zusätzlich Körper und A eine linear unabhängige Teilmenge von X , so ist A auch in X_Ψ linear unabhängig. Falls Ψ bijektiv ist, gilt auch die Umkehrung.*
- (c) *Ist A ein Erzeugendensystem von X_Ψ , dann ist A auch ein Erzeugendensystem von X . Falls Ψ surjektiv ist, gilt auch die Umkehrung.*
- (d) *Jede S -lineare Abbildung $X \rightarrow Y$ ist auch eine R -lineare Abbildung $X_\Psi \rightarrow Y_\Psi$. Sind R und S zusätzlich Körper, dann ist $(\text{Hom}_S(X, Y))_\Psi$ ein Unterraum des R -Vektorraumes $\text{Hom}_R(X_\Psi, Y_\Psi)$. Falls Ψ surjektiv ist, gilt die Gleichheit.*

2.9.5 Bemerkung. Um einzusehen, dass in Satz 2.9.4 jeweils die Umkehrung in (a), (b) und (c) und die Gleichheit in (d) nicht gelten muss, betrachte man $\mathbb{C}$ als Vektorraum über $\mathbb{R}$ mit Ψ als die kanonische Einbettung von $\mathbb{R}$ in $\mathbb{C}$.

2.9.6 Definition. Wenn in der Situation von 2.9.1 $R = \mathbb{R}$, $S = \mathbb{C}$ gilt und Ψ die kanonische Einbettung von $\mathbb{R}$ in $\mathbb{C}$ ist, dann heißt X_Ψ die *reelle Strukturierung* (im Englischen *realification*) von X und wird auch mit $X_{\mathbb{R}}$ bezeichnet.

2.9.7 Bemerkung. Sei X ein Vektorraum über $\mathbb{C}$. Da die komplexe Konjugation $\bar{\cdot}: \mathbb{C} \rightarrow \mathbb{C}$, $z \mapsto \bar{z}$ ein Ringhomomorphismus ist, ist mit $\Psi := \bar{\cdot}$ der sogenannte zu X *konjugiert-komplexe Vektorraum* $\bar{X} := X_\Psi$ über $\mathbb{C}$ definiert, das heißt, für alle $z \in \mathbb{C}$, $x \in X$ gilt dann $z \cdot_{\bar{X}} x = \bar{z} \cdot_X x$. Des Weiteren hat man dann $X_{\mathbb{R}} = (\bar{X})_{\mathbb{R}}$.

2.9.8. Sei $(Y, \|\cdot\|_Y)$ ein normierter Vektorraum über $\mathbb{C}$. Setze $X := Y_{\mathbb{R}}$. Bezeichne ι die identische Abbildung von der Menge X auf die Menge Y . Setze $\|x\|_X := \|\iota(x)\|_Y$ für alle $x \in X$. Dann ist $(X, \|\cdot\|_X)$ ein normierter Vektorraum über $\mathbb{R}$. Nach Definition von X existiert für jedes $\lambda \in \mathbb{C}$ und $x \in X$ auch das mit λx bezeichnete Element $\iota^{-1}(\lambda \iota(x))$; mit dieser Bezeichnung gilt also $\iota(\lambda x) = \lambda \iota(x)$. Wegen $x + y = \iota^{-1}(\iota(x) + \iota(y))$ gilt auch $\iota(x + y) = \iota(x) + \iota(y)$.

Sei γ die gemäß Satz 2.9.4(d) vorhandene lineare Einbettung

$$\gamma: (L(Y))_{\mathbb{R}} \hookrightarrow L(X), \quad T \mapsto \iota^{-1} \circ T \circ \iota.$$

Wegen $\iota(B_X) = B_Y$ hat man für alle $T \in \mathcal{L}(Y)$: $\|T\| = \sup_{y \in B_Y} \|Ty\|_Y = \sup_{x \in B_X} \|T\iota x\|_Y = \sup_{x \in B_X} \|\iota\gamma(T)x\|_Y = \sup_{x \in B_X} \|\gamma(T)x\|_X = \|\gamma(T)\|$. Somit ist die Abbildung $(\mathcal{L}(Y))_{\mathbb{R}} \rightarrow \mathcal{L}(X)$, $T \mapsto \gamma(T)$ eine Isometrie.

2.9.9 Lineare Funktionale (WERNER [291](2002), Seite 96, Lemma III.1.3; LI [198](2003), Seite 5, Satz 1.1.6; SCHAEFER [255](1967), Seite 32). Sei X ein Vektorraum über $\mathbb{C}$. Bezeichne r die identische Abbildung von X auf $X_{\mathbb{R}}$. Für alle $x^{\#} \in (X^{\#})_{\mathbb{R}}$ setze

$$(\operatorname{Re} x^{\#})(x) := \operatorname{Re} (x^{\#}(r^{-1}(x))) \quad \text{für alle } x \in X_{\mathbb{R}}.$$

Die dadurch definierte Abbildung

$$\operatorname{Re}: (X^{\#})_{\mathbb{R}} \rightarrow (X_{\mathbb{R}})^{\#}$$

ist linear und bijektiv. Dabei ist ihre Umkehrabbildung wie folgt erklärt:

$$\operatorname{Re}^{-1}(x^{\#})(x) := x^{\#}(r(x)) - ix^{\#}(r(ix)) \quad \text{für alle } x \in X.$$

Falls X ein topologischer Vektorraum über $\mathbb{C}$ ist, bildet die Abbildung Re den Unterraum $(X^*)_{\mathbb{R}}$ surjektiv auf $(X_{\mathbb{R}})^*$ ab und falls X ein normierter Vektorraum über $\mathbb{C}$ ist, ist diese Zuordnung auch isometrisch. Der somit für topologische Vektorräume X über $\mathbb{C}$ erklärte — im normierten Fall isometrische — Isomorphismus

$$(X^*)_{\mathbb{R}} \rightarrow (X_{\mathbb{R}})^*$$

wird ebenfalls mit Re bezeichnet.

2.9.10. Ist X ein Vektorraum über $\mathbb{R}$ und existiert eine $\mathbb{R}$ -lineare Abbildung $j: X \rightarrow X$ mit $j^2 = -\operatorname{Id}_X$, so ist es in Umkehrung zu der reellen Strukturierung eines Vektorraumes über $\mathbb{C}$ möglich, eine komplexe Vektorraumstruktur auf X wie folgt zu definieren.

Für jedes $a \in \mathbb{R}$ sind die zu a bestimmten Links-Multiplikationen L_a (Definition 2.1.18) Elemente aus dem Ring $L(X)$. Mit dem unitalen Ringhomomorphismus

$$\Psi: \mathbb{C} \rightarrow L(X), \quad z \mapsto L_{\operatorname{Re}(z)} + L_{\operatorname{Im}(z)} \circ j$$

und der Eigenschaft, dass X per $L(X) \times X \rightarrow X$, $(T, x) \mapsto T(x)$ ein unitaler $L(X)$ -Modul ist, ist X_{Ψ} nach 2.9.1 ein unitaler $\mathbb{C}$ -Modul.

2.9.11 Definition. (a) Der soeben in 2.9.10 definierte Vektorraum X_{Ψ} über $\mathbb{C}$ wird mit $X_{\mathbb{C}}$ bezeichnet und heißt die *komplexe Strukturierung von X* (bezüglich j). Eine Abbildung auf X , die von der Form des soeben dabei verwendeten j ist, wird auch als eine *komplexe Struktur auf X* bezeichnet.

(b) Sei X ein (normierter) Vektorraum (bzw. Algebra) über $\mathbb{R}$. Dann heißt X *komplex strukturierbar* oder auch *vom komplexen Typus*, falls ein (normierter) Vektorraum (bzw. Algebra) Y über $\mathbb{C}$ existiert, so dass X isomorph zu $Y_{\mathbb{R}}$ ist. Ist dabei speziell im Fall, dass X normiert ist, die besagte Isomorphie isometrisch, so heißt X *isometrisch komplex strukturierbar*.

(c) (INGELSTAM [145](1964)) Sei A eine (normierte) Algebra über $\mathbb{R}$. A heißt *vom reellen Typus*, falls A nicht vom komplexen Typus ist. A heißt *vom stark reellen Typus* (im Englischen *strictly real type*), wenn gilt:

$$-a^2 \text{ ist quasi-invertierbar} \quad \text{für alle } a \in A.$$

An dieser Stelle sei bemerkt, dass die Algebren vom stark reellen Typus genau den bei GELFAND [106](1941) definierten *reellen Ringen* entsprechen.

(d) (INGELSTAM [145](1964), Seite 253). Ein nichtassoziativer Ring $(R, +, \cdot)$ heißt vom *komplexen Typus*, wenn ein Gruppenendomorphismus j auf $(R, +)$ existiert mit $j^2 = -\text{Id}_R$ und

$$j(rs) = j(r) \cdot s = r \cdot j(s) \quad \text{für alle } r, s \in R.$$

(Siehe auch 2.1.28.)

Ist, wieder an 2.9.10 anschließend, andererseits X ein Vektorraum über $\mathbb{C}$, dann ist offensichtlich $x \mapsto ix$ eine komplexe Struktur auf $X_{\mathbb{R}}$, und es gilt: $X = (X_{\mathbb{R}})_{\mathbb{C}}$.

Damit gilt:

2.9.12 Satz. *Sei X ein Vektorraum über $\mathbb{R}$. Dann sind äquivalent:*

- (a) *Es gibt eine $\mathbb{R}$ -lineare Abbildung $j: X \rightarrow X$ mit $j^2 = -\text{Id}_X$.*
- (b) *X ist komplex strukturierbar.*

2.9.13 Korollar. *Ist X ein komplex strukturierbarer Vektorraum über $\mathbb{R}$, dann gilt $X = (X_{\mathbb{C}})_{\mathbb{R}}$.*

2.9.14. Manchmal ist die folgende Schreibweise praktisch: Sei X ein Vektorraum über $\mathbb{C}$, dann ist, falls nicht ausdrücklich etwas anderes vereinbart ist, mit $X_{\mathbb{C}}$ der Raum X_{Ψ} mit $\Psi = \text{Id}_{\mathbb{C}}$ gemeint, mit anderen Worten, wieder X selbst.

2.9.15 Satz (INGELSTAM [145](1964), Seite 249). *Sei A eine Algebra über $\mathbb{R}$ mit Eins. Dann ist A vom reellen Typus, falls sie vom stark reellen Typus ist.*

2.9.16 Satz (INGELSTAM [145](1964), Seite 249). *Sei A eine (normierte) Algebra über $\mathbb{R}$. Dann sind äquivalent:*

- (a) *Es gibt eine $\mathbb{R}$ -lineare (beschränkte) Abbildung $j: A \rightarrow A$ mit $j^2 = -\text{Id}_A$ und*

$$j(ab) = j(a) \cdot b = a \cdot j(b) \quad \text{für alle } a, b \in A.$$
- (b) *A ist komplex strukturierbar.*

Besitzt A eine Eins, dann ist zu (a) und (b) zusätzlich äquivalent:

- (c) *Es existiert ein Element a im Zentrum von A mit $a^2 = -e$.*

2.9.17 Bemerkung (INGELSTAM [145](1964)). Ist A eine Algebra über $\mathbb{R}$, dann besitzt A^1 keine komplexe Strukturierung.

2.9.18 Satz (KAPLANSKY [170](1949), Seite 400). *Sei $(A, \|\cdot\|)$ eine normierte Algebra über $\mathbb{R}$, welche komplexe Skalare in der Weise zulässt, dass die Multiplikation mit i stetig ist. Dann ist*

$$\|a\|_{\mathbb{C}} := \max_{\lambda \in S_{\mathbb{C}}} \|\lambda a\| \quad \text{für alle } a \in A,$$

eine äquivalente Norm, die A zu einer normierten Algebra über $\mathbb{C}$ macht.

2.9.19 Bemerkung. Ein endlichdimensionaler Vektorraum über $\mathbb{R}$ ist genau dann komplex strukturierbar, wenn seine Dimension n geradzahlig ist, also $n \in 2\mathbb{N}$ (DIEUDONNÉ [67](1952)).

Jeder unendlichdimensionale Vektorraum über $\mathbb{R}$ ist komplex strukturierbar (INGELSTAM [145](1964), Seite 249).

(ROSENTHAL [249](1984)). Sei X ein Banachraum über $\mathbb{R}$ und $j \in \mathcal{L}(X)$. Der im Folgenden verwendete Begriff *schiefhermitesch* wird weiter unten in 2.11.22 definiert. Dann sind äquivalent:

- (a) Es gilt $j^2 = -\text{Id}_X$ und für alle $\alpha, \beta \in \mathbb{R}$ mit $\alpha^2 + \beta^2 = 1$ ist $\alpha \text{Id}_X + \beta j \in \mathcal{L}(X)$ isometrisch.
- (b) Es gilt $j^2 = -\text{Id}_X$ und j ist schiefhermitesch (siehe 2.11.22).
- (c) j ist eine schiefhermitesche (siehe 2.11.22) Isometrie (wobei man a priori j nicht als linear anzunehmen braucht).
- (d) j ist eine Isometrie (die man a priori nicht als linear anzunehmen braucht) und für alle $\alpha, \beta \in \mathbb{R}$ mit $\alpha^2 + \beta^2 = 1$ gilt $\|\alpha \text{Id}_X + \beta j\| = 1$.

Genau dann existiert ein $j \in \mathcal{L}(X)$, das eine (und damit alle) der vorstehenden Äquivalenzen erfüllt, wenn X isometrisch komplex strukturierbar ist (Definition 2.9.11).

Mit DIEUDONNÉ [67](1952) hat man, dass es auf dem sogenannten *James-Raum* J (siehe MEGGINSON [211](1998), Kapitel 4.5 und Übung 4.45) kein $j \in \mathcal{L}(X)$ mit $j^2 = -\text{Id}_X$ existiert und mithin J also ein Beispiel eines unendlichdimensionalen Banachraumes über $\mathbb{R}$ ist, der nicht vollständig komplex strukturierbar ist.

2.9.20 Streng komplexe Strukturen I (GÄHLER und GÄHLER [105](1974)). (a) Sei X ein Vektorraum über $\mathbb{R}$, der mit einer komplexen Struktur j ausgestattet ist. Um mit den Elementen von X ähnlich wie mit den komplexen Zahlen rechnen zu können, muss neben j noch eine weitere lineare Abbildung $k \in L(X)$ gegeben sein, die die Rolle der Konjugiertenbildung übernimmt; als Eigenschaften von solch einem k wird gefordert, dass sie Periode 2 hat, das heißt, $k \circ k = \text{Id}_X$, und, dass

$$j \circ k = -k \circ j \quad (2.35)$$

gilt. Das Paar (j, k) heißt dann eine *streng komplexe Struktur* auf X , und weiter nennt man dann X einen *streng komplexen Vektorraum*.

Sei nun X ein streng komplexer Vektorraum mit der zugehörigen streng komplexen Struktur (j, k) . Per

$$\text{Re}(x) := \frac{x + k(x)}{2} \quad \text{und} \quad \text{Im}(x) := \frac{-j(x - k(x))}{2}$$

für alle $x \in X$, sind zwei Abbildungen $X \rightarrow X$ definiert, die man als *Realteil-* bzw. *Imaginärteilmapping* auffassen kann. Es gelten die folgenden Gleichungen: $\text{Id}_X = \text{Re} + j \circ \text{Im}$, $\text{Re} \circ \text{Re} = \text{Re}$, $\text{Re} \circ \text{Im} = \text{Im}$, $\text{Im} \circ \text{Re} = 0$, $\text{Im} \circ \text{Im} = 0$, $\text{Re} \circ j \circ \text{Re} = 0$, $\text{Re} \circ j \circ \text{Im} = 0$, $\text{Im} \circ j \circ \text{Re} = \text{Re}$, $\text{Im} \circ j \circ \text{Im} = \text{Im}$, $k \circ \text{Re} = \text{Re}$, $k \circ j \circ \text{Im} = -j \circ \text{Im}$. Folglich gilt $\text{Re}(X) = \text{Im}(X)$ (Dazu: $x \in \text{Im}(X) \Leftrightarrow -\frac{1}{2}j(z - k(z))$ für ein $z \in X$. Es ist $z = j(y)$ für ein $y \in X$. Also $x = -\frac{1}{2}j(j(y) + jk(y)) = \frac{1}{2}(y + k(y)) \in \text{Re}(X)$) und für die reellen Teilräume $\text{Re}(X)$ und $(j \circ \text{Im})(X)$ gilt $X = \text{Re}(X) \dot{+} (j \circ \text{Im})(X)$.

(b) Sei X ein Vektorraum über $\mathbb{R}$ und j eine komplexe Struktur auf X . GÄHLER und GÄHLER, ebd., Seite 65, zeigen, dass durch eine geeignete Wahl einer Basis von X die Abbildungsmatrix der komplexen Struktur j eine besonders einfache Gestalt annimmt. Als unmittelbare Folgerung haben sie die Aussage, dass ein reeller Untervektorraum U von X mit $X = U \dot{+} j(U)$ existiert. Definiert man dann die Abbildung $k: X \rightarrow X$, $x = u + v \mapsto u - v$, wobei $u \in U$ und $v \in j(U)$ sein soll, so ist k wegen $jk(x) = jk(u + v) = j(u - v) = -(jv - ju) = -k(ju + jv) = -kj(u + v) = -kj(x)$ derart, dass (j, k) eine streng komplexe Struktur auf X ist. Da dann zum Beispiel auch $(j, -k)$ eine streng komplexe Struktur auf X ist, ist diese nicht eindeutig bestimmt.

2.9.21 (KLINGENBERG und KLEIN [186](1972), Seite 218). Sei X ein Vektorraum über $\mathbb{R}$ und Z das direkte Produkt $X \times X$. Setze

$$j: Z \rightarrow Z, \quad (x, y) \mapsto (-y, x).$$

Dann ist offenbar $j^2 = -\text{Id}_Z$, so dass es also auf Z per $Z_{\mathbb{C}}$ die Struktur eines Vektorraumes über $\mathbb{C}$ gibt mit $Z = (Z_{\mathbb{C}})_{\mathbb{R}}$.

2.9.22 Definition (KLINGENBERG und KLEIN [186](1972), Seite 218). Den soeben definierten Vektorraum $Z_{\mathbb{C}}$ nennt man die *Komplexifizierung von X* und bezeichnet ihn mit $X_{(\mathbb{C})}$.

2.9.23 Definition. Sei X ein Vektorraum über $\mathbb{C}$. Dann heißt jeder Vektorraum Y über $\mathbb{R}$ für den $Y_{(\mathbb{C})} \simeq X$ gilt, eine *reelle Form des Vektorraumes X* über $\mathbb{C}$.

2.9.24 Bemerkung. Auf $X_{(\mathbb{C})}$ gilt für alle $\alpha, \beta \in \mathbb{R}, x, y \in X$:

$$\begin{aligned} (\alpha + \beta i)(x, y) &= \alpha(x, y) + \beta j(x, y) = \alpha(x, y) + \beta(-y, x) \\ &= (\alpha x - \beta y, \alpha y + \beta x). \end{aligned}$$

Also zum Beispiel: $\lambda(e, 0) = (\text{Re}(\lambda)e, \text{Im}(\lambda)e)$ für $\lambda \in \mathbb{C}$.

2.9.25 Bemerkung (KLINGENBERG und KLEIN [186](1972), Seite 219). Sei X ein Vektorraum über $\mathbb{R}$. Die Abbildung $k: X \rightarrow X_{(\mathbb{C})\mathbb{R}}, x \mapsto (x, 0)$, ist eine injektive $\mathbb{R}$ -lineare Abbildung. Daher ist $k(X) \simeq X$, und man identifiziert X mit dem Teilraum $k(X)$ des Vektorraumes $X_{(\mathbb{C})\mathbb{R}}$ über $\mathbb{R}$.

Wegen $ix = (0, x)$ für alle $x \in X$ ist dann $(x, y) = x + iy$ für alle $x, y \in X$, so dass man auch $X_{(\mathbb{C})\mathbb{R}} = X \oplus iX$ schreiben kann. Jedes $z \in X_{(\mathbb{C})}$ besitzt also eine kartesische Darstellung $z = x + iy$, das heißt, jedes $z \in X_{(\mathbb{C})}$ lässt sich eindeutig schreiben als $z = x + iy$, und man bezeichnet $\text{Re}(z) := x$ als Realteil von z und $\text{Im}(z) := y$ als Imaginärteil von z (bezüglich der Komplexifizierung).

(GOLDHABER und EHRLICH [115](1970), Kapitel 3.4). Ist R ein Ring, X ein R -Linksmodul und Y ein R -Rechtsmodul, dann wird das *Tensorprodukt* von X und Y über R mit $X \otimes_R Y$ bezeichnet; es ist ein $\mathbb{Z}$ -Modul. Für den kommutativen Fall siehe zum Beispiel auch KOWALSKY und MICHLER [188](2003).

Ist S ein kommutativer Ring mit Eins, die mit e_S bezeichnet sei, und R ein Unterring von S mit der gleichen Eins, und ist M ein freier R -Modul mit Basis $b_\nu, \nu \in I$, dann ist $S \otimes_R M$ ein freier S -Modul mit Basis $e_S \otimes b_\nu, \nu \in I$. Der R -Modul M ist R -isomorph zu einem Untermodul von $S \otimes_R M$, aufgefasst als ein R -Modul. Man sagt, der freie R -Modul M würde zu einem freien S -Modul *gelifet*. Speziell für $R = \mathbb{R}$ und $S = \mathbb{C}$ hat man:

Ist X ein (normierter) Vektorraum über $\mathbb{R}$, dann ist $X_{(\mathbb{C})}$ (isometrisch) isomorph zu dem Vektorraum $\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} X$ über $\mathbb{C}$. Ist $(b_\nu)_{\nu \in I}$ eine Basis von X , dann ist $(1 \otimes b_\nu)_{\nu \in I}$ eine Basis von $\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} X$. Identifiziere X als Teilmenge von $\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} X$ per $x \mapsto 1 \otimes x$. Sei Z wie in 2.9.21. Die $\mathbb{R}$ -lineare Abbildung $(x, y) \mapsto x + iy$ (genauer $(x, y) \mapsto 1 \otimes x + i \otimes y$ mit für $w, z \in \mathbb{C}: w(z \otimes x) := (wz) \otimes x$) ist dann eine Bijektion von Z auf $\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} X$, durch die die Produkttopologie von Z auf $\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} X$ übertragen wird.

2.9.26. Seien X und Y zwei Vektorräume über $\mathbb{R}$. Sei $T \in L(X, Y)$. Dann ist die *kanonische komplex-lineare Fortsetzung* von T die Abbildung

$$X_{(\mathbb{C})} \rightarrow X_{(\mathbb{C})}, x + iy \mapsto T(x) + iT(y);$$

sie wird mit $T_{(\mathbb{C})}$ bezeichnet.

Sei $n \in \mathbb{N}^\times$. Seien X und Y zwei n -dimensionale Vektorräume über $\mathbb{R}$. Sei $T \in L(X, Y)$. Sei $B := \{b_1, \dots, b_n\}$ eine Basis von X und $C := \{c_1, \dots, c_n\}$ eine Basis von Y . Sei $A = (a_{ij})$ die Matrix von T bezüglich der Basen B und C . Dann gilt für alle $i = 1, \dots, n$: $T_{(\mathbb{C})}(1 \otimes b_i) = 1 \otimes T(b_i) = 1 \otimes \sum_{j=1}^n a_{ji} c_j = \sum_{j=1}^n a_{ji} 1 \otimes c_j$, das heißt, A ist auch die Matrix von $T_{(\mathbb{C})}$ bezüglich der Basen $B_{(\mathbb{C})} := \{1 \otimes b_1, \dots, 1 \otimes b_n\}$ und $C_{(\mathbb{C})} := \{1 \otimes c_1, \dots, 1 \otimes c_n\}$.

2.9.27 Bemerkung und Definition. (a) Sei A eine Algebra über $\mathbb{R}$. Dann ist $\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} A$ mit dem kanonischen Produkt $(w \otimes x) \cdot (z \otimes y) := wz \otimes xy$, $w, z \in \mathbb{C}$, $x, y \in A$, eine Algebra über $\mathbb{C}$ (McCRIMMON [210](2004), Seite 69). Dementsprechend versteht man $A_{(\mathbb{C})}$ mit dem Produkt

$$(a, b)(c, d) := (ac - bd, ad + bc) \quad \text{für alle } a, b, c, d \in A$$

durch das $A_{(\mathbb{C})}$ eine Algebra über $\mathbb{C}$ ist, die wieder mit $A_{(\mathbb{C})}$ bezeichnet wird. Hat A eine Eins e , dann ist $(e, 0)$ die Eins in $A_{(\mathbb{C})}$, das heißt in Tensorschreibweise, dass dann $1 \otimes e$ die Eins in $\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} A$ ist.

(b) (KAPLANSKY [170](1949), Seite 400 und RICKART [246](1960), Seite 6). Sei $(A, \|\cdot\|)$ eine normierte Algebra über $\mathbb{R}$. Dann ist $A_{(\mathbb{C})\mathbb{R}}$ ausgestattet mit der Norm

$$\|(a, b)\|_{1\mathbb{R}} := \|a\| + \|b\| \quad \text{für alle } a, b \in A$$

eine normierte Algebra über $\mathbb{R}$, deren Multiplikation mit den komplexen Zahlen stetig ist. $\|\cdot\|_{1\mathbb{R}}$ ist genau dann vollständig, wenn die Norm von A vollständig ist.

Renormiert man $(A_{(\mathbb{C})\mathbb{R}}, \|\cdot\|_{1\mathbb{R}})$ gemäß Satz 2.9.18, also setzt man

$$\|(a, b)\|_{1\mathbb{C}} := \max_{\lambda \in S_{\mathbb{C}}} \|\lambda(a, b)\|_{1\mathbb{R}} \quad \text{für alle } a, b \in A,$$

dann ist $(A_{(\mathbb{C})}, \|\cdot\|_{1\mathbb{C}})$ eine normierte Algebra über $\mathbb{C}$. Wegen $\|(a, 0)\| = \sqrt{2}\|a\|$ für alle $a \in A$ ist dann A zwar per $a \mapsto (a, 0)$ isometrisch isomorph in dem Vektorraum $(A_{(\mathbb{C})\mathbb{R}}, \frac{1}{\sqrt{2}}\|\cdot\|_{1\mathbb{C}})$ über $\mathbb{R}$ eingebettet, aber $\frac{1}{\sqrt{2}}\|\cdot\|_{1\mathbb{C}}$ ist im Allgemeinen keine Algebranorm mehr.

(c) Sei $(A, \|\cdot\|)$ eine normierte Algebra über $\mathbb{R}$. Setze $\|\cdot\|_1 := \frac{1}{\sqrt{2}}\|\cdot\|_{1\mathbb{C}}$. Dann ist $A_{(\mathbb{C})}$ mit der Norm

$$\|a\| := \sup \left\{ \|ab\|_1 : b \in B_{(A_{(\mathbb{C})}, \|\cdot\|_1)} \right\} \quad \text{für alle } a \in A_{(\mathbb{C})}$$

eine normierte Algebra über $\mathbb{C}$ und A ist kanonisch isometrisch isomorph in $A_{(\mathbb{C})\mathbb{R}}$ eingebettet.

(d) Sei $(A, \|\cdot\|)$ eine normierte Algebra über $\mathbb{R}$. Auf $A_{(\mathbb{C})\mathbb{R}}$ sind die zwei Normen $\|\cdot\|_{1\mathbb{R}}$ und $\|\cdot\|_{1\mathbb{C}}$ von (b) und die Norm $\|\cdot\|$ von (c) äquivalent.

(e) (LI [198](2003), Seite 15). Sei $(A, \|\cdot\|)$ eine unitale, normierte Algebra über $\mathbb{R}$. Anstelle der ℓ_1 -Norm in (b) kann man $A_{(\mathbb{C})\mathbb{R}}$ auch mit einer ℓ_p -Norm, $1 \leq p < \infty$,

$$\|(a, b)\|_{p\mathbb{R}} := \left(\|a\|^p + \|b\|^p \right)^{1/p} \quad \text{für alle } a, b \in A$$

oder mit

$$\|(a, b)\|_{\infty\mathbb{R}} := \max \left\{ \|a\|, \|b\| \right\} \quad \text{für alle } a, b \in A$$

ausstatten und wieder gemäß Satz 2.9.18 zu den entsprechenden $\|\cdot\|_{p\mathbb{C}}$, $1 \leq p \leq \infty$ renormieren. Die Normen

$$\|(a, b)\|_p := \frac{1}{c_p} \|(a, b)\|_{p\mathbb{C}} \quad \text{für alle } a, b \in A$$

mit $c_p := \sup \left\{ \left(|\cos \theta|^p + |\sin \theta|^p \right)^{1/p} : \theta \in \mathbb{R} \right\}$ für $1 \leq p < \infty$ und $c_{\infty} := 1$, machen dann ebenfalls $A_{(\mathbb{C})}$ zu einer normierten Algebra über $\mathbb{C}$, so dass A kanonisch isometrisch isomorph in $A_{(\mathbb{C})\mathbb{R}}$ eingebettet ist. Alle Normen $\|\cdot\|_p$, $1 \leq p \leq \infty$, sind auf $A_{(\mathbb{C})}$ zu $\|\cdot\|_1$ äquivalent.

(f) Die Normen $\|\cdot\|_p$, $1 \leq p \leq \infty$, gebrauchen zur Konstruktion nur einen normierten Vektorraum über $\mathbb{R}$. Somit sind für alle normierten Vektorräume X über $\mathbb{R}$ die Räume $(X_{(\mathbb{C})}, \|\cdot\|_{p\mathbb{C}})$ normierte Vektorräume über $\mathbb{C}$, mit der Eigenschaft, dass X kanonisch isometrisch isomorph in $X_{(\mathbb{C})\mathbb{R}}$ eingebettet ist.

(g) (BONSALL und DUNCAN [29](1971), Seite 69). Sei A eine normierte Algebra über $\mathbb{R}$. Sei p das Minkowski-Funktional der absolutkonvexen Hülle (über $\mathbb{R}$) von $\{a \in A : \|a\| < 1\} \times \{0\}$, dann ist $A_{(\mathbb{C})}$ mit p als Norm eine normierte Algebra über $\mathbb{C}$. Die Abbildung $a \mapsto (a, 0)$ ist ein isometrischer Monomorphismus von A nach $(A_{(\mathbb{C})\mathbb{R}}, p)$. p heißt *Minkowski-Norm* und wird auch mit $\|\cdot\|_M$ bezeichnet.

(h) Auf $A_{(\mathbb{C})\mathbb{R}}$ sind die Normen $\|\cdot\|_M$ und $\|\cdot\|_{1\mathbb{R}}$ äquivalent. Allgemein gilt: Ist A mit einer vollständigen Norm versehen, dann sind alle vollständigen Algebranormen auf $A_{(\mathbb{C})}$, bei der A kanonisch isometrisch isomorph in $A_{(\mathbb{C})\mathbb{R}}$ eingebettet ist, äquivalent.

2.9.28. Sei X ein Vektorraum über $\mathbb{R}$. Jede der in 2.9.27 erklärten Norm erzeugt auf $X_{(\mathbb{C})\mathbb{R}}$ seine Produkttopologie.

2.9.29 Definition (INGELSTAM [144](1962), Seiten 23 und 24).

(a) Sei X ein normierter Vektorraum über $\mathbb{K}$. Ein Punkt des Randes einer konvexen Menge M in X heißt ein *Vertex nach Ingelstam*, wenn keine gerade Linie (also ein eindimensionaler affiner Raum über $\mathbb{K}$) durch diesen Punkt eine Tangente dieser Menge ist. (Siehe auch INGELSTAM [145](1964), Seite 260, für den Fall, dass X eine unitale Algebra ist, die Menge M die Einheitssphäre ist und die dann mögliche äquivalente Formulierung mittels des Funktional

$$x \mapsto \phi(x) := \nabla_{\{x\}} \|\cdot\|(e);$$

bezüglich des Symbols ∇ siehe 4.1.4.). Siehe auch 4.1.9.

(b) Eine Banachalgebra über $\mathbb{K}$ mit Eins hat die *Vertex-Eigenschaft nach Ingelstam*, wenn für alle Algebranormen $\|\cdot\|$ mit $\|e\| = 1$, die die Topologie definieren, die Eins ein Vertex nach Ingelstam der Einheitssphäre ist.

2.9.30 Satz (INGELSTAM [144](1962), Seiten 22 und 29). *Sei A eine Banachalgebra über $\mathbb{K}$ mit Eins. Dann gilt:*

- (a) *Ist $\mathbb{K} = \mathbb{R}$ und A vom stark reellen Typus, dann hat A die Vertex-Eigenschaft nach Ingelstam.*
- (b) *Ist $\mathbb{K} = \mathbb{C}$, dann hat A die Vertex-Eigenschaft nach Ingelstam.*

2.9.31 Konvexe Funktionen (VAN TIEL [277](1984)). Sei X ein Vektorraum über $\mathbb{K}$. Sei C eine konvexe Teilmenge von X . Eine Abbildung $f: C \rightarrow \mathbb{R} \cup \{\infty\}$ heißt *konvex*, wenn für alle $x, y \in C$ und $0 < \lambda < 1$ die Ungleichung

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y)$$

erfüllt ist. Eine Funktion $f: C \rightarrow \mathbb{R}$ heißt *konvex*, wenn sie, aufgefasst als eine Abbildung $C \rightarrow \mathbb{R} \cup \{\infty\}$, konvex ist. Eine Abbildung $f: C \rightarrow \overline{\mathbb{R}}$ heißt *konvex*, wenn für alle $x, y \in C$, $\alpha \in \mathbb{R}$ mit $f(x) < \alpha$, $\beta \in \mathbb{R}$ mit $f(y) < \beta$, $0 < \lambda < 1$ die Ungleichung

$$f(\lambda x + (1 - \lambda)y) < \lambda \cdot \alpha + (1 - \lambda) \cdot \beta$$

erfüllt ist. Zum Beispiel ist jede Halbnorm konvex. Allgemeiner gilt: Ist A eine konvexe, nicht leere Teilmenge von X , so ist $x \mapsto \text{dist}(x, A)$ konvex (ebd., Seite 66). Des Weiteren ist für jede bilineare, positive, reellwertige Bilinearform B auf X die Abbildung $x \mapsto B(x, x)$ konvex.

Eine Abbildung $f: C \rightarrow \mathbb{R} \cup \{\infty\}$ ist genau dann konvex, wenn sie, aufgefasst als eine Abbildung $f: C \rightarrow \overline{\mathbb{R}}$, konvex ist.

Beweis. Sei $f: C \rightarrow \mathbb{R} \cup \{\infty\}$ konvex. Seien $x, y \in C$, sei $0 < \lambda < 1$, $\alpha \in \mathbb{R}$ mit $f(x) < \alpha$ und $\beta \in \mathbb{R}$ mit $f(y) < \beta$. Dann gilt $f(\lambda x + (1 - \lambda)y) < \lambda \cdot \alpha + (1 - \lambda) \cdot \beta$, das heißt, f , aufgefasst als $f: C \rightarrow \overline{\mathbb{R}}$, ist konvex.

Sei nun umgekehrt f eine Abbildung $C \rightarrow \mathbb{R} \cup \{\infty\}$, die, aufgefasst als eine Abbildung $C \rightarrow \overline{\mathbb{R}}$, konvex sei. Seien $x, y \in C$ und $0 < \lambda < 1$. Dann ist zu zeigen, dass die Ungleichung $f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y)$ erfüllt ist. Falls $f(x)$ oder $f(y)$ gleich ∞ ist, ist diese Ungleichung erfüllt. Bleibt der Fall zu betrachten, wo weder $f(x)$ noch $f(y)$ gleich ∞ ist: Sei $\varepsilon > 0$. Dann gilt $f(\lambda x + (1 - \lambda)y) < \lambda \cdot (f(x) + \varepsilon) + (1 - \lambda) \cdot (f(y) + \varepsilon) = \lambda \cdot f(x) + (1 - \lambda) \cdot f(y) + \varepsilon$. Folglich hat man die besagte Ungleichung vorliegen und f ist konvex. $\square$

Sei $f: C \rightarrow \overline{\mathbb{R}}$ eine konvexe Abbildung. Die mit $\text{dom}_{\text{eff}}(f)$ bezeichnete Menge $\{x \in X : f(x) < \infty\}$ heißt der *effektive Definitionsbereich* von f . Die Abbildung f heißt *proper*, wenn $\text{dom}_{\text{eff}}(f) \neq \emptyset$ und $\text{ran}(f) \subseteq \mathbb{R} \cup \{\infty\}$ gilt.

Der effektive Definitionsbereich von f ist eine konvexe Menge. Jede propere, konvexe Funktion f kann aufgefasst werden als eine Erweiterung einer konvexen Funktion $g: \text{dom}_{\text{eff}}(f) \rightarrow \mathbb{R}$ per $f(x) = g(x)$ für alle $x \in \text{dom}_{\text{eff}}(f)$ und $f(x) = \infty$ sonst. In diesem Sinne kann man jede konvexe Funktion $f: C \rightarrow \overline{\mathbb{R}}$ als eine konvexe Abbildung $X \rightarrow \overline{\mathbb{R}}$ auffassen.

Sei $f: C \rightarrow \overline{\mathbb{R}}$ eine Abbildung. Dann sind die folgenden Aussagen äquivalent:

- (a) f ist konvex.
- (b) $\text{epi}(f)$, aufgefasst als Teilmenge von $X_{\mathbb{R}} \times \mathbb{R}$, ist konvex.
- (c) $\{(x, y) \in X \times \mathbb{R} : y > f(x)\}$, aufgefasst als Teilmenge von $X_{\mathbb{R}} \times \mathbb{R}$, ist konvex.
- (d) Für jeden eindimensionalen, affinen Untervektorraum U von $X_{\mathbb{R}}$ ist die Abbildung $f \upharpoonright C \cap U$ konvex.

Beweis. (a) $\Rightarrow$ (c): Sei f konvex. Setze $A := \{(x, y) \in X \times \mathbb{R} : y > f(x)\}$. Seien $(x_1, y_1), (x_2, y_2) \in A$. Also $y_1 > f(x_1)$ und $y_2 > f(x_2)$. Sei $0 < \lambda < 1$. Nach Voraussetzung also $\lambda \cdot y_1 + (1 - \lambda) \cdot y_2 > f(\lambda x_1 + (1 - \lambda)x_2)$, das heißt, der Punkt $\lambda \cdot (x_1, y_1) + (1 - \lambda) \cdot (x_2, y_2) = (\lambda x_1 + (1 - \lambda)x_2, \lambda y_1 + (1 - \lambda)y_2)$ ist in A .

(c) $\Rightarrow$ (a): Sei $A := \{(x, y) \in X_{\mathbb{R}} \times \mathbb{R} : y > f(x)\}$ konvex. Sei $f(x) < \alpha$, $\alpha \in \mathbb{R}$, $f(y) < \beta$, $\beta \in \mathbb{R}$. Also $(x, \alpha), (y, \beta) \in A$. Sei $0 < \lambda < 1$. Nach Voraussetzung ist der Punkt $\lambda \cdot (x, \alpha) + (1 - \lambda) \cdot (y, \beta)$ in A , das heißt, $f(\lambda x + (1 - \lambda)y) < \lambda \cdot \alpha + (1 - \lambda) \cdot \beta$.

(b) $\Rightarrow$ (c): Sei $\text{epi}(f) \subseteq X_{\mathbb{R}} \times \mathbb{R}$ konvex. Setze $A := \{(x, y) \in X \times \mathbb{R} : y > f(x)\}$. Seien $(x, \alpha), (y, \beta) \in A$ und sei $0 < \lambda < 1$. Also $\alpha > f(x)$ und $\beta > f(y)$. Seien $\alpha_0, \beta_0 \in \mathbb{R}$ mit $\alpha > \alpha_0 \geq f(x)$ und $\beta > \beta_0 \geq f(y)$. Also $(x, \alpha_0), (y, \beta_0) \in \text{epi}(f)$ und nach Voraussetzung somit $(\lambda x + (1 - \lambda)y, \lambda \alpha_0 + (1 - \lambda)\beta_0) \in \text{epi}(f)$; das heißt, $f(\lambda x + (1 - \lambda)y) \leq \lambda \alpha_0 + (1 - \lambda)\beta_0 < \lambda \alpha + (1 - \lambda)\beta$, also $(\lambda x + (1 - \lambda)y, \lambda \alpha + (1 - \lambda)\beta) = \lambda \cdot (x, \alpha) + (1 - \lambda) \cdot (y, \beta) \in A$.

(c) $\Rightarrow$ (b): Sei $A := \{(x, y) \in X_{\mathbb{R}} \times \mathbb{R} : y > f(x)\}$ konvex. Seien $(x, \alpha), (y, \beta) \in \text{epi}(f)$ und $0 < \lambda < 1$. Sei $\varepsilon > 0$ beliebig. Dann gilt $\alpha + \varepsilon > f(x)$ und $\beta + \varepsilon > f(y)$, das heißt, $(x, \alpha + \varepsilon), (y, \beta + \varepsilon) \in A$. Nach Voraussetzung gilt $f(\lambda x + (1 - \lambda)y) < \lambda(\alpha + \varepsilon) + (1 - \lambda)(\beta + \varepsilon) = \lambda \alpha + (1 - \lambda)\beta + \varepsilon$. Folglich gilt $f(\lambda x + (1 - \lambda)y) \leq \lambda \alpha + (1 - \lambda)\beta$, das heißt, $\lambda \cdot (x, \alpha) + (1 - \lambda) \cdot (y, \beta) \in \text{epi}(f)$.

(a) $\Leftrightarrow$ (d): Klar nach Definition von (a). $\square$

2.9.32 Unterhalbstetige $\overline{\mathbb{R}}$ -wertige Abbildungen (CHOQUET [45](1969), Seite 339). Sei X ein topologischer Raum und $f: X \rightarrow \overline{\mathbb{R}}$ eine Abbildung. Dann sind die folgenden drei Aussagen äquivalent:

- (a) f ist unterhalbstetig auf X . (Definition in 1.2.9)
- (b) Für jedes Netz (x_ν) in X und jedes $x \in X$ mit $x_\nu \rightarrow x$ gilt die Ungleichung $f(x) \leq \liminf_\nu f(x_\nu)$.
- (c) $\text{epi}(f)$, aufgefasst als Teilmenge von $X \times \mathbb{R}$, ist abgeschlossen.

Beweis. (a) $\Leftrightarrow$ (b): Klar nach Definition.

(c) $\Rightarrow$ (a): Sei $\text{epi}(f)$ in $X \times \mathbb{R}$ abgeschlossen. Sei $\lambda \in \mathbb{R}$ beliebig, aber fest. Dann ist die Menge $A := \{x \in X : f(x) \leq \lambda\} \times \{\lambda\} = \text{epi}(f) \cap \{(x, y) \in X \times \mathbb{R} : y = \lambda\}$ wegen $\{(x, y) \in X \times \mathbb{R} : y = \lambda\} = (X \times \mathbb{R}) \setminus ((X \times \{y \in \mathbb{R} : y < \lambda\}) \cup (X \times \{y \in \mathbb{R} : y > \lambda\}))$ abgeschlossen in $X \times \mathbb{R}$. Bezeichne p die kanonische Projektion $X \times \mathbb{R} \rightarrow X$. Dann ist $p(A) = \{x \in X : f(x) \leq \lambda\}$. Sei $x_0 \in \{x \in X : f(x) > \lambda\}$. Also $(x_0, \lambda) \notin A$. Da A abgeschlossen ist, ist (x_0, λ) kein Berührungspunkt von A , das heißt, es gibt in $X \times \mathbb{R}$ eine offene Menge G mit $(x_0, \lambda) \in G$ und $G \cap A = \emptyset$. Da p eine offene Abbildung ist, ist $p(G)$ offen in X . Wegen $G \cap A = \emptyset$ ist auch $p(G) \cap p(A) = \emptyset$. Somit ist x_0 per $x_0 \in p(G)$ ein innerer Punkt von $\{x \in X : f(x) > \lambda\}$. Folglich ist für jedes $\lambda \in \mathbb{R}$ die Menge $\{x \in X : f(x) > \lambda\}$ offen. Bleibt zu zeigen, dass die beiden Mengen $\{x \in X : f(x) > -\infty\}$ und $\{x \in X : f(x) > \infty\}$ offen sind. Die letztere dieser beiden Mengen ist leer, also offen. Die Menge $A := \{x \in X : f(x) \leq -\infty\} \times \{-\infty\} = \text{epi}(f) \cap (X \times \{-\infty\})$ ist wegen $X \times \{-\infty\} = (X \times \overline{\mathbb{R}}) \setminus (X \times \{y \in \overline{\mathbb{R}} : y > -\infty\})$ abgeschlossen in $X \times \overline{\mathbb{R}}$. Bezeichne p die kanonische Projektion $X \times \overline{\mathbb{R}} \rightarrow X$. Dann ist $p(A) = \{x \in X : f(x) \leq -\infty\}$. Ganz analog wie hier eben für alle $\lambda \in \mathbb{R}$ gezeigt, folgert man nun für $\lambda = -\infty$, dass die Menge $\{x \in X : f(x) > -\infty\}$ offen ist. Damit ist f als unterhalbstetig auf X gezeigt.

(a) $\Rightarrow$ (c): Sei f unterhalbstetig auf X . Die Menge $\{x \in X : f(x) \leq \lambda\}$ ist also für jedes $\lambda \in \overline{\mathbb{R}}$ abgeschlossen. Dies ist äquivalent dazu, dass für jedes $\lambda \in \overline{\mathbb{R}}$ gilt: Ist $(x_\nu)_{\nu \in I}$ ein Netz in X mit $f(x_\nu) \leq \lambda$ für alle $\nu \in I$ und konvergiert dieses Netz gegen ein $x \in X$, so ist $f(x) \leq \lambda$. Sei nun $(x_\nu, y_\nu)_{\nu \in I}$ ein Netz in $\text{epi}(f)$, welches gegen ein $(x, y) \in X \times \mathbb{R}$ konvergiert. Also $x_\nu \rightarrow x$, $y_\nu \rightarrow y$ und $f(x_\nu) \leq y_\nu$ für alle $\nu \in I$. Sei $\varepsilon > 0$. Dann existiert ein $\mu \in I$, so dass $y_\nu \leq y + \varepsilon$ für alle $\nu \geq \mu$ gilt. Somit konvergiert das Netz $(x_\nu)_{\nu \geq \mu}$ gegen x mit $f(x_\nu) \leq y + \varepsilon$ für alle $\nu \geq \mu$. Nach Voraussetzung also $f(x) \leq y + \varepsilon$. Es folgt $f(x) \leq y$, das heißt, $(x, y) \in \text{epi}(f)$ und $\text{epi}(f)$ ist abgeschlossen in $X \times \mathbb{R}$. $\square$

Ist X ein topologischer Vektorraum über $\mathbb{C}$, so ist (a) auch äquivalent zu:

(d) $\text{epi}(f)$, aufgefasst als Teilmenge von $X_{\mathbb{R}} \times \mathbb{R}$, ist abgeschlossen.

Beweis. Zum Beweis ersetze man im Beweis von der Äquivalenz von (a) und (c) das X durch $X_{\mathbb{R}}$. $\square$

Es sei angemerkt, dass, da f genau dann oberhalbstetig ist, wenn $-f$ unterhalbstetig ist, die in (b) genannte Ungleichung im Fall der Oberhalbstetigkeit $\limsup_{\nu} f(x_\nu) \leq f(x)$ lautet.

2.10 Präduale

2.10.1 Definition. Sei X ein Banachraum über $\mathbb{K}$. X heißt ein *Prädual* eines Banachraumes Y über $\mathbb{K}$, wenn der Dualraum X^* von X isometrisch isomorph zu Y ist, in Zeichen also, wenn $X^* \cong Y$ gilt. X heißt ein *eindeutiges Prädual*, wenn für jeden Banachraum Y über $\mathbb{K}$ gilt: $X^* \cong Y^*$ impliziert $X \cong Y$. X heißt ein *stark eindeutiges Prädual* (im Englischen *strongly unique predual*), wenn für jeden Banachraum Y über $\mathbb{K}$ jeder isometrische Isomorphismus von X^* nach Y^* schwach*-schwach*-stetig ist.

2.10.2 Satz von Dixmier-Ng (HOLMES [137](1975), Seite 211). *Sei X ein Banachraum über $\mathbb{K}$. Falls eine Hausdorff'sche, lokal konvexe Topologie τ auf X existiert, in der B_X τ -kompakt ist, dann ist $\{\ell \in X^\# : \ell \upharpoonright B_X \text{ ist } \tau\text{-stetig}\}$ ein abgeschlossener Unterraum von X^* und ein Prädual von X .*

2.10.3 Lemma (HOLMES [137](1975), Seite 212). *Sei Y ein Prädual eines Banachraumes X über $\mathbb{K}$, etwa vermittelt dem isometrischen Isomorphismus T von X auf Y^* . Dann ist*

der Unterraum $U := (T^* \circ i_Y)(Y)$ von X^* ein Prädual von X ; genauer ist die Abbildung $i_{X,U}$ ein isometrischer Isomorphismus von X auf U^* .

Beweis. Sei $T: X \rightarrow Y^*$ der nach Voraussetzung existierende isometrische Isomorphismus von X auf Y^* . Setze $U := (T^* \circ i_Y)(Y)$. Dann gilt für alle $x \in X$ und $y \in Y$:

$$\begin{aligned} \langle y, Tx \rangle_{(Y, Y^*)} &= \langle Tx, i_Y(y) \rangle_{(Y^*, Y^{**})} \\ &= \langle x, (T^* \circ i_Y)(y) \rangle_{(X, X^*)} \\ &= \langle x, (T^* \circ i_Y)(y) \rangle_{(X, U)} \\ &= \langle (T^* \circ i_Y)(y), i_{X,U}(x) \rangle_{(U, U^*)} \\ &= \langle U|(T^* \circ i_Y)(y), i_{X,U}(x) \rangle_{(U, U^*)} \\ &= \langle y, (U|(T^* \circ i_Y))^* \circ i_{X,U}(x) \rangle_{(Y, Y^*)}. \end{aligned}$$

Da, wie in 2.4.23 bemerkt wurde, das Dualsystem (Y, Y^*) in Y^* trennend ist, folgt für alle $x \in X$ die Gleichung $T(x) = (U|(T^* \circ i_Y))^* \circ i_{X,U}(x)$. Da T surjektiv ist, ist auch $(U|(T^* \circ i_Y))^*$ surjektiv. Da $U|(T^* \circ i_Y)$ nach Definition von U surjektiv ist, ist $(U|(T^* \circ i_Y))^*$ auch eine injektive Abbildung von U^* auf Y^* . Somit ist $i_{X,U} = (U|(T^* \circ i_Y))^{*-1} \circ T$. Da i_Y und mit T auch T^{-1} isometrisch sind, ist $i_{X,U}$ auch isometrisch. $\square$

2.10.4 Satz von Dixmier-Goldberg-Ruston (HOLMES [137](1975), Seite 213). *Sei X ein Banachraum über $\mathbb{K}$. Dann sind die folgenden Aussagen äquivalent:*

- (a) *Es existiert ein Prädual von X .*
- (b) *Es existiert ein die Punkte von X trennender Unterraum U von X^* , so dass $B_X \sigma(X, U)$ -kompakt ist.*
- (c) *X^* enthält einen Norm-bestimmenden, abgeschlossenen Unterraum U , wobei der Vektorraum U keinen echten, abgeschlossenen, die Punkte von X trennenden Unterraum enthält.*
- (d) *Es gibt eine Norm-1-Projektion von X^{**} auf $i_X(X)$ mit schwach*-abgeschlossenem Kern. (Folglich ist also $i_X(X)$ in X^{**} in der schwach*-Topologie Norm-1-komplementiert.)*

Beweis. (a) $\Rightarrow$ (b): Sei Y ein Prädual von X . Nach dem Lemma 2.10.3 gibt es einen Unterraum U von X^* , so dass $i_{X,U}$ ein isometrischer Isomorphismus von X auf U^* ist. Nach Bemerkung 2.5.16(d) trennt U die Punkte von X und $i_{X,U}$ ist ein Homöomorphismus von $(X, \sigma(X, U))$ auf $\text{ran}(i_{X,U})$ versehen mit der Relativtopologie von $\sigma(U^*, U)$. Somit gilt $i_{X,U}(\text{cl}(\sigma(X, U); B_X)) = \text{cl}(\sigma(U^*, U); i_{X,U}(B_X)) = (i_{X,U}(B_X))^{\bullet\bullet}$. Da $(i_{X,U}(B_X))^{\bullet\bullet}$ die kleinste, absolutkonvexe $\sigma(U^*, U)$ -abgeschlossene Menge in U^* ist, die $i_{X,U}(B_X) \subseteq B_{U^*}$ umfasst, ist sie eine Teilmenge von B_{U^*} . Da nach dem Satz von Alaoglu-Bourbaki B_{U^*} $\sigma(U^*, U)$ -kompakt ist, ist $\text{cl}(\sigma(X, U); B_X)$ $\sigma(X, U)$ -kompakt. Die restlichen Beweise a.a.O. $\square$

2.10.5 Bemerkungen (HOLMES [137](1975), Seite 213). Sei X ein Banachraum über $\mathbb{K}$. Ist dann zum Beispiel $V \subseteq X^*$ ein Prädual von X , so ist bezüglich des Dualsystems (X^*, X^{**}) V° ein die Aussage (d) des Satzes 2.10.4 erfüllender Unterraum von X^{**} . Ist umgekehrt W ein die Aussage (d) des Satzes 2.10.4 erfüllender Unterraum von X^{**} , so ist der bezüglich des Dualsystems (X^*, X^{**}) gebildete Unterraum W° von X^* ein Prädual von X .

2.10.6 Satz. *Sei X ein Banachraum über $\mathbb{K}$. Dann sind die folgenden Aussagen äquivalent:*

- (a) *X ist ein eindeutiges Prädual.*
- (b) *Jedes Prädual von X^* ist isometrisch isomorph zu X .*
- (c) *Jeder Banachraum Y über $\mathbb{K}$, der X als Prädual seines Dualraumes Y^* hat, ist isometrisch isomorph zu X .*

Beweis. Mnemotechnisch entspricht $*$ $\cong$ der Zeichenkette *ist Prädual von* und $\cong Y^*$ entspricht der Zeichenkette *mit Prädual Y* . $\square$

2.10.7 Satz. *Sei X ein Banachraum über $\mathbb{K}$. Dann sind die folgenden Aussagen äquivalent:*

- (a) *X ist ein stark eindeutiges Prädual.*
- (b) *Für jeden Banachraum Y über $\mathbb{K}$ und für jeden isometrischen Isomorphismus T von X^* nach Y^* existiert ein isometrischer Isomorphismus S von Y nach X , so dass T gleich der dualen Abbildung S^* von S ist.*
- (c) *Für jeden Banachraum Y über $\mathbb{K}$ und für jeden isometrischen Isomorphismus T von X^* nach Y^* ist $T^*(i_Y(Y))$ gleich $i_X(X)$.*
- (d) *Für jeden Banachraum Y über $\mathbb{K}$ und für jeden isometrischen Isomorphismus T von X^* nach Y^* ist die Abbildung $T^* \circ i_Y$ ein isometrischer Isomorphismus von Y nach $i_X(X)$.*
- (e) *X ist ein Prädual von einem Banachraum Y über $\mathbb{K}$ und die kanonischen Einbettungen von X und von jedem anderen Prädual von Y in Y^* stimmen überein.*
- (f) *Die kanonische Projektion π_{X^*} ist die einzige kontraktive Projektion von X^{***} auf X^* mit schwach*-abgeschlossenen Kern.*

Beweis. (a) $\Leftrightarrow$ (f): HARMAND, WERNER und WERNER [126](1993), Seiten 121 und 122.

(c) $\Rightarrow$ (b): GODEFROY [111](1978), Theorem 12.

(d): GODEFROY [112](1979), Seite 56, Definition. $\square$

2.10.8 Bemerkung. (HORN [140](1987)). Gemäß Satz 2.10.7 sagt man: Ein Banachraum X hat genau dann ein stark eindeutiges Prädual $Y \subseteq X^*$, wenn Y der einzige abgeschlossene Unterraum von X^* ist, welcher ein Prädual von X in der kanonischen Dualität ist. Besonders aufgefasst als diesen Unterraum von X^* , wird ein stark eindeutiges Prädual von X mit X_* bezeichnet. Nichtsdestotrotz ist es auch üblich, ein Prädual von X , das weder explizit als ein eindeutiges Prädual oder als ein stark eindeutiges Prädual von X noch als ein Unterraum von X^* vorausgesetzt wird, mit X_* zu bezeichnen. Manchmal ist auch praktisch, bei einem vorgegebenen Banachraum X , eine Menge zu definieren, diese Menge als X_* zu bezeichnen und dann im Nachhinein zu zeigen, dass diese besagte Menge ein Prädual von X ist. Siehe zum Beispiel 3.6.5, 3.7.36 und 3.7.40.

PALMER [232](2001), Seite 893 drückt es wie folgt aus: Ein Prädual Y von X ist genau dann ein stark eindeutiges Prädual, wenn es nur genau einen abgeschlossenen Unterraum von X^* gibt, der die Rolle von $i_Y(Y)$ übernehmen kann.

GODEFROY [111](1978) bezeichnet ein stark eindeutiges Prädual im Französischen als *un unique préduel normique* oder auch einfach nur als *un unique préduel*. Derselbe

Autor bezeichnet in GODEFROY [113](1983) und in GODEFROY [114](1984) mit *un unique préduel* ein „nur“ noch eindeutiges Prädual.

2.10.9 Korollar. *Sei X ein Banachraum über $\mathbb{K}$, der ein starkes Prädual X_* habe. Sei $T \in \mathcal{L}(X)$ ein isometrischer Isomorphismus. Dann ist $T = \sigma(X, X_*) - \sigma(X, X_*)$ -stetig.*

Beweis. Nach Satz 2.10.7(c) ist $T^*(i_{X_*}(X_*)) = i_{X_*}(X_*)$. Nach MEGGINSON [211](1998), Seite 290, Theorem 3.1.18, ist T^* ein isometrischer Isomorphismus. Somit ist $T^* \upharpoonright i_{X_*}(X_*)$ ein isometrischer Isomorphismus $i_{X_*}(X_*) \rightarrow i_{X_*}(X_*)$. Setze $S: X_* \rightarrow X_*$ als $S(x_*) := (i_{X_*})^{-1} \circ T^*(i_{X_*}(x_*))$. S ist also ein isometrischer Isomorphismus. Nun gilt für alle $x_* \in X_*$ und $x \in X$: $\langle S^*x, x_* \rangle = \langle x, Sx_* \rangle = \langle x, (i_{X_*})^{-1} \circ T^*\widehat{x_*} \rangle = \langle x, T^*\widehat{x_*} \rangle = \langle Tx, \widehat{x_*} \rangle = \langle Tx, x_* \rangle$. Da das Dualsystem (X, X_*) links trennend ist, folgt $S^*x = Tx$ für alle $x \in X$, das heißt, $T = S^*$. Nach MEGGINSON [211](1998), Seite 287, Theorem 3.1.11, ist somit $T = \sigma(X, X_*) - \sigma(X, X_*)$ -stetig. $\square$

2.10.10 Bemerkung (PALMER [232](2001), Seite 839; LI [198](2003), Seite 59). Schatten zeigte 1957: Sei H ein unendlichdimensionaler Hilbertraum über $\mathbb{K}$. Dann hat $K(H)$ kein Prädual. (Beweis per Satz von Alaoglu-Bourbaki und Satz von Krein-Milman.)

2.10.11. Siehe DUNFORD und SCHWARTZ [78](1958), Seite 458, Übung 4, für genau welche σ -finiten positiven Maßräume (S, Σ, μ) der Funktionenraum $L_1(S, \Sigma, \mu)$ ein Prädual besitzt. (Zum Beispiel *nicht* für $S = [0, 1]$, da $\text{ex}(L_1[0, 1]) = \emptyset$.)

2.10.12 (BOURGIN [36](1983), Seite 164, Theorem 5.7.6). Sei X ein Banachraum über $\mathbb{K}$, der die Radon-Nikodým-Eigenschaft aufweist. Dann ist X ein stark eindeutiges Prädual von X^* . Siehe auch 4.1.8.

2.11 Stützunktional, Kugelsystem, numerischer Bereich und Hermitezität

2.11.1 Definition. Sei X ein topologischer Vektorraum über $\mathbb{K}$. Sei A eine Teilmenge von X . Ein $x^* \in X^*$, $x^* \neq 0$, heißt ein *Stützunktional* (im Englischen *support functional*) für A , wenn ein $y \in A$ existiert mit $\text{Re } x^*y = \sup\{\text{Re } x^*x : x \in A\}$. In diesem Fall heißt y ein *Stützpunkt* von A und x^* *stützt* A bei y . (Im Fall von $\mathbb{K} = \mathbb{R}$ ist also ein $x^* \in X^{**}$ genau dann ein Stützunktional von A , wenn es auf A sein Supremum annimmt.) Ein $x^* \in X^*$, $x^* \neq 0$, heißt ein *Betragsstützunktional* (im Englischen *modulus support functional*) für A , wenn ein $y \in A$ existiert mit $|x^*y| = \sup\{|x^*x| : x \in A\}$. In diesem Fall heißt y ein *Betragsstützpunkt* von A .

Sei X ein normierter Vektorraum über $\mathbb{K}$. Ein $x^* \in X^*$ heißt ein (bei x) *Norm-annehmendes Funktional*, wenn es ein $x \in B_X$ mit $\|x^*\| = |x^*x|$ gibt. (Offensichtlich ist ein $x^* \in X^*$ genau dann Norm-annehmend, wenn es ein $x \in B_X$ mit $\|x^*\| = x^*(x)$ gibt.)

Seien X und Y zwei normierte Vektorräume über $\mathbb{K}$, die ein Dualsystem $(X, Y; B)$ bilden. Betrachte das Annihilatorsystem $(B_X, B_Y; \Xi; \mathbb{K})$ mit der durch

$$\Xi(x, y) := B(x, y) - 1 \quad \text{für alle } (x, y) \in B_X \times B_Y$$

definierten Systemabbildung. (Offensichtlich ist dieses System trennend.) Es wird im Folgenden als das *Kugelannihilatorsystem* oder einfach auch nur als das *Kugelsystem* des Dualsystems $(X, Y; B)$ bezeichnet. Falls von einem Kugelsystem eines normierten Raumes Z über $\mathbb{K}$ die Rede ist, so ist stets das Kugelsystem des Dualsystems (Z, Z^*) gemeint. Möchte man die Tatsache betonen, dass sich ein Kugelsystem auf die jeweiligen Einheitskugeln bezieht, kann man es genauer als ein *Einheitskugelannihilatorsystem* (im

Englischen also etwa *unit ball annihilator system*) bezeichnen. Anstelle des allgemein für Annihilatorsysteme üblichen Annihilatorsymbols $\perp$ ist für Kugelsysteme im Folgenden durchweg das Zeichen $\dashv$ als Annihilatorsymbol reserviert. (Bei Kugelsystemen eines normierten Vektorraumes über $\mathbb{K}$ wird an derselben Position und mit derselben Wirkung wie das hier verwendete Annihilatorsymbol $\dashv$ in der Literatur der Buchstabe f oder das Primzeichen $\prime$ verwendet.) Durch Vorsetzen des Wortes *Kugel* (oder bei Bedarf auch genauer des Wortes *Einheitskugel*) vor den Begriffen aus 2.1.4 kann man bequem deutlich machen, dass man sich auf ein Kugelsystem bezieht. So wird etwa der Annihilator in Kugelsystemen *Kugelannihilator* genannt. Oder: Für $x \dashv y$ sagt man, dass x und y *kugelorthogonal* (*zueinander*) seien. Anstelle von *kugelannihilatorabgeschlossen* sagt man einfacher *kugelabgeschlossen*.

Die *Stützfunktion* Φ_X (im Englischen auch *spherical image map* oder *duality map*) von S_X in die Potenzmenge der Einheitskugel von X^* ist definiert durch

$$\Phi_X(x) = x^\dashv \quad \text{für alle } x \in S_X.$$

Ist für ein $x \in S_X$ die Menge $\Phi_X(x)$ einelementig, so heißen X und die Norm von X *glatt* (im Englischen *smooth*) bei x und x heißt dann ein *glatter Punkt* (im Englischen *smooth point* oder auch *point of smoothness*) von B_X ; siehe auch 4.1.11. Falls $\Phi_X(x)$ für alle $x \in S_X$ einelementig ist, heißt X und seine Norm *glatt* und vereinbarungsgemäß soll dann Φ_X als eine Abbildung von S_X nach S_{X^*} aufgefasst werden.

Die Elemente von $\Phi_X(x)$ für ein $x \in S_X$ heißen die *x-unitalen Zustände* von X . Für fest gewähltes $x \in S_X$ wird die Menge $\Phi_X(x)$ auch *Zustandsraum* (im Englischen *state space*) von X genannt. Ist speziell A eine unital Algebra über $\mathbb{K}$, so heißen die Elemente von $\Phi_A(e)$ die *unitalen Zustände* von A .

Ein $x \in S_X$ heißt ein *Vertex* von B_X , falls $\Phi_X(x)$ die Punkte von X trennt, mit anderen Worten, falls das Dualsystem $(X, \text{span } \Phi_X(x))$ in X trennend ist.

2.11.2 Bemerkung. Sei X ein normierter Vektorraum über $\mathbb{K}$. Für $x \in S_X$ ist $\Phi_X(x)$ gleich dem Urbild von 1 unter $\hat{x} \upharpoonright S_{X^*}$, dem auf S_{X^*} eingeschränkten Auswertungsfunktional von x .

2.11.3 Bemerkung (MEGGINSON [211] (1998), Seiten 271-273). Sei X ein normierter Raum und $x^* \in X^*$. Dann ist x^* genau dann ein Stützunktional für B_X , wenn x^* ein von Null verschiedenes Norm-annehmendes Funktional ist. Für $x \in S_X$ sind die folgenden Aussagen äquivalent:

- (a) $x \dashv x^*$.
- (b) $x^* \in \Phi_X(x)$.
- (c) $x^* \in S_{X^*}$ ist ein Stützunktional für B_X bei x .
- (d) $x^* \in S_{X^*}$ ist ein Stützunktional für B_X bei x mit $\text{Im } x^*(x) = 0$.
- (e) $x^* \in S_{X^*}$ ist ein bei x Norm-annehmendes Funktional mit $|x^*(x)| = x^*(x)$.

2.11.4 Subdifferential (PHELPS [238] (1993); GÖPFERT, RIEDRICH und TAMMER [119] (2009), Seite 127). Sei X ein topologischer Vektorraum über $\mathbb{K}$. Sei A eine Teilmenge von X , $x_0 \in A$ und $f: A \rightarrow \mathbb{R}$ eine Funktion. Dann heißt die Menge

$$\partial f(x_0) := \left\{ x^* \in X^* : \text{Re } x^*(x - x_0) \leq f(x) - f(x_0) \text{ für alle } x \in A \right\}$$

das *Subdifferential* von f bei x_0 .

Sei X ein normierter Vektorraum über $\mathbb{K}$. Dann gilt für alle $x \in X^\times$ die Gleichung

$$\partial\|\cdot\|(x) = \left\{x^* \in S_{X^*} : \operatorname{Re} x^*(x) = \|x\|\right\}.$$

Des Weiteren gilt $\partial\|\cdot\|(0) = B_{X^*}$.

Beweis. Sei $x \in X^\times$ und $x^* \in S_{X^*}$ mit $\operatorname{Re} x^*(x) = \|x\|$. Dann gilt für alle $y \in X$: $\operatorname{Re} x^*(y - x) \leq \|x^*\| \|y - x\| = \|y - x\|$, das heißt, $x^* \in \partial\|\cdot\|(x)$.

Sei nun umgekehrt $x^* \in \partial\|\cdot\|(x)$. Also gilt für alle $y \in X$: $\operatorname{Re} x^*(y - x) \leq \|y - x\|$. Insbesondere gilt diese Ungleichung für $y = 0$ und $y = 2x$, das heißt, man hat sowohl $-\operatorname{Re} x^*(x) \leq -\|x\|$, also $\|x\| \leq \operatorname{Re} x^*(x)$, als auch $\operatorname{Re} x^*(x) \leq \|x\|$. Somit $\operatorname{Re} x^*(x) = \|x\|$. Dies eingesetzt gibt $\operatorname{Re} x^*(y) - \|x\| \leq \|y - x\|$, also $\operatorname{Re} x^*(y) \leq \|y\|$ für alle $y \in X$, heißt, $\|\operatorname{Re} x^*\| = 1$, wegen 2.9.9 also $\|x^*\| = 1$.

$\partial\|\cdot\|(0) = \{x^* \in X^* : \operatorname{Re} x^*(x) \leq \|x\| \text{ für alle } x \in X\}$. $\operatorname{Re} x^*(x) \leq \|x\|$ für alle $x \in X$ heißt insbesondere auch $\operatorname{Re} x^*(-x) \leq \| -x \|$ für alle $x \in X$, das heißt, $-\|x\| \leq \operatorname{Re} x^*(x)$ für alle $x \in X$. Somit $\partial\|\cdot\|(0) = \{x^* \in X^* : \|\operatorname{Re} x^*\| \leq 1\} = B_{X^*}$, wobei das letzte Gleichheitszeichen wegen 2.9.9 gilt. $\square$

Insbesondere hat man also

$$\partial\|\cdot\| \upharpoonright S_X = \Phi_X.$$

Siehe auch PHELPS [238](1993), Seite 27, Beispiel 2.26, und SIMONS [261](2008), Seite 49, für einige Eigenschaften von $\frac{1}{2}\|\cdot\|^2$.

2.11.5 Bishop-Phelps (MEGGINSON [211](1998), Seiten 275, 278). Sei X ein Banachraum über $\mathbb{K}$ und A eine abgeschlossene, konvexe Teilmenge von X . Dann gilt:

(a) Die Menge der Stützpunkte von A ist dicht im Rand ∂A von A (Stützpunkt-Theorem von Bishop-Phelps).

(b) Falls zusätzlich sowohl A beschränkt und nicht leer ist als auch $X \neq \{0\}$ gilt, ist die Menge der Stützfunktionale für A dicht in X^* (Stützfunktional-Theorem von Bishop-Phelps oder meist schlichtweg *der Satz von Bishop-Phelps*).

2.11.6 James' Theorem (MEGGINSON [211](1998), Seite 121). Sei X ein normierter Vektorraum über $\mathbb{K}$. Als eine einfache Konsequenz des Satzes von Hahn-Banach gilt: X reflexiv $\Rightarrow$ Alle $x^* \in X^*$ sind Norm-annehmend. JAMES hat 1964 gezeigt, dass, falls X ein Banachraum über $\mathbb{K}$ ist, auch die Umkehrung gilt.

2.11.7 Kugelannihilatoren. Sei X ein normierter Vektorraum über $\mathbb{K}$. Für jedes $x \in S_X$ ist die Menge $\Phi_X(x)$ eine nicht leere, konvexe und schwach*-kompakte Seite von B_{X^*} .

Beweis. Sei $x \in S_X$. Nach dem Satz von Hahn-Banach ist $\Phi_X(x)$ nicht leer. Die Konvexität ist klar. Zur schwach*-Kompaktheit sei $(x_\nu^*) \subseteq \Phi_X(x)$ ein Netz, das gegen ein $x^* \in X^*$ schwach*-konvergiert. Das ist äquivalent dazu, dass für jedes $y \in X$ das Netz $(x_\nu^*(y))$ gegen $x^*(y)$ konvergiert; insbesondere ist $x^*(x) = 1$, also $\|x^*\| \geq 1$. Angenommen, $\|x^*\| > 1$. Dann gäbe es ein $z \in S_X$ mit $x^*(z) > 1$, im Widerspruch zu $x_\nu^*(z) \leq 1$ für alle ν . Also $x^* \in \Phi_X(x)$ und $\Phi_X(x)$ ist schwach*-abgeschlossen. Nach dem Satz von Alaoglu-Bourbaki ist $\Phi_X(x)$ schwach*-kompakt. Bleibt zu zeigen, dass eine Seite vorliegt. Seien dazu $a, b \in B_{X^*}$ und $\lambda \in \mathbb{R}, 0 < \lambda < 1$. Falls dann $\lambda a + (1 - \lambda)b \in \Phi_X(x)$ ist, insbesondere also $\lambda a(x) + (1 - \lambda)b(x) = 1$ gilt, muss wegen $|a(x)| \leq 1$ und $|b(x)| \leq 1$ und da 1 ein Extrempunkt von $S_{\mathbb{K}}$ ist, $a(x) = b(x) = 1$ gelten, das heißt, $a, b \in \Phi_X(x)$. Somit ist $\Phi_X(x)$ eine Seite von B_{X^*} . $\square$

Folglich ist wegen 1.1.19(a) für jede nicht leere Teilmenge $M \subseteq S_X$ die Menge $M^\perp$ eine konvexe, schwach*-kompakte Seite von B_{X^*} . Insbesondere stimmen gemäß 2.4.40(b) die $(n - w^*)$ - und die $(n - \overline{w^*})$ -Halbstetigkeiten für Φ_X überein.

Wie eben zeigt man: Für jede nicht leere Teilmenge $N \subseteq S_{X^*}$ ist die Menge $N_\perp$ eine konvexe, abgeschlossene Seite von B_X . Es sei an James' Theorem (siehe 2.11.6) erinnert. Siehe auch 4.7.71.

2.11.8 Semi-exponierte Seiten (EDWARDS und RÜTTIMANN [86](2001)). Sei (X, τ) ein lokal konvexer Hausdorffraum über $\mathbb{K}$. Sei C eine abgeschlossene, konvexe Teilmenge von X . Die Menge aller abgeschlossenen Seiten von C wird mit $\text{Faces}(\tau; C)$ bezeichnet. Für jede nicht leere Familie von Elementen aus $\text{Faces}(\tau; C)$ ist deren Durchschnitt wieder ein Element aus $\text{Faces}(\tau; C)$.

Auf $\text{Faces}(\tau; C)$ ist die Mengeninklusionsrelation eine partielle Ordnung, und mit dieser versehen, ist $\text{Faces}(\tau; C)$ ein vollständiger Verband. $\text{Faces}(\tau; C)$ wird auch stets als eine punktierte Menge mit der leeren Menge als Basispunkt aufgefasst.

Eine Teilmenge E von C heißt eine τ -exponierte Seite (im Englischen τ -exposed face), wenn ein $x^* \in X^*$ und ein $t \in \mathbb{R}$ existiert, so dass gilt: $\text{Re } x^*(y) = t$ für alle $y \in E$ und $\text{Re } x^*(y) < t$ für alle $y \in C \setminus E$. Man bemerke, dass die leere Menge und C τ -exponierte Seiten von C sind. Die Menge aller τ -exponierten Seiten von C wird mit $\text{Expfaces}(\tau; C)$ bezeichnet. Es gilt

$$\text{Expfaces}(\tau; C) \subseteq \text{Faces}(\tau; C).$$

Beweis. Sei $E \subseteq C$ eine nicht leere Teilmenge von C . Existiere ein $x^* \in X^*$ und ein $t \in \mathbb{R}$ mit (1.) $\text{Re } x^*(y) = t$ für alle $y \in E$ und (2.) $\text{Re } x^*(y) < t$ für alle $y \in C \setminus E$. Dann ist zu zeigen, dass E eine abgeschlossene Seite von C ist. (a) Es ist $E = \{y \in C : \text{Re } x^*(y) = t\} = (\text{Re } x^* \upharpoonright C)^{-1}(t)$ und E ist abgeschlossen. (b) (i) Seien $a, b \in C$, sei $\lambda \in \mathbb{R}$, $0 < \lambda < 1$ und $\lambda a + (1 - \lambda)b \in E$. $\Rightarrow \text{Re } x^*(\lambda a + (1 - \lambda)b) = t \Rightarrow \lambda \text{Re } x^*(a) + (1 - \lambda) \text{Re } x^*(b) = t$. Nach Voraussetzung ist $\text{Re } x^*(a) \leq t$ und $\text{Re } x^*(b) \leq t$. Wäre $\text{Re } x^*(a) < t$ oder $\text{Re } x^*(b) < t$, so würde wegen $\lambda, (1 - \lambda) > 0$ gelten: $\lambda \text{Re } x^*(a) + (1 - \lambda) \text{Re } x^*(b) < \lambda t + (1 - \lambda)t = t$. Also $\text{Re } x^*(a) = \text{Re } x^*(b) = t$, $a, b \in E$ und E ist eine extremale Teilmenge von C . (ii) Seien $a, b \in E$. Sei $\lambda \in \mathbb{R}$, $0 < \lambda < 1$. Betrachte $x := \lambda a + (1 - \lambda)b$. $\Rightarrow \text{Re } x^*(x) = \lambda \text{Re } x^*(a) + (1 - \lambda) \text{Re } x^*(b) = \lambda t + (1 - \lambda)t = t$ und E ist konvex. $\square$

Jeder Durchschnitt von einer endlichen, aber nicht leeren Familie von Elementen aus $\text{Expfaces}(\tau; C)$ ist wieder ein Element von $\text{Expfaces}(\tau; C)$. Der Durchschnitt von einer beliebigen, aber nicht leeren Familie von Elementen von $\text{Expfaces}(\tau; C)$ heißt eine τ -semi-exponierte Seite (im Englischen τ -semi-exposed face) von C . Die Menge aller τ -semi-exponierten Seiten von C wird mit $\text{Semiexpfaces}(\tau; C)$ bezeichnet. Es gilt

$$\text{Expfaces}(\tau; C) \subseteq \text{Semiexpfaces}(\tau; C) \subseteq \text{Faces}(\tau; C).$$

Für jede beliebige, nicht leere Familie von Elementen aus $\text{Semiexpfaces}(\tau; C)$ ist ihr Durchschnitt wieder ein Element von $\text{Semiexpfaces}(\tau; C)$. Folglich ist die Menge $\text{Semiexpfaces}(\tau; C)$, versehen mit der durch die Mengeninklusion induzierten partiellen Ordnung, ein vollständiger Verband. Das Infimum einer Familie von Elementen aus $\text{Semiexpfaces}(\tau; C)$ stimmt überein mit dem von ihr in $\text{Faces}(\tau; C)$ gebildeten Infimum.

Sei nun X ein Banachraum über $\mathbb{K}$. Betrachte das Kugelsystem von X . Dann sind genau die kugelabgeschlossenen Teilmengen von B_X die Elemente von $\text{Semiexpfaces}(\|\cdot\|; B_X)$ und genau die kugelabgeschlossenen Teilmengen von B_{X^*} die Elemente von $\text{Semiexpfaces}(w^*; B_{X^*})$. Die Abbildungen $F \mapsto F^\perp$ und $F \mapsto F_\perp$ zwischen $\text{Semiexpfaces}(\|\cdot\|; B_X)$ und $\text{Semiexpfaces}(w^*; B_{X^*})$ sind Antiordnungsisomorphismen und jeweils die Inverse der anderen. Die Bildung des Kugelrechtsannihilators wird in der englischen Literatur die *facear operation*, die des Kugellinkssannihilators die *pre-facear operation* genannt.

2.11.9 Bemerkung. Sei X ein normierter Vektorraum über $\mathbb{K}$. Betrachte das Kugelsystem von X . Es sei an die Definition 2.5.1 der Absolutpolaren erinnert. Für jede Teilmenge $A \subseteq B_X$ gilt die Inklusion

$$A^\perp \subseteq A^\bullet \cap S_{X^*}.$$

2.11.10 (PALMER [231](1994), Seite 261). Sei A eine unitale Algebra über $\mathbb{K}$. Bezeichnet $\bar{A}$ die Vervollständigung von A , dann kann man $\Phi_A(e)$ kanonisch mit $\Phi_{\bar{A}}(e)$ identifizieren.

2.11.11 Exponierte Punkte (GILES [109](1982), Seiten 190, 195, 212; FONF et al. in [158](2001), Seiten 628, 635; DAY [61](1973), Seite 105; PHELPS [238] (1993), Seite 91).

Sei X ein topologischer Vektorraum über $\mathbb{K}$ und A eine Teilmenge von X . Ein $a \in A$ heißt ein *exponierter Punkt* (im Englischen *exposed point* und ehemals auch mal *accessible point*) von A , wenn ein $x^* \in X^*$ mit

$$\operatorname{Re} x^*(a) = \sup \operatorname{Re} x^*(A) > \operatorname{Re} x^*(y) \quad \text{für alle } y \in A \setminus \{a\}$$

existiert. In diesem Fall sagt man, x^* *exponiere den Punkt* $a \in A$ oder auch, $a \in A$ *werde mittels* x^* *exponiert* (im Englischen *exposed by* x^*). Die Menge aller exponierten Punkte von A wird mit $\operatorname{expo}(A)$ bezeichnet. Es gilt stets $\operatorname{expo}(A) \subseteq \operatorname{ex}(A)$. Die umgekehrte Inklusion gilt aber im Allgemeinen nicht, wie das folgende Beispiel zeigt: Betrachte $X := \mathbb{R} \oplus_p \mathbb{R}$ für ein p mit $1 \leq p < \infty$ und $A := B_X \cup ([-1, 1] \times [-1, 0])$. Dann sind zwar die Punkte $(-1, 0)$ und $(1, 0)$ beide Extrempunkte von A , aber keiner von ihnen ist ein exponierter Punkt von A .

Sei X ein normierter Vektorraum über $\mathbb{K}$ und $x \in S_X$ ein glatter Punkt von B_X . Bezeichne x^* das eindeutig bestimmte Element von $x^\perp$. Dann ist $\operatorname{Re} \hat{x}$, der Realteil des Auswertungsfunktional $\hat{x}$, eine Funktion auf der w^* -kompakten Menge B_{X^*} , die nur an dem Punkt x^* ihr Maximum annimmt. Das heißt, x^* ist ein durch $\operatorname{Re} \hat{x}$ exponierter Punkt von B_{X^*} . Salopp ausgedrückt: Jeder glatte Punkt x von B_X exponiert sein Stützfunktional $x^\perp$.

Der Satz von STRASZEWICZ-KLEE lautet: Ist X ein normierter Vektorraum über $\mathbb{K}$ und A eine kompakte, konvexe Teilmenge von X , so gilt

$$A = \operatorname{cl} \operatorname{co} \operatorname{expo} A.$$

Im Gegensatz zum Satz von Krein-Milman, der die Gleichung

$$A = \operatorname{cl} \operatorname{co} \operatorname{ex} A \tag{2.36}$$

allgemeiner für Hausdorff'sche topologische Vektorräume X über $\mathbb{K}$ behauptet, deren Dualsystem (X, X^*) in X trennend ist (und dabei ebenfalls A als kompakt und konvex voraussetzt), gilt der Satz von Straszewicz-Klee für solche X im Allgemeinen nicht.

(BOURGIN [36](1983), Seite 40). Sei C eine abgeschlossene, konvexe Teilmenge von X . Dann sagt man, C habe die *Krein-Milman-Eigenschaft* (im Englischen *Krein-Milman property*; abgekürzt *KMP*), wenn für jede abgeschlossene, beschränkte, konvexe Teilmenge A von C die Gleichung (2.36) gilt. C besitzt genau dann die Krein-Milman-Eigenschaft, wenn jede abgeschlossene, beschränkte, konvexe Teilmenge von C einen Extrempunkt besitzt. Falls C die Radon-Nikodým-Eigenschaft aufweist, so besitzt C auch die Krein-Milman-Eigenschaft (ebd., Seite 49, Theorem 3.3.6 von Lindenstrauss). Insbesondere besitzt also jeder Banachraum über $\mathbb{K}$, der die Radon-Nikodým-Eigenschaft aufweist, die Krein-Milman-Eigenschaft, und für Banachräume über $\mathbb{K}$, die ein Prädual haben, gilt auch die Umkehrung, siehe ebd., Seite 111, Theorem 4.4.1, und GILES [109](1982), Seite 211.

Nach BOURGIN, [36](1983), Seite 40, Satz 3.1.3, besitzt jeder reflexive Banachraum über $\mathbb{K}$ die Krein-Milman-Eigenschaft.

Sei nun X ein normierter Vektorraum über $\mathbb{K}$ und A eine Teilmenge von X . Ein $a \in A$ heißt *stark exponierter Punkt von A* , wenn ein $x^* \in X^*$ existiert, welches a exponiert, und für jede Folge $(a_n)_{n \in \mathbb{N}}$ in A mit $\operatorname{Re} x^*(a_n) \rightarrow \operatorname{Re} x^*(a)$, $n \rightarrow \infty$, folgt $\|a_n - a\| \rightarrow 0$, $n \rightarrow \infty$. Ein $x^* \in X^*$ stark exponiert genau dann ein $a \in A$, wenn es a exponiert und die Durchmesser der Scheiben $S(A, x^*, \alpha)$ mit $\alpha \rightarrow 0$, $\alpha > 0$, gegen Null konvergieren — also genau dann, wenn es a exponiert und die Scheiben $S(A, x^*, \alpha)$, $\alpha > 0$, eine Umgebungsbasis von a in A darstellen. Eine weitere dazu äquivalente Aussage ist, dass x^* das a exponiert und gilt:

$$\forall \varepsilon > 0 \exists \delta > 0 : (x \in A \text{ und } \|x - a\| \geq \delta) \Rightarrow (\operatorname{Re} x^*(x) \leq \operatorname{Re} x^*(a) - \varepsilon).$$

Die Menge der stark exponierten Punkte von A wird mit $\operatorname{str\,expo}(A)$ bezeichnet. Ist X ein Banachraum über $\mathbb{K}$ und A eine abgeschlossene, beschränkte und konvexe Teilmenge von X , die die Radon-Nikodým-Eigenschaft aufweist, so gilt

$$A = \operatorname{cl\,co\,str\,expo}\,A.$$

Ist speziell $x \in S_X$, so ist x genau dann ein stark exponierter Punkt von B_X , wenn ein $x^* \in \Phi_X(x)$ existiert, so dass für jede Folge $(x_n)_{n \in \mathbb{N}}$ in B_X mit $\lim_{n \rightarrow \infty} \operatorname{Re} x^*(x_n) = 1$ die Gleichung $\lim_{n \rightarrow \infty} \|x_n - x\| = 0$ folgt. Es sei an die in 2.6.14 angemerkte Bezeichnungsweise von BOURGIN [36](1983) erinnert.

Sei X ein Banachraum über $\mathbb{K}$ und K eine nicht leere, abgeschlossene, beschränkte, konvexe Teilmenge von X . Für jeden Punkt von K gilt dann das folgende Implikationschema:

$$\begin{array}{ccccc} \text{stark exponierter Punkt} & \Rightarrow & \text{stark extremer Punkt} & \Leftrightarrow & \text{Beulpunkt} \\ \downarrow & & \downarrow & & \\ \text{Vertex} & \Rightarrow & \text{exponierter Punkt} & \Rightarrow & \text{Extrempunkt} \\ & & \downarrow & & \\ & & \text{Stützpunkt} & & \end{array}$$

2.11.12 Minimale Elemente von $\operatorname{Faces}(\tau; C)$. Sei X ein Banachraum über $\mathbb{K}$ mit $X \neq \{0\}$. Sei C eine beschränkte, abgeschlossene, konvexe, nicht leere Teilmenge von X . Betrachte den Verband $\operatorname{Faces}(\|\cdot\|; C)$. Sei F ein minimales Element von $\operatorname{Faces}(\|\cdot\|; C)^\times$. Dann ist F einelementig.

Beweis. Angenommen, F sei nicht einelementig. F besitzt dann also zwei verschiedene Elemente $x, y \in F$. Nach dem Satz von Hahn-Banach existiert ein $x^* \in X^*$ mit $\operatorname{Re} x^*(x) < \operatorname{Re} x^*(y)$. Setze $\delta := \operatorname{Re} x^*(y) - \operatorname{Re} x^*(x)$ und t als die reelle Zahl $\sup \operatorname{Re} y^*(F)$. Nach dem Satz von Bishop-Phelps existiert ein $y^* \in X^*$ mit $\|y^* - x^*\| < \frac{\delta}{2 \cdot \max\{\|x\|, \|y\|\}}$, so dass für ein $z \in F$ die Gleichung $\operatorname{Re} y^*(z) = t$ erfüllt ist. Es gilt $\operatorname{Re} y^*(x) - \operatorname{Re} x^*(x) = \operatorname{Re} (y^* - x^*)(x) \leq |(y^* - x^*)(x)| \leq \|y^* - x^*\| \cdot \|x\| < \delta/2$ und analog erhält man $\operatorname{Re} x^*(y) - \operatorname{Re} y^*(y) < \delta/2$. Folglich hat man $\operatorname{Re} y^*(y) - \operatorname{Re} y^*(x) = \delta - (\operatorname{Re} x^*(y) - \operatorname{Re} y^*(y)) - (\operatorname{Re} y^*(x) - \operatorname{Re} x^*(x)) > 0$, also $\operatorname{Re} y^*(x) < \operatorname{Re} y^*(y)$. Setze $E := \{w \in F : \operatorname{Re} y^*(w) = t\}$. Da nun $w \in F \setminus E$ äquivalent zu $\operatorname{Re} y^*(w) < t$ ist, ist E eine Norm-exponierte Seite von F , insbesondere also eine abgeschlossene Seite von F und damit auch eine abgeschlossene Seite von C , die aber wegen $x \in F \setminus E$ von F verschieden ist. Das heißt, $\emptyset < E < F$, im Widerspruch zur Minimalität von F . $\square$

Somit bilden genau die Extrempunkte von C die minimalen Elemente von $\operatorname{Faces}(\|\cdot\|; C)^\times$.

Sei nun X ein Banachraum über $\mathbb{K}$ und betrachte den Verband $\operatorname{Faces}(w^*; B_{X^*})$. Da B_{X^*} nach dem Satz von Alaoglu-Bourbaki w^* -kompakt ist,

enthält nach dem Satz von Krein-Milman jedes Element von $\text{Faces}(w^*; B_{X^*})^\times$ einen Extrempunkt von B_{X^*} . Damit bilden genau die Extrempunkte von B_{X^*} die minimalen Elemente von $\text{Faces}(w^*; B_{X^*})^\times$.

2.11.13 Schwach Hahn-Banach-glatt (SMITH und SULLIVAN [263](1977), Seite 135; GILES, GREGORY und SIMS [108](1978), Seite 105). Sei X ein normierter Vektorraum über $\mathbb{K}$ und $x \in S_X$. Es sei an die in 2.5.16(b) erwähnte Zerlegung (2.31) von X^{***} erinnert. Man betrachte die Kugelsysteme der Dualsysteme (X, X^*) , (X^*, X^{**}) und (X^{**}, X^{***}) . Offensichtlich gilt $\hat{x}^\perp \cap \hat{X}^\perp = \emptyset$. Es gilt die Gleichung

$$i_{X^*}(x^\perp) = \hat{x}^\perp \cap i_{X^*}(X^*). \quad (2.37)$$

Beweis. $i_{X^*}(x^\perp) = \{x^{***} \in S_{X^{***}} : x^{***} = \widehat{x^*} \text{ für ein } x^* \in S_{X^*} \text{ mit } x^*(x) = 1\} = \{x^{***} \in S_{X^{***}} : x^{***} = \widehat{x^*} \text{ für ein } x^* \in S_{X^*} \text{ mit } \hat{x}(x^*) = 1\} = \{x^{***} \in S_{X^{***}} : x^{***} = \widehat{x^*} \text{ für ein } x^* \in S_{X^*} \text{ mit } x^{***}(\hat{x}) = 1\} = \{x^{***} \in S_{i_{X^*}(X^*)} : x^{***}(\hat{x}) = 1\} = \hat{x}^\perp \cap i_{X^*}(S_{X^*})$. $\square$

X heißt *schwach Hahn-Banach-glatt* bei x , wenn $i_{X^*}(x^\perp) = \hat{x}^\perp$ gilt. Ist X bei allen $x \in S_X$ schwach Hahn-Banach-glatt, so heißt X *schwach Hahn-Banach-glatt*. Offensichtlich ist X genau dann schwach Hahn-Banach-glatt bei $x \in S_X$, wenn $\hat{x}^\perp \subseteq \hat{x}^\perp \cap i_{X^*}(S_{X^*})$ gilt, das heißt, genau dann, wenn gilt: Ist $x^{***} \in S_{X^{***}}$ und $x^{***} = x^* + y$ mit den nach der genannten Zerlegung eindeutig bestimmten Elementen $x^* \in i_{X^*}(X^*)$, $y \in i_{X^*}(X)^\perp$ und dann $x^*(\hat{x}) = 1$, $\|x\| = \|x^*\| = 1$, so gilt $y = 0$. Mit anderen Worten ist X genau dann schwach Hahn-Banach-glatt, wenn für jedes auf S_X Norm-annehmende Funktional $x^* \in S_{X^*}$ das Funktional $i_{X^*}(x^*)$ die einzige — mithin eindeutige — Hahn-Banach-Fortsetzung von x^* in X^{**} ist.

Die folgenden drei Aussagen sind für jedes $x \in S_X$, X ein Banachraum über $\mathbb{K}$, äquivalent:

- (a) X ist bei x schwach Hahn-Banach-glatt.
- (b) Φ_X ist $(n - \overline{w})$ -oberhalbstetig bei x und $\Phi_X(x)$ ist schwach-kompakt.
- (c) Die schwache und die schwach*-Topologie stimmen auf der Einheitssphäre des Dualraumes auf Punkten von $\Phi_X(x)$ überein.

2.11.14 Bemerkung (RUDIN [251](1973), Seite 93). Seien X und Y zwei normierte Vektorräume über $\mathbb{K}$. Zu jedem $T \in \mathcal{L}(X, Y)$ gibt es genau ein $S \in \mathcal{L}(Y^*, X^*)$ mit $\langle Tx, y^* \rangle = \langle x, Sy^* \rangle$ für alle $x \in X$, $y^* \in Y^*$. Mit der Bezeichnung von 2.4.9 ist dieses S genau die duale Abbildung T^* von T .

2.11.15 Definition (BONSALL und DUNCAN [29](1971), Seiten 15, 81 und 85).

- (a) Sei X ein normierter Vektorraum über $\mathbb{K}$ mit einem fest gewählten Element aus S_X , welches mit 1 bezeichnet sei. Dann heißt für $x \in X$ die Menge

$$V(1; x) := V_X(1; x) := \{x^*(x) : 1 \dashv x^*\} \quad (\equiv 1^\perp(x) = \Phi_X(1)(x))$$

der *numerische Wertebereich* von x bezüglich 1 .

- (b) Sei X ein normierter Vektorraum über $\mathbb{K}$. Dann setze:

$$\Pi := \Pi_X := \{(x, x^*) \in S_X \times S_{X^*} : x \dashv x^*\}.$$

- (c) Sei X ein normierter Vektorraum über $\mathbb{K}$ und $T \in \mathcal{L}(X)$. Dann heißt die Menge

$$V_{\text{spatial}}(T) := \{\langle Tx, x^* \rangle \in \mathbb{K} : (x, x^*) \in \Pi_X\},$$

wobei $\langle \cdot, \cdot \rangle$ die kanonische Bilinearform von X ist, der *räumliche numerische Wertebereich* (im Englischen *spatial numerical range*) von T .

- (d) Sei A eine unitale Algebra über $\mathbb{K}$ und $a \in A$. Mit L_a als die durch a bestimmte Links-Multiplikation in A , also $L_a \in \mathcal{L}(A)$, heißt die Menge

$$V(a) := V_A(a) := V_{\text{spatial}}(L_a)$$

der *numerische Wertebereich* von a (oder auch genauer: der *algebraische numerische Wertebereich* von a). Es ist also (c) mit X als A und T als L_a anzuwenden und man hat somit

$$V(a) = \{ \langle ab, b^* \rangle \in \mathbb{K} : (b, b^*) \in \Pi_A \},$$

wobei $\langle \cdot, \cdot \rangle$ die kanonische Bilinearform von A ist. Die Zahl

$$v(a) := \sup \{ |\lambda| : \lambda \in V(a) \}$$

heißt der *numerische Radius* von a .

- (e) Sei X ein normierter Vektorraum über $\mathbb{K}$. Ein *LUMER-Produkt* auf X ist eine Abbildung $[\cdot, \cdot] : X \times X \rightarrow \mathbb{K}$ mit

- (i) Für jedes feste $y \in X$ ist $x \mapsto [x, y]$ eine lineare Abbildung auf X .
- (ii) $[x, x] > 0$, wenn $x \neq 0$.
- (iii) $|[x, y]|^2 \leq [x, x][y, y]$ für alle $x, y \in X$.

(BONSALL und DUNCAN, ebd., nennen solch ein Produkt ein Semiskalarprodukt.)

- (f) Sei X ein normierter Vektorraum über $\mathbb{K}$ mit einem LUMER-Produkt $[\cdot, \cdot]$ und $T \in \mathcal{L}(X)$. Dann heißt die Menge

$$W(T) := W_{[\cdot, \cdot]}(T) := \{ [Tx, x] \in \mathbb{K} : [x, x] = 1 \}$$

der *numerische Wertebereich* von T bezüglich $[\cdot, \cdot]$.

2.11.16. Die in Definition 2.11.15 verwendete Idee, anstelle einer unitalen Algebra allgemeiner ein Paar $(X, \mathbf{1})$, bestehend aus einem normierten Vektorraum X und einem Norm-1-Element $\mathbf{1}$ aus X , zu verwenden, findet sich weiter ausgeführt bei MARTÍNEZ MORENO, MENA JURADO, PAYÁ ALBERT und RODRÍGUEZ PALACIOS [207](1985).

2.11.17 Satz (BONSALL und DUNCAN [29](1971), Seiten 15 und 16). *Sei A eine unitale Algebra über $\mathbb{K}$. Dann gilt:*

- (a) $V(a)$ ist für jedes $a \in A$ eine kompakte, konvexe und nicht leere Teilmenge von $\mathbb{K}$.
- (b) $V(a + b) \subseteq V(a) + V(b)$ für alle $a, b \in A$.
- (c) $V(e) = \{1\}$.
- (d) $V(\lambda + \mu a) = \lambda + \mu \cdot V(a)$ für alle $\lambda, \mu \in \mathbb{K}, a \in A$.
- (e) $V(a) = V(e; a)$ für alle $a \in A$. Konkret heißt das:

$$V(a) = e^\perp(a) \quad \left(\equiv \{ \ell(a) : e \dashv \ell \} = \{ \ell(a) : \ell \in \Phi_A(e) \} \right)$$

für alle $a \in A$.

2.11.18. Sei X ein normierter Vektorraum über $\mathbb{K}$ und $T \in \mathcal{L}(X)$. Für jedes $(x, x^*) \in \Pi_X$ setze $\ell_{x, x^*} : \mathcal{L}(X) \rightarrow \mathbb{K}, R \mapsto \langle Rx, x^* \rangle$. Für jedes $(x, x^*) \in \Pi_X$ ist dann offenbar $\ell_{x, x^*} \in \Phi_{\mathcal{L}(X)}(\text{Id}_X)$, also $\ell_{x, x^*}(T) \in V(\text{Id}_X; T)$ und nach Satz 2.11.17(e) somit $\ell_{x, x^*}(T) \in V(T)$. Das heißt, es gilt $V_{\text{spatial}} \subseteq V(T)$. Genauer gilt:

2.11.19 Satz (BONSALL und DUNCAN [29](1971), Seite 84). *Sei X ein normierter Vektorraum über $\mathbb{K}$ und $T \in \mathcal{L}(X)$. Dann gilt: $cl(co V_{\text{spatial}}(T)) = V(T)$.*

2.11.20 Bemerkung und Definition (LUMER [203](1961), Seite 31). *Sei X ein Vektorraum über $\mathbb{K}$. Ist $[\cdot, \cdot]$ ein LUMER-Produkt auf X , dann ist $x \mapsto [x, x]^{1/2}$ eine Norm auf X .*

Trägt X eine Norm $\|\cdot\|$, dann heißt ein LUMER-Produkt $[\cdot, \cdot]$ die Norm bestimmend, wenn $\|x\| = [x, x]^{1/2}$ für alle $x \in X$ gilt.

Zu jeder Norm auf X gibt es ein die Norm bestimmendes LUMER-Produkt.

2.11.21 Satz (BONSALL und DUNCAN [29](1971), Seite 86). *Sei X ein normierter Vektorraum über $\mathbb{K}$, $[\cdot, \cdot]$ ein die Norm bestimmendes LUMER-Produkt auf X und $T \in \mathcal{L}(X)$. Dann gilt:*

$$(a) \quad W(T) \subseteq V_{\text{spatial}}(T).$$

$$(b) \quad cl\left(co W(T)\right) = cl\left(co V_{\text{spatial}}(T)\right) = V(T).$$

(Siehe aber auch 3.4.9.)

2.11.22 Hermitezität in normierten Vektorräumen. *Sei X ein normierter Vektorraum über $\mathbb{C}$ und $1 \in S_X$. Dann heißt ein $x \in X$ 1-hermitesch, wenn $V(1; x) \subseteq \mathbb{R}$. Die Menge aller 1-hermiteschen $x \in X$ wird mit $X_{1,\text{herm}}$ bezeichnet. Ein $x \in X$ ist also genau dann 1-hermitesch, wenn $\ell(x)$ für alle 1-unitalen Zustände ℓ von X reell ist. Wenn keine Missverständnisse möglich sind, wird als Bezeichnung auch einfach nur X_{herm} verwendet.*

*Ist A eine unitale Algebra über $\mathbb{C}$, so wird, falls nicht etwas anderes ausdrücklich vereinbart ist, für 1 die Eins gesetzt und in diesem Fall sagt man anstelle von e-hermitesch einfach nur *hermitesch*. Somit ist ein $a \in A$ genau dann hermitesch, wenn $\ell(a)$ für alle unitalen Zustände ℓ von A reell ist, und gemäß Satz 2.11.17(e) ist dies genau dann der Fall, wenn der numerische Wertebereich von a reell ist. Ein $a \in A$ heißt *schiefhermitesch*, falls ia hermitesch ist; für Weiteres siehe ROSENTHAL [249](1984).*

Als eine Menge eines Positivitätskonzeptes (siehe 2.3.32) hat man hier:

$$\text{Pos}(V; A) := \{a \in A : V(a) \subseteq \mathbb{R}^+\}.$$

2.11.23 Satz. *Sei A eine unitale Algebra über $\mathbb{C}$. Dann gelten die folgenden Aussagen, wobei bei (b), (c) und (d) A als eine unitale Banachalgebra über $\mathbb{C}$ vorausgesetzt sei.*

$$(a) \quad A_{\text{herm}} \cap iA_{\text{herm}} = \{0\} \text{ und } A_{\text{herm}} \text{ ist ein Vektorraum über } \mathbb{R}.$$

$$(b) \quad A_{\text{herm}} \text{ ist schwach abgeschlossen (also auch abgeschlossen und damit ein Banachraum über } \mathbb{R}).$$

$$(c) \quad \text{Pos}(V; A) \text{ ist im Banachraum } A_{\text{herm}} \text{ über } \mathbb{R} \text{ ein abgeschlossener punktierter spitzer Kegel mit der Eins als inneren Punkt von } \text{Pos}(V; A).$$

$$(d) \quad A_{\text{herm},(\mathbb{C})}, \text{ kanonisch mit der Norm von } A \text{ versehen, ist ein Banachraum über } \mathbb{C}.$$

$$(e) \quad A_{\text{herm},(\mathbb{C})} \text{ ist genau dann eine Algebra über } \mathbb{C}, \text{ wenn mit jedem } x \in A_{\text{herm}} \text{ auch stets } x^2 \text{ in } A_{\text{herm}} \text{ liegt.}$$

Beweis. (a) BONSALL und DUNCAN [29](1971), Seite 49. (b) PALMER [231] (1994), Seite 265. (c) BONSALL und DUNCAN [29](1971), Seite 48 und DORAN und BELFI [77](1986), Seite 193. (d) BONSALL und DUNCAN [29](1971), Seite 50. (e) BONSALL und DUNCAN [29](1971), Seite 59, Theorem 3. $\square$

2.11.24 (DORAN und BELFI [77](1986), Seite 193). Gilt in der Situation von Satz 2.11.23 die Gleichung $A = A_{\text{herm}} + iA_{\text{herm}}$, so heißt A eine *Vidav-Algebra*.

2.11.25. Sei A eine unitale Algebra über $\mathbb{C}$ mit $A \neq \{0\}$. Nach Satz 2.11.23(a) ist A_{herm} kein Untervektorraum von A über $\mathbb{C}$. Ist demgemäß von A_{herm} als einem Vektorraum die Rede, so wird er im Allgemeinen stillschweigend als ein Vektorraum über $\mathbb{R}$ angenommen. Entsprechendes gilt, wenn von A_{herm} als einem Untervektorraum die Rede ist. Wenn kein Missverständnis möglich ist, kann man daher insbesondere von A_{herm} als einen Untervektorraum von A reden, auch wenn, genau genommen, A_{herm} als ein Untervektorraum von $A_{\mathbb{R}}$ gemeint ist.

2.11.26 (MCCRIMMON [210](2004), Seite 27). Eine komplexe Zahl z ist genau dann reell, wenn $\exp(itz) \in S_{\mathbb{C}}$ gilt für alle $t \in \mathbb{R}$. In diesem Sinne sind gemäß dem folgenden Satz genau die hermiteschen Elemente einer unitalen Banachalgebra über $\mathbb{C}$ als „reell“ aufzufassen.

2.11.27 Satz (BONSALL und DUNCAN [29](1971), Seite 46). *Sei A eine unitale Banachalgebra über $\mathbb{C}$ und $T \in A$. Dann sind die folgenden neun Aussagen äquivalent:*

- (a) T ist hermitesch. (Siehe 2.11.22.)
- (a.i) $V(T) \subseteq \mathbb{R}$.
- (a.ii) $\text{Im } V(T) = \{0\}$.
- (a.iii) $\text{Re } V(iT) = \{0\}$.
- (a.iv) $\max \text{Re } V(iT) = \max \text{Re } V(-iT) = \{0\}$.
- (a.v) $\langle TS, S^* \rangle \in \mathbb{R}$ für alle $S \in A$, $S^* \in A^*$ mit $S^*(S) = \|S^*\| \|S\|$.
- (b) T ist hermitesch nach Vidav, das heißt,

$$\lim_{\substack{t \in \mathbb{R} \\ t \rightarrow 0}} \frac{1}{t} (\|e + iT\| - 1) = 0$$

oder in der Schreibweise mit der Landau-Symbolik:

$$\|e + iT\| = 1 + o(t) \quad \text{mit } t \in \mathbb{R}, t \rightarrow 0.$$

(c) $\|\exp(itT)\| \leq 1$ für alle $t \in \mathbb{R}$.

(d) $\|\exp(itT)\| = 1$ für alle $t \in \mathbb{R}$.

Beweis. (c) $\Rightarrow$ (d): Gelte (c) und sei $t \in \mathbb{R}$. Setze $S := \exp(itT)$. Nach dem noch zu folgenden Beispiel 4.1.3 existiert S^{-1} und ist gleich $\exp(i(-t)T)$, also $\|S^{-1}\| \leq 1$. Wegen $\|S^{-1}\| = \sup_{y \in A^\times} \frac{\|S^{-1}y\|}{\|y\|} = \sup_{x \in A^\times} \frac{\|S^{-1}Sx\|}{\|Sx\|} \geq \frac{1}{\|S\|}$, wäre bei $\|S\| < 1$ also $\|S^{-1}\| > 1$, Widerspruch.

Bezüglich (a.v) bemerke man nur: Sei $S \in A$, $S^* \in A^*$ mit $S^*(S) = \|S^*\| \|S\| \neq 0$. Setze $R := \frac{S}{\|S\|}$ und $R^* := \frac{S^*}{\|S^*\|}$. Dann ist $R^*(R) = 1$, also $(R, R^*) \in \Pi_A$ und es ist $\langle TS, S^* \rangle = \langle TR, R^* \rangle \|S\| \|S^*\|$. $\square$

2.11.28 Bemerkung (DORAN und BELFI [77](1986), Seite 178). Vidav bemerkte für $x \in \mathbb{C}$:

$$x \in \mathbb{R} \iff |1 + i\lambda x| = 1 + o(\lambda) \quad \text{mit } \lambda \in \mathbb{R}, \lambda \rightarrow 0.$$

2.11.29 Lemma (UPMEIER [279](1985), Seite 115). *Sei X ein Banachraum über $\mathbb{K}$ und $T \in GL(X)$. Dann sind folgende Aussagen äquivalent:*

(a) T ist eine Isometrie.

(b) $\|T\| = \|T^{-1}\| \leq 1$.

(c) $\|T\| = \|T^{-1}\| = 1$.

2.11.30 Korollar. Ist X ein Banachraum über $\mathbb{C}$ und $T \in \mathcal{L}(X)$, dann gilt $\exp(itT) \in G\mathcal{L}(X)$ für alle $t \in \mathbb{R}$, und daher sind äquivalent:

(a) T ist hermitesch.

(b) $\exp(itT)$ ist eine Isometrie für alle $t \in \mathbb{R}$.

2.11.31 (BONSALL und DUNCAN [29](1971), Seite 38). Sei A eine unitale Algebra über $\mathbb{C}$. Dann ist die Eins von A ein Vertex von B_A (Definition 2.11.1). Siehe auch 3.5.33.

2.12 Spektrum

2.12.1 Definition (RICKART [246](1960), Seiten 30 und 32). Sei A eine Algebra über $\mathbb{K}$ und $a \in A$. Sei B eine Algebra über $\mathbb{L}$, $\mathbb{L} = \mathbb{R}$ oder $\mathbb{L} = \mathbb{C}$, mit Eins e . Ist A eingebettet in B , dann setze man

$$\sigma_{A,B}(a) := \{\lambda \in \mathbb{L} : \lambda e - a \text{ ist in } B \text{ nicht invertierbar}\}.$$

Das *Spektrum* $\sigma(a)$ von a ist definiert als:

- $\sigma_{A,A}(a)$, falls $\mathbb{K} = \mathbb{C}$ und A eine Eins hat;
- $\sigma_{A,A^1}(a)$, falls $\mathbb{K} = \mathbb{C}$ und A keine Eins hat;
- $\sigma_{A,A_{(\mathbb{C})}}(a)$, falls $\mathbb{K} = \mathbb{R}$ und A eine Eins hat;
- $\sigma_{A,(A_{(\mathbb{C})})^1}(a)$, falls $\mathbb{K} = \mathbb{R}$ und A keine Eins hat.

Soll betont werden, dass sich das Spektrum auf die Algebra A bezieht, dann wird anstelle von $\sigma(a)$ die Bezeichnung $\sigma_A(a)$ verwendet.

2.12.2 Bemerkung. Sei A eine Algebra über $\mathbb{C}$. Für $a \in A$ setze $\varsigma(a) := \{\lambda \in \mathbb{C} : \lambda \neq 0 \text{ und } \frac{1}{\lambda} \notin G^q(A)\}$. Hat A eine Eins und ist $a \in G(A)$, dann ist $\sigma(a) = \varsigma(a)$. Hat A keine Eins oder ist $a \notin G(A)$, dann ist $\sigma(a) = \varsigma(a) \cup \{0\}$.

Beweis. PALMER [231](1994), Seite 196; RICKART [246](1960). □

2.12.3 Definition. Sei A eine Algebra über $\mathbb{K}$ und $a \in A$. Dann ist der *Spektralradius* $\rho(a)$ von a definiert als das in $\overline{\mathbb{R}}$ genommene Supremum aller $|\lambda|$ mit $\lambda \in \sigma(a)$.

2.12.4 Bemerkungen. (a) Sei A eine Algebra über $\mathbb{K}$ und $a \in A$. Hat A keine Eins, dann gilt $\sigma_A(a) = \sigma_{A^1}(a)$; das heißt insbesondere für den Fall $\mathbb{K} = \mathbb{R}$, dass $\sigma_{A,(A_{(\mathbb{C})})^1}(a) = \sigma_{A,(A^1)_{(\mathbb{C})}}(a)$ ist. Hat allerdings A bereits eine Eins, dann müssen die beiden Mengen $\sigma_A(a)$ und $\sigma_{A^1}(a)$ nicht gleich sein!

(b) Sei A eine Algebra über $\mathbb{C}$, versehen mit einer Algebrahalbnorm p , für die im Fall, dass A eine Eins hat, $p(e) \neq 0$ gelte, oder sei (A, p) eine normierte Algebra über $\mathbb{K}$. Dann hat jedes $x \in A$ ein nicht leeres Spektrum und es gilt

$$\lim_{n \rightarrow \infty} (p(x^n))^{1/n} \leq \rho(x) \quad \text{für alle } x \in A;$$

im Fall, dass A eine normierte Algebra über $\mathbb{K}$ ist, existiert für jedes $x \in A$ ein $\lambda \in \sigma(x)$ mit

$$\lim_{n \rightarrow \infty} \|x^n\|^{1/n} \leq |\lambda|.$$

(c) Sei A eine Banachalgebra über $\mathbb{K}$. Sei $x \in A$. Dann ist $\sigma(x)$ eine kompakte Menge und es gilt

$$\lim_{n \rightarrow \infty} \|x^n\|^{1/n} = \rho(x).$$

(d) Ist A eine Algebra über $\mathbb{K}$ ohne Eins, dann ist $0 \in \sigma(a)$ für alle $a \in A$. Manche Autoren nennen im Fall solch einer Algebra die Menge $\sigma(a)$ auch das *Quasi-Spektrum* von a und bezeichnen es mit $\sigma'(a)$, während dann in dessen $\sigma(a)$ die Null nicht enthalten ist.

Beweis. (a) RICKART [246](1960), Seite 32, Theorem 1.6.9(a); LARSEN [196] (1973), Seite 55. (b) Bezüglich der Halbnormen-Aussage siehe PALMER [231] (1994), Seite 210, Theorem 2.2.2. Für die Norm-Aussage siehe RICKART [246] (1960), Seite 28 (Siehe auch INGELSTAM [145](1964), Seite 244.). Die Aussage für $\mathbb{K} = \mathbb{R}$ wird per Komplexifizierung erhalten; für einen reell-intrinsischen Beweis siehe KULKARNI und LIMAYE [190](1980). (c) RICKART [246](1960), Seite 30. (d) RICKART [246](1960), Seite 32. $\square$

2.12.5 Bemerkung. Sei A eine Algebra über $\mathbb{K}$. Je nachdem, ob $\mathbb{K}$ gleich $\mathbb{R}$ oder $\mathbb{C}$ ist, und ob A eine Eins e hat oder nicht, liegt einer von vier Fällen vor. Diese seien selbsterklärend mit $(\mathbb{R}, e)$, $(\mathbb{R}, -)$, $(\mathbb{C}, e)$ und $(\mathbb{C}, -)$ bezeichnet. Sei x ein beliebiges Element von A . Nachstehend ist für jeden der vier Fälle das Spektrum von x angeführt.

$$\begin{aligned} (\mathbb{C}, e) & : \quad \sigma_{A,A}(x) = \{\lambda \in \mathbb{C} : \lambda e - x \notin G(A)\}; \\ (\mathbb{C}, -) & : \quad \sigma_{A,A^1}(x) = \{\lambda \in \mathbb{C} : (-x, \lambda) \notin G(A^1)\}; \\ (\mathbb{R}, e) & : \quad \sigma_{A,A(\mathbb{C})}(x) = \{\lambda \in \mathbb{C} : \lambda \otimes e - 1 \otimes x \notin G(\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} A)\}; \\ (\mathbb{R}, -) & : \quad \sigma_{A,(A(\mathbb{C}))^1}(x) = \{\lambda \in \mathbb{C} : (-1 \otimes x, \lambda) \notin G((\mathbb{C}_{\mathbb{R}} \otimes_{\mathbb{R}} A)^1)\}. \end{aligned}$$

2.12.6 Bemerkung (RICKART [246](1960), Seiten 28 und 31). Sei A eine Algebra über $\mathbb{K}$ und $a \in A$. Es gilt für

$$\begin{aligned} \mathbb{K} = \mathbb{C} & : \quad \sigma_{A_{\mathbb{R}}}(a) = \sigma(a) \cup \overline{\sigma(a)}; \\ \mathbb{K} = \mathbb{R} & : \quad \sigma(a) \text{ ist selbstadjungiert}; \\ \mathbb{K} = \mathbb{R}, A \text{ hat eine Eins} & : \quad \operatorname{Re}(\sigma(a)) = \sigma_{A,A}(a); \\ \mathbb{K} = \mathbb{R}, A \text{ hat keine Eins} & : \quad \operatorname{Re}(\sigma(a)) = \sigma_{A,A^1}(a); \\ \mathbb{K} = \mathbb{R} & : \quad \forall \lambda \in \mathbb{R}^{\times} : \lambda \in \sigma(a) \Leftrightarrow \frac{1}{\lambda} a \notin G^q(A). \end{aligned}$$

2.12.7 Definition. Sei A eine Algebra über $\mathbb{K}$. Dann seien als zwei weitere Konzepte von Positivität die folgenden beiden Mengen definiert:

$$\begin{aligned} \operatorname{Pos}(V\sigma; A) & := \{a \in A : V(a) \subseteq \mathbb{R}, \sigma(a) \subseteq \mathbb{R}^+\} \text{ und} \\ \operatorname{Pos}(\sigma; A) & := \{a \in A : \sigma(a) \subseteq \mathbb{R}^+\}. \end{aligned}$$

2.12.8 Satz (BONSALL und DUNCAN [29](1971), Seite 19). *Sei A eine unitale Banachalgebra über $\mathbb{K}$. Dann gilt für den numerischen Wertebereich (Definition 2.11.15):*

$$\sigma(a) \cap \mathbb{K} \subseteq V(a) \subseteq \|a\| \cdot B_{\mathbb{K}} \quad \text{für alle } a \in A.$$

Beweis. Sei $\lambda \in \sigma(a) \cap \mathbb{K}$. Also ist $\lambda - a$ nicht in A invertierbar. (Im Fall von $\mathbb{K} = \mathbb{R}$ ist per Definition $\lambda - a$ nicht in $A_{(\mathbb{C})}$ invertierbar. Dann aber erst recht nicht in A .) Das heißt, $\lambda - a$ ist nicht links-invertierbar oder $\lambda - a$ ist nicht rechts-invertierbar. Also ist $A \cdot (\lambda - a)$ ein properes Linksideal von A oder es ist $(\lambda - a) \cdot A$ ein properes Rechtsideal von A . Bezeichne $\mathfrak{p}$ solch ein properes einseitige Ideal von A . Kein Element von $\mathfrak{p}$ ist invertierbar. (Denn, angenommen doch, etwa $b = c(\lambda - a)$, dann wäre $b^{-1} = (\lambda - a)^{-1}c^{-1}$, $(\lambda - a)^{-1} = b^{-1}c$, Widerspruch.) Bekanntlich ist ein Element x einer unitalen Banachalgebra über $\mathbb{K}$ invertierbar, wenn $\|x - e\| < 1$ gilt (GELFAND [106](1941), §2, Hilfssatz 1; NAGUMO [220](1936), §2, Hilfssatz 2). Somit gilt hier für jedes $b \in \mathfrak{p}$ die Ungleichung $\|b - e\| \geq 1$. Insbesondere gilt diese Ungleichung für alle $b \in \text{cl } \mathfrak{p}$. Das heißt, die Eins e von A liegt nicht im Abschluss des Untervektorraumes $\mathfrak{p}$ von A . Nach einer Folgerung aus dem Satz von Hahn-Banach (TAYLOR und LAY [275](1980), Seite 136, Theorem 3.4) existiert ein $\ell \in A^*$ mit $\|\ell\| = 1$, $\ell(b) = 0$ für alle $b \in \mathfrak{p}$ und $\ell(e) = 1$. Das heißt, es ist $(e, \ell) \in \Pi_A$ (siehe Definition 2.11.15). Da insbesondere $\lambda - a$ ein Element von $\mathfrak{p}$ ist, gilt also $\ell(\lambda - a) = 0$. Da $\lambda \in \mathbb{K}$, gilt $\ell(\lambda - a) = \ell(\lambda e) - \ell(a) = \lambda - \ell(a)$, also $\lambda = \ell(a)$. Wegen $\ell(a) = \langle a, \ell \rangle = \langle ae, \ell \rangle$ mithin $\lambda \in V(a)$ und es gilt $\sigma(a) \cap \mathbb{K} \subseteq V(a)$. Sei $\lambda \in V(a)$. Somit $\lambda = \langle ab, \ell \rangle$ für ein $(b, \ell) \in \Pi_A$. Also $|\lambda| = |\ell(ab)| \leq \|\ell\| \|a\| \|b\| = \|a\|$ und somit gilt $\lambda \in \|a\| \cdot B_{\mathbb{K}}$. $\square$

2.12.9 Satz (BONSALL und DUNCAN [29](1971), Seite 88). *Sei X ein Banachraum über $\mathbb{K}$ und $T \in \mathcal{L}(X)$. Dann gilt: $\sigma(T) \subseteq \text{cl}(V_{\text{spatial}}(T))$, falls $\mathbb{K} = \mathbb{C}$.*

2.12.10 Satz (BONSALL und DUNCAN [29](1971), Seite 53). *Sei A eine unitale Banachalgebra über $\mathbb{C}$. Dann gilt:*

$$V(h) = \text{co } \sigma(h) \subseteq \mathbb{R} \quad \text{für alle } h \in A_{\text{herm}}.$$

2.12.11 Satz (INGELSTAM [145](1964)). *Sei A eine Banachalgebra über $\mathbb{R}$. Dann sind äquivalent:*

- (a) A ist vom stark reellen Typus.
- (b) $\sigma(a) \subseteq \mathbb{R}$ für alle $a \in A$.
- (c) A^1 ist vom stark reellen Typus.

2.12.12 Satz (BONSALL und DUNCAN [30](1973), Seite 207). *Sei A eine unitale Banachalgebra über $\mathbb{C}$. Für alle $a \in A_{\text{herm}}$ gilt $V(a) = \text{co } \sigma(a)$ (siehe Satz 2.11.17). Somit sind für alle $a \in A$ die folgenden beiden Aussagen äquivalent:*

- (a) $a \in \text{Pos}(V; A)$. (Definition 2.11.22)
- (b) $a \in \text{Pos}(V\sigma; A)$. (Definition 2.12.7)

2.12.13 (PALMER [231](1994), Seiten 189 und 225). *Sei A eine Algebra über $\mathbb{C}$. Dann ist das Jacobson-Radikal $\mathfrak{J}(A)$ von A gleich der Menge $\{y \in A : \rho(xy) = 0 \text{ für alle } x \in A\}$. Sei $\mathfrak{a}$ ein Ideal von A . Dann sind die folgenden Aussagen äquivalent:*

- (a) $\sigma(x + y) = \sigma(x)$ für alle $x \in A$, $y \in \mathfrak{a}$;
- (b) $\sigma(y) = \{0\}$ für alle $y \in \mathfrak{a}$;
- (c) $\rho(x + y) = \rho(x)$ für alle $x \in A$, $y \in \mathfrak{a}$;
- (d) $\rho(y) = 0$ für alle $y \in \mathfrak{a}$.

Das größte Ideal $\mathfrak{a}$ von A , das eine (und damit alle) der vorstehenden Aussagen erfüllt, ist gleich dem Jacobson-Radikal von A . Somit ist zum Beispiel gemäß (d) das Jacobson-Radikal von A gleich dem größten Ideal von A , auf dem der Spektralradius identisch gleich Null ist.

2.12.14 Satz (RICKART [246](1960), Seiten 11 und 30). *Sei A eine Banachalgebra über $\mathbb{K}$. Dann sind äquivalent:*

$$(a) \quad \rho(a) = \|a\| \quad \text{für alle } a \in A.$$

$$(b) \quad \|a^2\| = \|a\|^2 \quad \text{für alle } a \in A.$$

2.12.15 Satz (SINCLAIR [262](1971), Seite 447, Satz 2). *Sei A eine unital Banachalgebra über $\mathbb{C}$. Dann gilt für alle $h \in A_{\text{herm}}$ und für alle $\lambda \in \mathbb{C}$ die Gleichung $\rho(h + \lambda e) = \|h + \lambda e\|$.*

Beweis. Für einen direkten Beweis des Spezialfalls $\lambda = 0$ siehe zum Beispiel BONSALL und DUNCAN [30](1973), Seite 57, Theorem 17; UPMEIER [279](1985), Seite 242, Lemma 14.30; PALMER [231](1994), Seite 267, Theorem 2.6.7(e); ISIDRO und STACHÓ [149](1984), Seite 245, Theorem 10.27 $\square$

2.12.16 Kommutativität (PALMER [231](1994), Seiten 238 und 311). Als eine Folge des Satzes von Liouville und des Satzes von Hahn-Banach hat man: Ist A eine normierte unital Algebra über $\mathbb{C}$, so ist sie genau dann kommutativ, wenn ein $k \in \mathbb{R}$ existiert mit $\|xy\| \leq k\|yx\|$ für alle $x, y \in A$.

Ist also A eine normierte Algebra über $\mathbb{C}$ und ist auf A^1 die Norm per $(a, \lambda) \mapsto \|a\| + |\lambda|$ für alle $x \in A, \lambda \in \mathbb{C}$ definiert (siehe dazu Satz 2.3.24), so ist A genau dann kommutativ, wenn ein $k \in \mathbb{R}$ existiert mit $\|xy\| \leq k\|yx\|$ für alle $x, y \in A^1$; hinreichend für die Kommutativität von A ist, dass ein $k \in \mathbb{R}$ existiert mit $\|x\|^2 \leq k\|x^2\|$ für alle $x \in A$.

2.12.17 Gelfand-Räume (PALMER [231](1994), Seite 303 für den Fall $\mathbb{K} = \mathbb{C}$; LI [198](2003), Seiten 28, 41 und 78 und GOODEARL [117](1982), Seite 95 für den Fall $\mathbb{K} = \mathbb{R}$; RICKART [246](1960), Seite 108 für den allgemeinen Fall $\mathbb{K}$).

Sei A eine Algebra über $\mathbb{K}$. Setze

$$\Gamma_A := \{\ell: A \rightarrow \mathbb{C}_{\mathbb{R}} : \ell \neq 0, \ell \text{ ein Algebra-Homomorphismus}\}, \quad \text{falls } \mathbb{K} = \mathbb{R}$$

und

$$\Gamma_A := \{\ell: A \rightarrow \mathbb{C} : \ell \neq 0, \ell \text{ ein Algebra-Homomorphismus}\}, \quad \text{falls } \mathbb{K} = \mathbb{C}.$$

Setze $\Gamma_A^0 := \Gamma_A \cup \{0\}$. Für nicht leeres Γ_A sei für jedes $x \in A$ die Abbildung $\hat{x}: \Gamma_A \rightarrow \mathbb{C}$, falls $\mathbb{K} = \mathbb{C}$, bzw. $\hat{x}: \Gamma_A \rightarrow \mathbb{C}_{\mathbb{R}}$, falls $\mathbb{K} = \mathbb{R}$, die Auswertung bei x , also $\hat{x}(\ell) = \ell(x)$ für alle $\ell \in \Gamma_A$; $\hat{x}$ heißt die *Gelfand-Transformierte* von x ; für Γ_A^0 wird für jedes $x \in A$ dann die Abbildung $\hat{x}$ per $\hat{x}(0) := 0$ fortgesetzt. Die bezüglich all dieser Abbildungen $\hat{x}$ definierte Initialtopologie auf Γ_A (bzw. Γ_A^0) heißt die *Gelfand-Topologie*; es ist die Topologie der punktweisen Konvergenz, im Fall $\mathbb{K} = \mathbb{C}$ also gleich $\sigma(A^\#, A)|_{\Gamma_A}$ (bzw. $\sigma(A^\#, A)|_{\Gamma_A^0}$) und im Fall $\mathbb{K} = \mathbb{R}$ gleich $\sigma(L(A, \mathbb{C}_{\mathbb{R}}), A)|_{\Gamma_A}$ (bzw. $\sigma(L(A, \mathbb{C}_{\mathbb{R}}), A)|_{\Gamma_A^0}$); sie wird hier symbolisch auch als $\sigma(\Gamma_A, A)$ (bzw. $\sigma(\Gamma_A^0, A)$) geschrieben. Äquivalent dazu erhält man die Gelfand-Topologie, wenn man die Menge $\mathbb{C}^A$ bzw. $(\mathbb{C}_{\mathbb{R}})^A$ aller komplexwertigen Funktionen auf A mit der Produkttopologie ausstattet und dann Γ_A mit der Relativtopologie versieht. Für nicht leeres Γ_A heißt die mit der Gelfand-Topologie $\sigma(\Gamma_A, A)$ versehene Menge Γ_A der *Gelfand-Raum* von A ; $(\Gamma_A^0, \sigma(\Gamma_A^0, A))$ heißt der *erweiterte Gelfand-Raum* von A . Beide Räume sind Hausdorff'sch. Die durch $x \mapsto \hat{x}$ definierte Abbildung $\hat{\cdot}: A \rightarrow C(\Gamma_A, \mathbb{C})$ im Fall $\mathbb{K} = \mathbb{C}$ bzw. $\hat{\cdot}: A \rightarrow C(\Gamma_A, \mathbb{C}_{\mathbb{R}})$ im Fall $\mathbb{K} = \mathbb{R}$ ist ein Algebra-Homomorphismus und heißt der *Gelfand-Homomorphismus*; er ist unital, falls die Algebra eine Eins hat.

Falls A eine Banachalgebra über $\mathbb{C}$ ist, so gilt $\|\ell\| \leq 1$ für alle $\ell \in \Gamma_A$ und $\sigma(\Gamma_A, A) = \sigma(A^*, A)|_{\Gamma_A}$, das heißt in Worten, dass die Gelfand-Topologie die relativierte schwach*-Topologie ist; ist A außerdem unital, so gilt $\|\ell\| = 1$ für alle $\ell \in \Gamma_A$.

Falls A eine kommutative Banachalgebra über $\mathbb{R}$ ist, so gilt einerseits für alle $\ell \in \Gamma_A$ die Ungleichung $\|\ell\| \leq 1$, wobei jeweils $\|\ell\| = 1$ gilt, falls A unital ist, und andererseits die Äquivalenz $\ell \in A^* \Leftrightarrow \ell(A) \subseteq \mathbb{R}$.

Sei A eine Banachalgebra über $\mathbb{C}$ oder eine kommutative Banachalgebra über $\mathbb{K}$. Dann ist Γ_A lokal kompakt und Γ_A^0 ist seine Ein-Punkt-Kompaktifizierung. Falls A eine Eins hat, so ist 0 ein isolierter Punkt von Γ_A^0 und der Gelfand-Raum Γ_A ist kompakt.

Ist A eine Algebra über $\mathbb{C}$ ohne Eins mit nicht leerem Γ_A , so ist die Restriktion von jedem $\ell \in \Gamma_{A^1}$ auf A ein Homöomorphismus von Γ_{A^1} auf Γ_A^0 .

Sei A eine Banachalgebra über $\mathbb{C}$ oder eine kommutative Banachalgebra über $\mathbb{K}$. Da jedes $\hat{x}$ bei $0 \in \Gamma_A^0$ verschwindet, definiert die Abbildung $x \mapsto \hat{x}$ den Algebra-Homomorphismus $\hat{\cdot}: A \rightarrow C_0(\Gamma_A, \mathbb{C})$ im Fall $\mathbb{K} = \mathbb{C}$ bzw. $\hat{\cdot}: A \rightarrow C_0(\Gamma_A, \mathbb{C}_{\mathbb{R}})$ im Fall $\mathbb{K} = \mathbb{R}$; es ist $\|\hat{\cdot}\| \leq 1$ und falls A unital ist, gilt $\|\hat{\cdot}\| = 1$.

Eine kommutative Banachalgebra über $\mathbb{R}$ mit nicht leerem Γ_A ist genau dann vom stark reellen Typus (siehe Definition 2.9.11(c)), wenn $\Gamma_A(A) \subseteq \mathbb{R}$ gilt.

Kapitel 3

Algebren mit Involutionen

3.1 Involutionen auf Vektorräumen

3.1.1 Streng komplexe Strukturen II. Zum Satz 2.9.12 hat man beobachtet: Ist X ein Vektorraum über $\mathbb{R}$, versehen mit einer komplexen Struktur j , so wird per $ix := j(x)$ eine komplexe skalare Multiplikation auf X eingeführt, und ist umgekehrt X ein Vektorraum über $\mathbb{C}$, so ist die durch $j(x) := ix$ definierte Abbildung eine komplexe Struktur auf dem X zugrunde liegenden reellen Vektorraum $X_{\mathbb{R}}$. Im besagten Satz wurde dann als Essenz festgehalten, dass genau die Vektorräume über $\mathbb{C}$ als die Vektorräume über $\mathbb{R}$, auf denen eine komplexe Struktur existiert, aufgefasst werden können.

In 2.9.20 wurde der Begriff der streng komplexen Struktur auf einem reellen Vektorraum eingeführt, um in komplexen Vektorräumen etwas Ähnliches wie die Konjugiertenbildung bei den komplexen Zahlen zur Verfügung zu haben. Dies setzte aber voraus, dass eine komplexe Struktur vorhanden ist, was aber, wie in Bemerkung 2.9.19 erwähnt, genau in den reellen Vektorräumen mit endlicher, ungeradzahlgiger Dimension eben nicht der Fall ist. Da die für streng komplexe Strukturen in 2.9.20(a) geforderte Gültigkeit der Gleichung (2.35) in komplexen Vektorräumen genau die Eigenschaft konjugiert-linear ausdrückt, ist das nachfolgende, etwas weniger restriktive Konzept, das auch die reellen Vektorräume berücksichtigt, auf denen keine komplexe Struktur existiert, der Involutionen naheliegend.

3.1.2 Definition (BONSALL und DUNCAN [30](1973), Seite 63). Eine *Involution* (genauer eine *Vektorraum-Involution*) auf einem Vektorraum X über $\mathbb{K}$ ist eine semilineare Abbildung $*$: $X \rightarrow X$, $x \mapsto x^*$ mit $(x^*)^* = x$.

3.1.3 Streng komplexe Strukturen III. Sei X ein Vektorraum über $\mathbb{K}$ derart, dass $X_{\mathbb{R}}$ mit einer streng komplexen Struktur (j, k) versehen ist. Dann ist k eine Vektorrauminvolution auf X . Zur Diskussion der Umkehrung dieser Aussage muss man den reellen von dem komplexen Fall getrennt betrachten. Ist X ein Vektorraum über $\mathbb{R}$ von endlicher, ungeradzahlgiger Dimension, so kann offensichtlich mangels einer komplexen Struktur keine streng komplexe Struktur vorliegen. Aber selbst, wenn X ein Vektorraum über $\mathbb{R}$ mit endlicher und geradzahlgiger oder unendlicher Dimension ist und X eine Involution $*$ besitzt, so kann man zwar für jedes $x \in X$ den Term $j(x^*)$ als $ix^* \in X_{\mathbb{C}}$ auffassen, da aber die Involution $*$ zwar $\mathbb{R}$ -linear, aber nicht konjugiert-linear auf $X_{\mathbb{C}}$ sein muss, muss die Involution $*$ nicht notwendig mit j eine streng komplexe Struktur $(j, *)$ bilden. Im Fall, dass X ein Vektorraum über $\mathbb{C}$ ist, der eine Involution $*$ trägt, gilt aber die erwähnte Umkehrung immer: Als Vektorraum über $\mathbb{C}$ ist $X_{\mathbb{R}}$ mit der per $j(x) := ix$ definierten komplexen Struktur versehen. Die Involution ist offenbar $\mathbb{R}$ -linear und für alle $x \in X$ gilt: $(j \circ *)(x) = j(x^*) = ix^* = -((-i)x^*) = -(ix)^* = -(x \circ j)(x)$, also $j \circ * = -* \circ j$, das heißt, $(j, *)$ ist eine streng komplexe Struktur auf $X_{\mathbb{R}}$.

Man kann also sagen: Während bei jedem Vektorraum mit einer streng komplexen Struktur die Konjugiertenbildung eine Involution ist, ist ein Vektorraum über $\mathbb{K}$, der mit einer Involution versehen ist, im Allgemeinen nur im Fall von $\mathbb{K} = \mathbb{C}$ stets auch mit einer streng komplexen Struktur versehen. Somit gilt: Genau die streng komplexen Vektorräume sind die Vektorräume über $\mathbb{C}$, die mit einer Involution versehen sind; dabei ist die jeweilige Konjugiertenbildung genau gleich der jeweiligen Involution.

3.1.4 Bemerkung und Definition. Sei X ein Vektorraum über $\mathbb{K}$ mit einer Involution $*$. Eine Teilmenge S von X heißt eine $*$ -Teilmenge, wenn sie abgeschlossen unter der Involution ist, das heißt, wenn $x \in S$ immer auch $x^* \in S$ impliziert. Man beachte, dass dann umgekehrt auch für jedes $x \in S$ per $y := x^*$ ein $y \in S$ mit $x = y^*$ existiert, das heißt, es gilt $S = S^*$, wobei hier S^* die Menge $\{s^* \in S : s \in S\}$ bezeichnet.

Ein $x \in X$ heißt *selbstadjungiert* bezüglich der Involution $*$ (oder auch $*$ -symmetrisch), wenn x ein Fixpunkt der Involution $*$ ist, das heißt, wenn $x = x^*$ gilt. Der *selbstadjungierte Teil* von X ist die Fixpunktmenge der Involution $*$, das heißt, die Menge aller selbstadjungierten $x \in X$, und wird mit $X_{(*)}$ bezeichnet. Mit $-*$ wird die mit (-1) multiplizierte Involution $*$ bezeichnet, also die Abbildung $-*: X \rightarrow X$, $x \mapsto x^{-*} := -x^*$. Sie ist ebenfalls eine Involution auf X und heißt *die zu der Involution $*$ gespiegelte Involution*. Ihre Fixpunktmenge wird mit $X_{(-*)}$ bezeichnet. Die Elemente von $X_{(-*)}$ heißen *antiselbstadjungiert* (oder auch *schief selbstadjungiert*, im Englischen *skew self-adjoint*).

3.1.5 Bemerkungen. Sei X ein Vektorraum über $\mathbb{K}$ mit einer Involution $*$. Dann gilt:

(a) $X_{(*)}$ und $X_{(-*)}$ sind Untervektorräume von $X_{\mathbb{R}}$. Ähnlich wie in 2.11.25 für die hermiteschen Elemente einer nichttrivialen unitalen Algebra über $\mathbb{C}$ gilt auch hier: Ist $\mathbb{K} = \mathbb{C}$ und $X \neq \{0\}$, so sind $X_{(*)}$ und $X_{(-*)}$ keine Untervektorräume von X . Ist also von $X_{(*)}$ oder $X_{(-*)}$ als einem Vektorraum die Rede, so wird er im Allgemeinen stillschweigend als ein Vektorraum über $\mathbb{R}$ angenommen; die Aussage für Untervektorräume ist entsprechend.

(b) Im Fall von $\mathbb{K} = \mathbb{C}$ gilt $X_{(-*)} = iX_{(*)}$. Diese Gleichung ist dahingehend von Bedeutung, da bisweilen Aussagen für den Fall $\mathbb{K} = \mathbb{C}$, in denen $iX_{(*)}$ verwendet wird, sich auf den $\mathbb{K} = \mathbb{R}$ -Fall als entsprechende, ähnliche Aussagen mit $X_{(-*)}$ übertragen. Siehe zum Beispiel Satz 3.5.27.

(c) $\frac{1}{2}(x + x^*) \in X_{(*)}$ und $\frac{1}{2}(x - x^*) \in X_{(-*)}$ für alle $x \in X$.

(d) $X_{\mathbb{R}} = X_{(*)} \dot{+} X_{(-*)}$.

(e) Im Fall von $\mathbb{K} = \mathbb{C}$ gilt gemäß Bemerkung 2.9.25: $X = X_{(*)}(\mathbb{C})$ (KLINGENBERG und KLEIN [186](1972), Seite 220, Punkt 6(iv)). Insbesondere ist also $X_{(*)}$ eine reelle Form des Vektorraumes X über $\mathbb{C}$ (Definition 2.9.23). Nach 3.1.3 ist per $j(x) := ix$ und $k := *$ das Paar (j, k) eine streng komplexe Struktur auf $X_{\mathbb{R}}$. Mithin gelten für $X_{\mathbb{R}}$ alle in 2.9.20(a) angegebenen Gleichungen, insbesondere gilt also auch $\operatorname{Re}(x) = \frac{1}{2}(x + x^*) \in X_{(*)}$ und $\operatorname{Im}(x) = \frac{1}{2i}(x - x^*) \in X_{(*)}$ für alle $x \in X$, womit also die Abbildungen Re und Im von 2.9.20(a) und von 2.9.25 jeweils übereinstimmen.

(f) Aus (c) folgen für alle $\ell \in X^{\#}$ die zwei nachstehenden Äquivalenzen: $\ell \upharpoonright X_{(*)} = 0 \Leftrightarrow \forall x \in X : \ell(x^*) = -\ell(x)$ und $\ell \upharpoonright X_{(-*)} = 0 \Leftrightarrow \forall x \in X : \ell(x^*) = \ell(x)$.

3.1.6 Bemerkung und Definition. Seien X und Y zwei Vektorräume über $\mathbb{K}$, die jeweils eine Involution tragen, die beide mit $*$ bezeichnet seien. Dann ist auf $L(X, Y)$ die Abbildung $*$: $T \mapsto * \circ T \circ *$ eine Involution. Für $T \in L(X, Y)$ heißt dann T^* die *Adjungierte* von T . Sie ist zu unterscheiden von der ebenfalls mit T^* bezeichneten dualen Abbildung von T und von der Hilbertraumadjungierten von T weiter unten. Aus dem jeweiligen Zusammenhang wird immer deutlich hervorgehen, was mit T^* gemeint ist.

Eine selbstadjungierte Abbildung $T: X \rightarrow Y$ wird auch $*$ -Abbildung genannt. Anschaulich sind die selbstadjungierten Abbildungen durch ihre Verträglichkeit mit den

beiden Involutionen (die eine von X , die andere von Y) charakterisiert: $T = T^* \Leftrightarrow T(x^*) = T(x)^*$ für alle $x \in X$. Insbesondere ist also jede Involution und per Definition jeder $*$ -Homomorphismus eine selbstadjungierte Abbildung.

Für den Fall $Y = \mathbb{K}$ sei als Involution von Y immer die komplexe Konjugation gemeint, wenn nicht ausdrücklich etwas anderes vereinbart ist.

3.1.7 Bemerkung. Sei X ein normierter Vektorraum über $\mathbb{K}$ mit einer Involution $*$. BONSALL und DUNCAN [30](1973), Seite 187, bemerken, dass der Dualraum X^* genau dann eine $*$ -Teilmenge von $X^\#$ ist, wenn die Involution auf X stetig ist.

3.1.8 Bemerkung und Definition. Sei X ein Vektorraum über $\mathbb{R}$. Dann ist

$$\begin{aligned} \bar{\cdot}: X_{(\mathbb{C})} &\rightarrow X_{(\mathbb{C})}, \\ (x, y) &\mapsto \overline{(x, y)} := (x, -y) \quad \text{für alle } x, y \in X \end{aligned}$$

eine Involution auf $X_{(\mathbb{C})}$; sie wird die *kartesische Involution* von $X_{(\mathbb{C})}$ genannt.

Beweis. Semilinearität: Sei $\lambda = \alpha + i\beta \in \mathbb{C}$ und seien $x, y \in X$. Nach Bemerkung 2.9.24 gilt dann: $\overline{\lambda \cdot (x, y)} = \overline{(\alpha + i\beta)(x, y)} = \overline{(\alpha x - \beta y, \alpha y + \beta x)} = (\alpha x - \beta y, -\alpha y - \beta x) = (\alpha x - (-\beta)(-y), \alpha(-y) + (-\beta)x) = (\alpha - i\beta)(x, -y) = \overline{\lambda} \cdot \overline{(x, y)}$. $\square$

Es gilt $X = X_{(\mathbb{C})}(\bar{\cdot})_{\mathbb{R}}$.

Sei k die Abbildung in Bemerkung 2.9.25. Dann bezeichnet man mit $k^{-1}(U)$ den Realteil eines Unterraumes U von $X_{(\mathbb{C})\mathbb{R}}$. Wegen $k(k^{-1}(U)) = U \cap k(X)$ identifiziert man $k^{-1}(U) \subseteq X$ mit $U \cap k(X)$, also mit der Menge der reellen Vektoren von $U \subseteq X_{(\mathbb{C})}$. Des Weiteren gilt damit $k^{-1}(U) = U_{(-)}$. In diesem Zusammenhang sei auch an 2.9.20 erinnert.

Sei A eine unitale Banachalgebra über $\mathbb{C}$. Dann ist die kartesische Involution des Banachraumes $A_{\text{herm},(\mathbb{C})}$ (siehe Satz 2.11.23(d)) stetig (BONSALL und DUNCAN [29](1971), Seite 50).

3.1.9 Definition. Eine *Konjugation auf einem normierten Vektorraum X über $\mathbb{K}$* ist eine Involution auf X , die eine Isometrie ist.

3.1.10 (KÖTHE [187](1966), Seite 431). Auf Vektorräumen hängen lineare Abbildungen mit Periode 2 und Projektionen eng zusammen: Sei X ein Vektorraum über $\mathbb{K}$. Für $T \in L(X)$ bezeichne $T^\perp$ die Abbildung $\text{Id}_X - T$. Man betrachte die Abbildung

$$\Lambda: L(X) \rightarrow L(X), \quad T \mapsto 2T - \text{Id}_X \quad (= T - T^\perp).$$

Die Abbildung Λ ist bijektiv und für $X \neq \{0\}$ nicht linear. Ihre Umkehrabbildung lautet

$$\Lambda^{-1}: L(X) \rightarrow L(X), \quad T \mapsto \frac{1}{2}(\text{Id}_X + T).$$

Sei $p \in L(X)$ eine Projektion. Dann ist

$$* := \Lambda(p)$$

eine lineare Abbildung mit Periode 2 und es gilt $X_{(*)} = p(X)$ und $\ker(p) = X_{(-*)}$. Siehe auch 3.5.35.

Beweis. $(2p - \text{Id}_X)^2 = 4p - 4p + \text{Id}_X = \text{Id}_X$. $x \in X_{(*)} \Leftrightarrow (2p - \text{Id}_X)x = x \Leftrightarrow 2px = 2x \Leftrightarrow px = x \Leftrightarrow x \in p(X)$. $x \in \ker(p) \Leftrightarrow 2p(x) - x = -x \Leftrightarrow x^* = -x \Leftrightarrow x \in X_{(-*)}$. $\square$

Ist umgekehrt $*$ eine lineare Abbildung auf X mit Periode 2, so ist

$$p := \Lambda^{-1}(*)$$

eine Projektion aus $L(X)$ mit $p(X) = X_{(*)}$ und $\ker(p) = X_{(-*)}$; man vergleiche mit der Abbildung Re aus 3.1.5(e).

Beweis. $(\frac{1}{2}(* + \text{Id}_X))^2 = \frac{1}{4}(\text{Id}_X + 2 \cdot * + \text{Id}_X) = \frac{1}{2}(* + \text{Id}_X)$. $x \in p(X) \Leftrightarrow px = x \Leftrightarrow \frac{1}{2}(x^* + x) = x \Leftrightarrow x^* + x = 2x \Leftrightarrow x^* = x \Leftrightarrow x \in X_{(*)}$. $x \in \ker(p) \Leftrightarrow \frac{1}{2}(x + x^*) = 0 \Leftrightarrow x^* = -x \Leftrightarrow x \in X_{(-*)}$. $\square$

Ist X ein Banachraum über $\mathbb{C}$, der mit einer stetigen Involution $*$ versehen ist, so ist $X_{\mathbb{R}}$ per $*$ mit einer stetigen linearen Abbildung mit Periode 2 versehen und man sieht, dass $X_{(*)}$ ein Banachraum über $\mathbb{R}$ ist.

3.1.11. Sei X ein Banachraum über $\mathbb{K}$. Ist $T \in \mathcal{L}(X)$ eine surjektive Isometrie mit Periode 2, so ist $p := \frac{1}{2}(\text{Id}_X + T)$ eine bikontraktive Projektion (zur Definition siehe 2.5.16(a)). Ist zum Beispiel X ein Banachraum über $\mathbb{C}$ mit einer Konjugation $*$, so ist $Re_{\mathbb{R}} := (x \mapsto \frac{1}{2}(x + x^*))$ eine bikontraktive Projektion aus $\mathcal{L}(X_{\mathbb{R}})$ mit $Re_{\mathbb{R}}(X_{\mathbb{R}}) = X_{(*)}$. Ein Banachraum Y über $\mathbb{R}$ heißt eine *reelle Form* eines Banachraumes X über $\mathbb{C}$, falls es auf X eine Konjugation $*$ gibt, so dass $Y = X_{(*)}$ gilt. Man beachte, dass jede solche reelle Form auch eine reelle Form des Y zugrunde liegenden Vektorraumes über $\mathbb{C}$ (Definition 2.9.23) ist.

Bezeichne $\mathcal{S}$ die Klasse der Banachräume, in denen jede bikontraktive Projektion p von der anfangs erwähnten Gestalt ist, das heißt, wenn die der Projektion p gemäß 3.1.10 zugeordnete lineare Abbildung $* = 2p - \text{Id}_X$ mit Periode 2 eine surjektive Isometrie ist. Dann ist $X \in \mathcal{S}$, wenn der Dualraum $X^* \in \mathcal{S}$ ist. Für alle $1 \leq p < \infty$ ist $L_p(\mu) \in \mathcal{S}$. FRIEDMAN und RUSSO [102](1987) zeigen, dass die weiter unten im Abschnitt 4.6 eingeführten JB^* -Tripel über $\mathbb{C}$ alle in $\mathcal{S}$ enthalten sind.

3.1.12 (LI [198](2003), Seite 2). Sei X ein normierter Vektorraum über $\mathbb{R}$. Sei $\|\cdot\|$ eine Norm auf $X_{(\mathbb{C})}$. Sei der Fall vorliegend, dass die kartesische Involution von $X_{(\mathbb{C})}$ isometrisch ist. Dann gilt: $\|x\| \leq \|x + iy\|$ für alle $x, y \in X$. Folglich hat man die Ungleichung $\max\{\|x\|, \|y\|\} \leq \|x + iy\|$ für alle $x, y \in X$.

3.1.13. LI [198](2003) definiert den Begriff der Komplexifizierung eines normierten Vektorraumes über $\mathbb{R}$ oder einer normierten Algebra über $\mathbb{R}$ wie folgt: Sei X ein normierter Vektorraum über $\mathbb{R}$ oder eine normierte Algebra über $\mathbb{R}$. Dann nennt LI den Raum $X_{(\mathbb{C})}$ eine Komplexifizierung von X , wenn $X_{(\mathbb{C})}$ eine Norm $\|\cdot\|$ trägt, mit der die kartesische Involution des Vektorraumes $X_{(\mathbb{C})}$ isometrisch ist und mit der X kanonisch isometrisch isomorph in $X_{(\mathbb{C})\mathbb{R}}$ eingebettet ist, das heißt, $\|\cdot\|_{(\mathbb{C})\mathbb{R}} \upharpoonright X$ die Norm von X ist. Da alle Normen, mit der $X_{(\mathbb{C})}$ solch eine Komplexifizierung von X ist, nach 3.1.12 äquivalent sind, ist solch eine Komplexifizierung von X bis auf die Normäquivalenz eindeutig. Dieser Begriff der Komplexifizierung wird hier im Folgenden als die *Komplexifizierung nach LI* genannt. Man beachte 3.2.16.

3.1.14. Sei $1 \leq p \leq \infty$ und gelten die Bezeichnungen von 2.9.27. Ist X ein normierter Vektorraum über $\mathbb{R}$, so ist die kartesische Involution von $(X_{(\mathbb{C})}, \|\cdot\|_p)$ isometrisch, also eine Konjugation. Ebenso hat man: Ist A eine normierte Algebra über $\mathbb{R}$ mit Eins, so ist die kartesische Involution von $(A_{(\mathbb{C})}, \|\cdot\|_p)$ und von $(A_{(\mathbb{C})}, \|\cdot\|_M)$ isometrisch, also eine Konjugation. Somit sind der Raum $(X_{(\mathbb{C})}, \|\cdot\|_p)$ und die Algebren $(A_{(\mathbb{C})}, \|\cdot\|_p)$ und $(A_{(\mathbb{C})}, \|\cdot\|_M)$ Komplexifizierungen nach LI. Anstelle von der kartesischen Involution spricht man dann auch von der kartesischen Konjugation. Man beachte 3.2.16.

3.1.15 Prädual und Komplexifizierung (LI [198](2003), Seite 4, Satz 1.1.5). Sei X ein Banachraum über $\mathbb{R}$. Sei $X_{(\mathbb{C})}$ eine Komplexifizierung nach LI. Liege der Fall vor, dass $X_{(\mathbb{C})}$ ein Prädual besitze; es sei mit $X_{(\mathbb{C})}^*$ bezeichnet. Des Weiteren sei die kartesische Involution $\bar{}$ von $X = \sigma(X_{(\mathbb{C})}, X_{(\mathbb{C})}^*)$ -stetig. Dann gilt:

(a) X ist $\sigma(X_{(\mathbb{C})}, X_{(\mathbb{C})}^*)$ -abgeschlossen in $X_{(\mathbb{C})}$.

(b) Es sei an 3.1.6 erinnert. $X_* := X_{(\mathbb{C})}^* \mathbb{R}(\bar{})$ ist ein abgeschlossener Untervektorraum von $X_{(\mathbb{C})}^* \mathbb{R}$, der ein Prädual von X ist und die Eigenschaft hat, dass $X_{(\mathbb{C})}^*$ eine Komplexifizierung nach LI von ihm ist. Kurz: $X_{(\mathbb{C})}^* = X_*(\mathbb{C})$.

3.1.16 Satz. Sei X , $X \neq \{0\}$, ein Vektorraum über $\mathbb{R}$ mit einer Involution $*$. Dann ist

$$*: (x, y) \mapsto (x^*, -y^*) \quad \text{für alle } x, y \in A$$

eine Involution auf dem Vektorraum $X_{(\mathbb{C})}$, aber keine Involution auf dem Vektorraum $X_{(\mathbb{C})} \mathbb{R}$. Die Abbildung

$$*: (x, y) \mapsto (x^*, y^*) \quad \text{für alle } x, y \in A$$

ist dagegen zwar keine Involution auf dem Vektorraum $X_{(\mathbb{C})}$, wohl aber eine auf dem Vektorraum $X_{(\mathbb{C})} \mathbb{R}$.

Beweis. Beide Abbildungen sind additiv: $((a, b) + (c, d))^* = (a + c, b + d)^* = (a^* + c^*, \pm b^* \pm d^*) = (a^*, \pm b^*) + (c^*, \pm d^*) = (a, b)^* + (c, d)^*$. Aber $((k + im)(a, b))^* = (ka - mb, kb + ma)^* = (ka^* - mb^*, \pm kb^* \pm ma^*)$ und $\overline{(k + im)(a, b)^*} = (k - im)(a^*, \pm b^*) = (ka^* - (-m)(\pm b^*), \pm kb^* + (-m)a^*) = (ka^* \pm mb^*, \pm kb^* - ma^*)$ sind im Fall von $\mathbb{K} = \mathbb{C}$ nur dann identisch, wenn man für das $\pm$ -Zeichen das Minus-Zeichen liest, es sich also um die erstgenannte Abbildung handelt. Im Fall von $\mathbb{K} = \mathbb{R}$ liegt die Identität beider Ausdrücke nur vor, wenn das $\pm$ -Zeichen als das Plus-Zeichen gelesen wird, es sich also um die zweitgenannte Abbildung handelt. $\square$

3.1.17. In 2.2.24 sind Definitionen und Aussagen für ein Sesquilinearsystem $(X, Y; B; Z)$ mit $Z = \mathbb{K}$ gemacht worden. Ist Z ein Vektorraum mit einer Involution $\bar{}$, dann kann man die Definitionen von 2.2.24 wörtlich für das Sesquilinearsystem $(X, Y; B; Z)$ formulieren und die Aussagen gelten entsprechend.

Sei $(X, X; B; Y)$ ein Bilinearsystem oder ein Sesquilinearsystem. Sei auf Y ein Positivitätskonzept $\text{Pos}(Y)$ erklärt. Dann heißt die Abbildung B *positiv*, wenn $B(x, x) \in \text{Pos}(Y)$ für alle $x \in X$ gilt; falls zusätzlich $B(x, x) \neq 0$ für alle $x \in X^\times$ gilt, so heißt die Abbildung B *positiv definit*. Die Abbildung B heißt ein *Y-wertiges Skalarprodukt*, falls B eine positiv definite, sesquilineare, hermitesche Abbildung ist. Über das System $(X, X; B; Y)$ ist gemäß 2.1.4 der Begriff der Orthogonalität auf X definiert.

3.2 Involutionen auf Ringen und Algebren

3.2.1 Definition. Eine *Involution* eines nichtassoziativen Ringes R ist ein Antiautomorphismus $*$: $R \rightarrow R$ mit $(r^*)^* = r$ für alle $r \in R$.

3.2.2 Bemerkung. Sei R ein nichtassoziativer Ring mit einer Involution $*$. Hat R eine Linkseins ℓ , dann ist wegen $a\ell^* = (\ell a^*)^* = a^{**}$ für alle $a \in R$ das Element ℓ^* eine Rechtseins in R . Somit hat R wegen 2.1.2 die Eins e mit $e = \ell = \ell^*$. Hat R eine Rechtseins $r \in R$, dann ist analog $e = r = r^*$ die Eins.

3.2.3. Ganz entsprechend wie in 3.1.4 werden für nichtassoziative Ringe R mit einer Involution $*$ die Begriffe $*$ -Teilmenge, selbstadjungiert und der selbstadjungierte Teil definiert. Während zwar ganz wie in 3.1.4 dann $R = R^*$ gilt, ist aber die Abbildung $x \mapsto -x^*$ für $R \neq \{0\}$ keine Involution auf R . Denn: $(xy)^{-*} = -(xy)^* = -y^*x^* = -y^{-*}x^{-*} \neq y^{-*}x^{-*}$.

3.2.4 Definition (BONSALL und DUNCAN [30](1973), Seite 64). Eine *Involution* (auch *Adjungierung* oder genauer eine $\mathbb{K}$ -Algebra-Involution) auf einer nichtassoziativen Algebra A über $\mathbb{K}$ ist eine Abbildung, die sowohl auf dem Vektorraum A als auch auf dem Ring A eine Involution ist. Eine (nichtassoziative) Algebra über $\mathbb{K}$ ausgestattet mit einer Involution heißt *involutive (nichtassoziative) Algebra* über $\mathbb{K}$ oder auch eine *(nichtassoziative) *-Algebra* über $\mathbb{K}$. Eine nichtassoziative *-Algebra über $\mathbb{K}$ heißt *symmetrisch*, wenn gilt:

$$-a^*a \text{ ist quasi-invertierbar} \quad \text{für alle } a \in A.$$

3.2.5 Lemma. Sei R ein Ring mit Operatorenbereich K , $K \in \{\mathbb{Z}, \mathbb{R}, \mathbb{C}\}$, versehen mit einer Involution $*$. Dann gilt für jedes $p \in \text{Idem}(R)$ die Gleichungskette

$$\{x \in R : px^*p = x\} = R_{(*)} \cap pRp = (pRp)_{(*)} = pR_{(*)}p. \quad (3.1)$$

Beweis. Es sei an 2.1.33 erinnert. Sei $p \in \text{Idem}(R)$. Setze $M := \{x \in R : px^*p = x\}$. Sei $x \in M$. Dann gilt $x = px^*p = ppx^*pp = pxp$ und $x = px^*p = (pxp)^* = x^*$. Also $M \subseteq R_{(*)} \cap pRp$. Ist umgekehrt $x \in R_{(*)} \cap pRp$, dann also $x = pyp$ für ein $y \in R$ und $x = x^*$; also $px^*p = pxp = ppypp = pyp = x$. Somit $M = R_{(*)} \cap pRp$.

Da pRp sowohl ein Unterring mit Operatorenbereich K als auch eine *-Teilmenge von R ist, ist $*$ auf pRp eine Involution; wird diese wieder mit $*$ bezeichnet, gilt also $M = (pRp)_{(*)}$.

Sei $x \in M$. Da, wie bereits oben bemerkt, dann $x = px^*p = x^*$ gilt, hat man $M \subseteq pR_{(*)}p$. Sei umgekehrt $x \in pR_{(*)}p$. Dann ist $x = pyp$ für ein $y \in R$ mit $y = y^*$. Also $px^*p = ppy^*pp = pyp = x$. Somit gilt $M = pR_{(*)}p$. $\square$

3.2.6. Hat die Algebra A in der Situation von 3.2.4 eine Eins, so ist die Algebra A gemäß Bemerkung 2.1.20(c) genau dann symmetrisch, wenn $e + a^*a$ für alle $a \in A$ invertierbar ist.

3.2.7 Bemerkung (RICKART [246](1960), Seite 178). Sei A eine Algebra über $\mathbb{K}$. Dann ist Id_A genau dann eine Involution auf A , wenn $\mathbb{K} = \mathbb{R}$ und A kommutativ ist.

3.2.8 Bemerkung und Definition (INGELSTAM [145](1964), Seite 264 und RICKART [246](1960), Seite 179). Sei A eine *-Algebra über $\mathbb{K}$. Dann ist auf A^1 (siehe Definition 2.3.19) eine Involution per

$$(a, z)^* := (a^*, \bar{z}) \quad \text{für alle } a \in A, z \in \mathbb{K}$$

definiert und die so erhaltene *-Algebra über $\mathbb{K}$ wird wieder mit A^1 bezeichnet. Wenn nicht ausdrücklich etwas anderes vereinbart ist, sei A^1 immer diese *-Algebra über $\mathbb{K}$. A ist eine *-Unteralgebra von A^1 .

3.2.9 Bemerkung und Definition (RICKART [246](1960), Seite 179). Sei A eine *-Algebra über $\mathbb{R}$. Dann ist die Abbildung

$$*: (x, y) \mapsto (x^*, -y^*) \quad \text{für alle } x, y \in A$$

eine Involution auf der Algebra $A_{(\mathbb{C})}$. Mit dieser Involution ist $A_{(\mathbb{C})}$ also eine *-Algebra über $\mathbb{C}$. Wenn nicht ausdrücklich etwas anderes vereinbart ist, dann sei mit $A_{(\mathbb{C})}$ immer diese *-Algebra gemeint. Diese Involution wird auch als die *kanonische Involution der Komplexifizierung einer *-Algebra über $\mathbb{R}$* angesprochen. Wenn die Involution von A auf A stetig ist, dann ist auch $*$ auf $A_{(\mathbb{C})}$ stetig. A ist eine *-Unteralgebra von $A_{(\mathbb{C})\mathbb{R}}$. Siehe auch 3.2.22.

Die Abbildung $(x, y) \mapsto (x^*, y^*)$ ist auf dem Ring $A_{(\mathbb{C})}$ und auf der Algebra $A_{(\mathbb{C})\mathbb{R}}$ eine Involution.

Beweis. Satz 3.1.16 und die Gleichungen $((a, b)(c, d))^* = (ac - bd, bc + ad)^* = (c^*a^* - d^*b^*, \pm c^*b^* \pm d^*a^*) = (c^*, \pm d^*)(a^*, \pm b^*) = (c, d)^*(a, b)^*$ zeigen für $\pm \equiv -$, dass $(x, y) \mapsto (x^*, -y^*)$ eine Algebra-Involution auf $A_{(\mathbb{C})}$ ist. Die Aussage für die Abbildung $(x, y) \mapsto (x^*, y^*)$ erhält man mit $\pm \equiv +$. $\square$

3.2.10 (RICKART [246](1960), Seite 179; NAIMARK [221](1972), Seite 184). Sei A eine $*$ -Algebra über $\mathbb{K}$. Dann transformiert die Involution $*$ jedes Linksideal in ein Rechtsideal und vice versa; dabei bleiben die Eigenschaften modular und maximal erhalten. Ist ein ein- oder beidseitiges Ideal von A eine $*$ -Teilmenge, so ist es automatisch ein beidseitiges Ideal von A ; es heißt dann ein $*$ -Ideal von A . Es folgt, dass das Jacobson-Radikal von A ein $*$ -Ideal von A ist.

3.2.11 Definition. Sei A eine $*$ -Algebra über $\mathbb{K}$. Dann heißt ein Element $a \in A$

- *normal*, wenn $a^*a = aa^*$;
- eine *Projektion*, wenn $a \in A_{(*)}$ und a idempotent ist;
- *unitär*, wenn A eine Eins besitzt und dann $a^*a = aa^* = e$ gilt.

Die Menge aller normalen Elemente (bzw. Projektionen, bzw. unitären Elemente) von A wird mit A_{normal} (bzw. $\text{Proj}(A)$, bzw. $A_{\text{unitär}}$) bezeichnet.

3.2.12 Bemerkungen. (a) (PEDERSEN [234](1989), Seite 171). Man beachte, dass $\mathcal{L}(H)_{\text{normal}}$, H ein Hilbertraum über $\mathbb{K}$, im Allgemeinen kein Vektorraum ist.

(b) Sei A eine $*$ -Algebra über $\mathbb{K}$ mit Eins. Gemäß Bemerkung 3.2.2 ist die Eins eine Projektion.

(c) Ist A eine $*$ -Algebra über $\mathbb{K}$ mit Eins, dann ist $A_{\text{unitär}}$ eine $*$ -Untergruppe (soll heißen, eine $*$ -Teilmenge und Untergruppe) von der Gruppe $G(A)$ aller Inversen von A .

3.2.13 Definition. Sei A eine $*$ -Algebra über $\mathbb{K}$. Dann gibt es folgendes Konzept von Positivität:

$$\text{Pos}(*\sigma; A) := \{x \in A_{(*)} : \sigma(x) \subseteq \mathbb{R}^+\}.$$

3.2.14 Definition. Sei A eine Algebra über $\mathbb{C}$. Dann ist das *reelle Jordanprodukt* definiert als die Abbildung $A \rightarrow A$, $(r, s) \mapsto \frac{1}{2}(rs + sr)$ und das *komplexe Jordanprodukt* definiert als die Abbildung $A \rightarrow A$, $(r, s) \mapsto \frac{1}{2i}(rs - sr)$.

3.2.15. Sei A eine $*$ -Algebra über $\mathbb{K}$. Dann ist im Allgemeinen $A_{(*)}$ keine $*$ -Algebra. $A_{(*)}$ ist aber immer abgeschlossen bezüglich des reellen und komplexen Jordanproduktes.

Insbesondere ist also die reelle Jordan-Algebra $(A_{(*)})^{\oplus}$ erklärt; siehe 2.3.16.

3.2.16 (LI [198](2003), Seite 78). Sei X ein Vektorraum über $\mathbb{R}$ und liege der Fall vor, dass $X_{(\mathbb{C})}$ eine Algebra ist. Dann ist die kartesische Involution $\bar{}$ des Vektorraumes $X_{(\mathbb{C})}$ zwar ein Automorphismus der Algebra $X_{(\mathbb{C})}$ (das heißt, $\overline{\overline{x}y} = \overline{x} \overline{y}$ für alle $x, y \in X_{(\mathbb{C})}$), aber im Allgemeinen kein Antiautomorphismus, also keine Involution auf der Algebra $X_{(\mathbb{C})}$. Im Allgemeinen bleibt somit die Definition der Komplexifizierung nach LI auf den Vektorraum-Fall beschränkt. Das heißt, wenn im Folgenden von einer Komplexifizierung nach LI einer Algebra A gesprochen wird, so bezieht sich das weiterhin auf A betrachtet als Vektorraum. Immerhin liegt folgender Satz vor (siehe auch Satz 3.5.29):

3.2.17 Satz (BONSALL und DUNCAN [30](1973), Seite 64). *Ist X ein Vektorraum über $\mathbb{R}$, so dass $X_{(\mathbb{C})}$ eine Algebra über $\mathbb{C}$ ist, und ist dann X als Unterraum über $\mathbb{R}$ von $X_{(\mathbb{C})\mathbb{R}}$ abgeschlossen bezüglich des reellen und komplexen Jordanproduktes, so ist die kartesische Involution des Vektorraumes $X_{(\mathbb{C})}$ eine Involution auf der Algebra $X_{(\mathbb{C})}$.*

3.2.18 Definition (BONSALL und DUNCAN [30](1973), Seite 64; ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seite 316).

- (a) Eine *normierte *-Algebra* ist eine normierte Algebra A mit einer Involution auf der Algebra A .
- (b) Eine **-normierte Algebra über $\mathbb{K}$* ist eine normierte Algebra über $\mathbb{K}$ mit einer Konjugation auf A — mit anderen Worten also eine normierte *-Algebra, deren Involution eine Isometrie ist.
- (c) Eine *involutive Banachalgebra* (oder auch *Banach-*-Algebra*) über $\mathbb{K}$ ist eine normierte *-Algebra über $\mathbb{K}$, deren Norm vollständig ist.
- (d) Eine **-normierte involutive Banachalgebra* (oder einfach **-Banachalgebra*) über $\mathbb{K}$ ist eine *-normierte Algebra über $\mathbb{K}$, deren Norm vollständig ist.

3.2.19 Bemerkungen (PALMER [232](2001), Seite 799). Man beachte, dass, außer wenn die Involution einer normierten *-Algebra stetig ist, die Vervollständigung einer normierten *-Algebra keine kanonische Involution aufweisen muss. Genauso muss, solange die Involution nicht stetig ist, der Abschluss einer *-Teilmenge nicht wieder eine *-Teilmenge sein.

Für eine Banachalgebra X über $\mathbb{K}$ mit einer stetigen Involution ist der Unterraum $X_{(*)}$ über $\mathbb{R}$ abgeschlossen in $X_{\mathbb{R}}$, aber im Allgemeinen nicht abgeschlossen bezüglich der Algebrenmultiplikation. (UPMEIER [279](1985), Seite 234.)

Jede Banach-*-Algebra über $\mathbb{C}$ mit Eins ist gleich der linearen Hülle ihrer unitären Elemente. (BONSALL und DUNCAN [30], Seite 66.) Dies gilt im Allgemeinen nicht für Banach-*-Algebren über $\mathbb{R}$, siehe LI [198](2003), Seite 38. (Siehe auch Satz 3.7.45(b).)

In jeder Banach-*-Algebra A über $\mathbb{C}$ mit Eins, deren Involution $*$ stetig ist, gilt $\exp(iA_{(*)}) \subseteq A_{\text{unitär}}$; dabei heißen die Elemente von $\exp(iA_{(*)})$ die *exponentiell unitären Elemente* von A ; siehe auch PALMER [232](2001), Seite 1174. Es sei daran erinnert, dass $iA_{(*)}$ bei $\mathbb{K} = \mathbb{C}$ gemäß Bemerkung 3.1.5(b) der schiefselbstadjungierte Teil von A ist.

3.2.20 Bemerkung. Ist A eine normierte *-Algebra über $\mathbb{K}$ ohne Eins, dann ist A^1 mit der Norm von Definition 2.2.18 eine normierte *-Algebra über $\mathbb{K}$, deren Involution stetig ist, wenn die Involution von A stetig ist. Insbesondere ist dieses A^1 eine *-Banachalgebra über $\mathbb{K}$, wenn A eine *-Banachalgebra über $\mathbb{K}$ ist.

3.2.21 Bemerkung. Sei A eine normierte *-Algebra über $\mathbb{K}$ ohne Eins. Dann kann wegen Bemerkung 3.2.2 die in Bemerkung 2.3.23 erwähnte Schwierigkeit nicht auftreten. Daher ist

$$\|(a, \lambda)\| := \sup\{\|ab + \lambda b\| : b \in B_A\} \quad \text{für alle } a \in A, \lambda \in \mathbb{K}$$

eine Norm auf A^1 mit der A^1 eine normierte *-Algebra über $\mathbb{K}$ ist.

3.2.22 Bemerkung (BONSALL und DUNCAN [29](1971), Seite 50). Sei $(A, \|\cdot\|)$ eine unitale Banachalgebra über $\mathbb{C}$. Es sei an Satz 2.11.23 erinnert. Sei X der mit der Norm von A versehene Banachraum $A_{\text{herm}} + iA_{\text{herm}}$ über $\mathbb{C}$. Dann ist die kartesische Involution auf X , also die Abbildung

$$*: X \rightarrow X, \quad h + ik \mapsto h - ik \quad \text{für alle } h, k \in A_{\text{herm}},$$

eine stetige Involution auf X ; sie wird in dieser Allgemeinheit im Folgenden als die *Vidav-Involution* bezeichnet. Im Fall dass $A_{\text{herm}, \mathbb{R}}$ eine kommutative Algebra ist, ist nach Bemerkung 3.2.7 die identische Abbildung eine Involution auf $A_{\text{herm}, \mathbb{R}}$ und die Vidav-Involution somit genau die kanonische Involution von $A_{\text{herm}, \mathbb{R}(\mathbb{C})}$, siehe 3.2.9.

3.2.23 (RICKART [246](1960), Seite 212). Sei A eine $*$ -Algebra über $\mathbb{C}$. Sei $\ell \in A^*$ ein selbstadjungiertes Funktional. Dann ist $[a, b]_\ell := \ell(b^*a)$, $a, b \in A$, eine hermitesche Sesquilinearform auf A , der aber im Allgemeinen noch die positive Definitheit fehlt, um ein Skalarprodukt zu sein.

$\mathfrak{N}_\ell := \{a \in A : [a, b]_\ell = 0 \text{ für alle } b \in A\}$ ist ein Linksideal in A und auf $A_\ell := A/\mathfrak{N}_\ell$ mit $[a] := a + \mathfrak{N}_\ell$ ist dann per $\langle [a], [b] \rangle := \ell(b^*a)$ ein Skalarprodukt definiert, falls $\ell(a^*a) \geq 0$ für alle $a \in A$. So motiviert, gelangt man zu einem weiteren Konzept von Positivität und zu einer Definition von positiven Funktionalen:

3.2.24 Definition. (a) (PALMER [231](1994)). Sei A eine $*$ -Algebra über $\mathbb{K}$. Dann setze:

$$\text{Pos}(*; A) := \left\{ \sum_{j=1}^n a_j^* a_j : a_1, \dots, a_n \in A, n \in \mathbb{N} \right\}.$$

(b) Ein $\ell \in A^\#$ heißt *positiv*, in Zeichen $\ell \geq 0$, wenn für alle $a \in \text{Pos}(*; A)$ die Ungleichung $\ell(a) \geq 0$ gilt.

(c) Sei A eine normierte $*$ -Algebra über $\mathbb{K}$. Ein *Zustand* (bezüglich der Involution $*$) von A ist ein selbstadjungiertes positives Funktional ℓ aus S_{A^*} .

(d) Sei A eine $*$ -Algebra über $\mathbb{K}$. Ein positives $\ell \in A^\#$ heißt eine *Spur* auf A , wenn für alle $x, y \in A$ die Gleichung $\ell(xy) = \ell(yx)$ gilt. (DIEUDONNÉ [68](1987), Band 2, Abschnitt 15.7.)

(e) Ein positives $\ell \in A^\#$ heißt *treu* (im Englischen *faithful*), wenn 0 das einzige Element x aus $\text{Pos}(*; A)$ mit $\ell(x) = 0$ ist.

3.2.25 Bemerkungen. (a) Sei A eine $*$ -Algebra über $\mathbb{C}$. Dann ist $\text{Pos}(*; A)$ ein punktierter spitzer konvexer Kegel im reellen Vektorraum $A_{(*)}$ und somit ist $(A_{(*)}, \leq_*)$ ein geordneter Vektorraum über $\mathbb{R}$.

(b) Die Selbstadjungiertheit in der Definition 3.2.24(c) wird gefordert, damit mit den Zuständen die Skalarprodukte gemäß 3.2.23 gebildet werden können. (Bezüglich weiterer Forderungen an ein Funktional in allgemeineren Situationen — wie zum Beispiel admissible, representable — siehe RICKART [246](1960), Seite 213.) Insbesondere ist die Selbstadjungiertheit bei C^* -Algebren über $\mathbb{R}$ von Bedeutung; siehe Satz 3.5.39.

(c) Ist A eine $*$ -Algebra über $\mathbb{K}$, dann ist ein $\ell \in A^\#$ genau dann positiv, wenn $\ell(a^*a) \geq 0$ für alle $a \in A$ gilt.

(d) Sei A eine $*$ -Algebra über $\mathbb{K}$. Man bemerke, dass jedes positive Funktional $\ell \in A^\#$ für alle $x \in A$ die Gleichung $\ell(x^*) = \ell(x)$ erfüllt; siehe 3.1.5(f).

(e) Die in Definition 3.2.24 definierten Funktionalen beziehen sich auf das Konzept $\text{Pos}(*; A)$. Positive Funktionalen — und damit auch Spuren und treue Funktionalen — kann man natürlich auch bezüglich jeden anderen Positivitäts-konzeptes definieren, wobei A auch nicht zwingend eine $*$ -Algebra sein muss, sondern zum Beispiel ein prägeordneter Vektorraum sein kann. Für ein Beispiel siehe 4.7.15.

3.2.26 Satz. In jeder normierten $*$ -Algebra A über $\mathbb{K}$ gilt für alle $x \in A$:

$$(a) \|x^*x\| = \|x\|^2 \Rightarrow (\|x\| = \|x^*\| \text{ und } \|xx^*\| = \|x\|^2).$$

$$(b) \|xx^*\| = \|x\|^2 \Rightarrow (\|x\| = \|x^*\| \text{ und } \|x^*x\| = \|x\|^2).$$

Beweis. Zur Isometrie in (a): $\|x\|^2 = \|x^*x\| \leq \|x^*\|\|x\| \Rightarrow \|x\| \leq \|x^*\|$. Analog zeigt man $\|x^*\| \leq \|x\|$. Die Isometrie in (b) zeigt man entsprechend. $\square$

3.2.27 Satz (DALES [56](2000), Seite 147). Sei A eine $*$ -Algebra über $\mathbb{C}$ und ℓ ein positives Funktional aus $A^\#$. Dann ist ℓ eingeschränkt auf A^2 selbstadjungiert. Demnach ist, falls A eine Eins hat, ℓ selbstadjungiert.

3.2.28 Satz. Sei A eine $*$ -Algebra über $\mathbb{K}$ und ℓ ein positives Funktional aus $A^\#$, welches eingeschränkt auf A^2 selbstadjungiert ist. Dann gilt:

(a) ℓ erfüllt auf A^2 die Cauchy-Schwartz-Bunyakovsky'sche Ungleichung:

$$|\ell(a^*b)|^2 \leq \ell(a^*a)\ell(b^*b) \quad \text{für alle } a, b \in A.$$

Falls A eine Eins hat, dann gilt zusätzlich:

(b) $\ell(A_{(*)}) \subseteq \mathbb{R}$.

(c) $|\ell(a)|^2 \leq \ell(e)\ell(a^*a)$ für alle $a \in A$.

Beweis. (a) Siehe RICKART [246](1960), Seiten 212 und 213.

(b) $\ell(e^*a) = \ell(a^*e)$ nach Voraussetzung. Also $\ell(a) = \overline{\ell(a^*)}$.

(c) Folgt direkt aus (a).

Für (b) und (c) siehe auch NAIMARK [221](1972), Seite 186. $\square$

3.2.29 Satz. Sei A eine $*$ -Banachalgebra über $\mathbb{K}$ mit Eins. Sei $\ell \in A^\#$ mit $\ell \geq 0$. Dann gilt:

$$|\ell(a)| \leq \ell(e)\|a\| \quad \text{für alle } a \in A_{(*)}.$$

Diese Ungleichung gilt für alle $a \in A$, falls die Ungleichung (c) vom Satz 3.2.28 gilt und in diesem Fall hat man $\|\ell\| = \ell(e)$.

Beweis. NAIMARK [221](1972), Seiten 188 und 189. $\square$

3.2.30 Bemerkung. Man beachte, dass der Satz 3.2.29 im Allgemeinen nicht für Banach- $*$ -Algebren gilt. Siehe etwa RUDIN [251](1973), Seite 289, Übung 13.

3.2.31 Satz. Sei A eine symmetrische Banach- $*$ -Algebra über $\mathbb{K}$ ohne Eins, deren Involution im reellen Fall $\mathbb{K} = \mathbb{R}$ zusätzlich noch als stetig vorausgesetzt wird. Dann ist auch A^1 symmetrisch.

Beweis. Die Beweisidee stammt von YOOD. Für $\mathbb{K} = \mathbb{R}$ siehe INGELSTAM, [145] (1964), Seite 264. Für $\mathbb{K} = \mathbb{C}$ siehe RICKART [246](1960), Seite 233. $\square$

3.2.32 Satz. Sei A eine $*$ -Algebra über $\mathbb{K}$. Dann sind äquivalent:

(a) A ist symmetrisch.

(b) $\sigma(a^*a) \subseteq \mathbb{R}^+$ für alle $a \in A$.

(c) $a^*a \in \text{Pos}(*\sigma; A)$ für alle $a \in A$. (Definition 3.2.13)

Beweis. RICKART [246](1960), Seite 233, Theorem 4.7.6, zeigt (a) $\Rightarrow$ (b) per Äquivalenzen. Siehe auch PALMER [230](1970), Seite 197. (c) ist Definition 3.2.13. $\square$

3.2.33 Definition (DORAN und BELFI [77](1986), Seite 128; LI [198](2003), Seite 41). Sei A eine $*$ -Algebra über $\mathbb{K}$. Die Involution von A heißt *hermitesch*, wenn $\sigma(s) \subseteq \mathbb{R}$ für alle $s \in A_{(*)}$ gilt und in diesem Fall heißt A eine *hermitesche Algebra*. A heißt *schiefhermitesch*, wenn $\sigma(s) \subseteq i\mathbb{R}$ für alle $s \in iA_{(*)}$ gilt (vgl. 3.1.4 und 3.1.5).

3.2.34 Satz (PALMER [232](2001), Seite 1010). Sei A eine $*$ -Algebra über $\mathbb{C}$. Dann sind äquivalent:

(a) Für alle $a \in A$ gilt: $\sigma(a^*a) \subseteq \mathbb{R}^+$.

(b) Für alle $a \in A^1$ gilt: $\sigma(a^*a) \subseteq \mathbb{R}^+$.

(c) A ist symmetrisch.

Diese äquivalenten Bedingungen implizieren, dass gilt:

(d) A ist hermitesch.

3.2.35 Satz (DORAN und BELFI [77](1986), Seiten 136 und 142). Sei A eine Banach- $*$ -Algebra über $\mathbb{C}$. Dann sind äquivalent:

(a) A ist symmetrisch.

(b) A ist hermitesch.

(c) $\rho(a)^2 \leq \rho(a^*a)$ für alle $a \in A$.

3.2.36. Die Äquivalenz von (a) und (b) im Satz 3.2.35 stammen von SHIRALI und FORD. Die Äquivalenz zu (c) stammt von PTÁK.

3.2.37 Satz (LI [198](2003), Seite 54). Sei A eine Banach- $*$ -Algebra über $\mathbb{R}$. Dann sind äquivalent:

(a) A ist symmetrisch.

(b) A ist hermitesch und schiefhermitesch.

(c) $-\lambda \notin \sigma(a^*a)$ für alle $\lambda \in \mathbb{R}^{+\times}$, $a \in A^1$.

3.2.38 (TOPPING [278](1979), Seite 11). Ist A eine $*$ -Algebra über $\mathbb{K}$, so ist für eine beliebige Teilmenge S von A die Kommutante S' zwar nach Satz 2.3.26(a) eine Algebra, aber im Allgemeinen keine $*$ -Teilmenge von A , also insbesondere keine $*$ -Algebra.

Beweis. Man betrachte dazu zum Beispiel die Algebra $M_{\mathbb{K}}(2)$ aller 2×2 -Matrizen über $\mathbb{K}$ und ihr Element $A := \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}$. Damit $A \begin{pmatrix} a & b \\ c & d \end{pmatrix} - \begin{pmatrix} a & b \\ c & d \end{pmatrix} A = \begin{pmatrix} 0 & 0 \\ a & b \end{pmatrix} - \begin{pmatrix} b & 0 \\ d & 0 \end{pmatrix} = \begin{pmatrix} -b & 0 \\ a-d & b \end{pmatrix}$ gleich Null ist, muss $b = 0$, $a = d$, $c \in \mathbb{K}$ gelten. Also $\{A\}' = \left\{ \begin{pmatrix} u & 0 \\ v & u \end{pmatrix} \in M_{\mathbb{K}}(2) : u, v \in \mathbb{K} \right\}$. Dann ist aber für alle $u \in \mathbb{K}$, $v \in \mathbb{K}^\times$ $\begin{pmatrix} u & 0 \\ v & u \end{pmatrix}^* = \begin{pmatrix} \bar{u} & \bar{v} \\ 0 & \bar{u} \end{pmatrix} \notin \{A\}'$, das heißt, $\{A\}'$ ist keine $*$ -Teilmenge von $M_{\mathbb{K}}(2)$. $\square$

Aber es gilt:

3.2.39 Satz (SUNDER [271](1987), Seite 12). Sei A eine $*$ -Algebra über $\mathbb{K}$. Ist S eine $*$ -Teilmenge von A , so ist auch ihre Kommutante S' eine $*$ -Teilmenge von A , insbesondere also wegen Satz 2.3.26(a) eine $*$ -Unteralgebra von A .

Beweis. Sei $x \in S'$. Dann gilt für alle $s \in S$: $xs = sx$, also $s^*x^* = x^*s^*$. Da nach einer Bemerkung in 3.1.4 $S = S^*$ mit $S^* := \{s^* \in S : s \in S\}$ gilt, hat man hier somit für alle $s \in S$: $sx^* = x^*s$, das heißt, $x^* \in S'$. $\square$

3.3 Cayley-Dickson-Algebren

3.3.1 Quaternionen (KLINGENBERG und KLEIN [186](1972), Seite 234).

(a) Betrachte $\mathbb{R}^4$ mit der kanonischen Basis (e_1, e_2, e_3, e_4) . Dann ist $\mathbb{R}^4$ genau mit der Multiplikation, die durch

e_1 ist ein neutrales Element und

$$e_2 e_2 := -e_1, \quad e_3 e_3 := -e_1, \quad e_2 e_3 := -e_3 e_2 := e_4$$

erklärt ist, eine nichtkommutative Divisionsalgebra über $\mathbb{R}$ mit e_1 als Eins. Sie heißt die *Hamilton'sche Algebra der Quaternionen* (über $\mathbb{R}$) oder auch die *Divisionsalgebra der reellen Quaternionen* und wird mit $\mathbb{H}$ bezeichnet.

Für $x = \sum_{i=1}^4 \alpha_i e_i$ und $y = \sum_{i=1}^4 \beta_i e_i$ gilt dann $xy = \sum_{i,j=1}^4 \alpha_i \beta_j e_i e_j$ für alle $\alpha_i, \beta_i \in \mathbb{R}, i = 1, \dots, 4$.

$\mathbb{R}$ ist per $x \mapsto (x, 0, 0, 0)$ in $\mathbb{H}$ eingebettet. Die Abbildung $\Psi: \mathbb{C}_{\mathbb{R}} \rightarrow \mathbb{H}, x + iy \mapsto (x, y, 0, 0)$ ist ein Schiefkörpermonomorphismus. Daher ist $\mathbb{H}_{\Psi}$ nach 2.9.1 ein zweidimensionaler Vektorraum über $\mathbb{C}$ mit $\mathbb{H} = \mathbb{H}_{\Psi, \mathbb{R}}$. Man beachte, dass in 2.9.1 die skalare Multiplikation als Linksmultiplikation definiert ist. Man bekommt eine andere Vektorraumstruktur über $\mathbb{C}$, wenn man die Skalarmultiplikation als Rechtsmultiplikation definiert.

(b) (SCHULZE [258](2005)). Bezeichne $M_{\mathbb{K}}(n)$ die Algebra über $\mathbb{K}$ aller $n \times n$ -Matrizen mit Einträgen aus $\mathbb{K}$. Für $u, v \in \mathbb{C}$ bezeichne $M(u, v)$ das Element aus $M_{\mathbb{C}}(2)$ der Form $\begin{pmatrix} u & v \\ -\bar{v} & \bar{u} \end{pmatrix}$. Sei $\tilde{\mathbb{H}} := \{M(u, v) : u, v \in \mathbb{C}\}$. Dann ist $\tilde{\mathbb{H}}$ eine Unteralgebra über $\mathbb{R}$ mit Eins $M(1, 0)$ von der Algebra $(M_{\mathbb{C}}(2))_{\mathbb{R}}$ über $\mathbb{R}$.

Mit $E_1 := M(1, 0)$, $E_2 := M(i, 0)$, $E_3 := M(0, 1)$ und $E_4 := M(0, i)$ ist (E_1, E_2, E_3, E_4) eine Basis von $\tilde{\mathbb{H}}$.

(c) Die Abbildung $T: \mathbb{H} \rightarrow \tilde{\mathbb{H}}, e_i \mapsto E_i$ ist ein $\mathbb{R}$ -Algebra-Isomorphismus. Folglich kann man $\mathbb{H}$ mit $\tilde{\mathbb{H}}$ identifizieren.

Für jedes $a = (a_1, a_2, a_3, a_4) \in \mathbb{H}$ ist $\bar{a} := (a_1, -a_2, -a_3, -a_4) \in \mathbb{H}$ die zu a konjugierte Quaternion und $|a| := (a_1^2 + a_2^2 + a_3^2 + a_4^2)^{1/2}$ ist der Betrag der Quaternion a .

Wegen $T(\bar{a}) = (\overline{T(a)})^t$ für alle $a \in \mathbb{H}$ ist die zu $A \in \tilde{\mathbb{H}}$ konjugierte Quaternion die Transponierte der konjugiert-komplexierten Matrix A . Wegen $|a|^2 = \det T(a)$ ist das Quadrat des Betrages von $A \in \tilde{\mathbb{H}}$ definiert als die Determinante der Matrix A . Das zu $a \in \mathbb{H}$ inverse Element ist $\bar{a}/|a|^2$. Dementsprechend ist für $M(u, v) \neq 0$ die Inverse von $M(u, v)$ gleich $\frac{M(\bar{u}, -v)}{\det M(u, v)}$.

Die Bildung der konjugierten Quaternionen ist eine Involution auf $\mathbb{H}$. Es gilt $a = \bar{\bar{a}} \Leftrightarrow a \in \mathbb{R}$, und $a\bar{a} = |a|^2$ für alle $a \in \mathbb{H}$.

3.3.2 Normierte Divisionsalgebren. Von Frobenius stammt das klassische Theorem, dass jede endlichdimensionale Divisionsalgebra über $\mathbb{C}$ isomorph zu $\mathbb{C}$ ist (also eindimensional ist, womit übrigens $\mathbb{H}_{\Psi}$ von 3.3.1(a) keine Divisionsalgebra sein kann). Im Jahr 1938 hat Mazur dieses Theorem verallgemeinert: Jede normierte Divisionsalgebra A über $\mathbb{C}$ besteht nur aus komplex-skalaren Vielfachen der Eins von A und somit existiert ein eindeutig bestimmter Isomorphismus (bezüglich sowohl der Algebra als auch des normierten Vektorraumes) von A auf die komplexen Zahlen; ist insbesondere A unital, so ist dieser Isomorphismus eine Isometrie. Für einen Beweis siehe zum Beispiel BONSALL und DUNCAN [30](1973), Seite 71. Für Literaturhinweise siehe PALMER [231](1994), Seite 211, wo man auch die folgende etwas schärfere Version findet: Jede Divisionsalgebra A über $\mathbb{C}$, auf der eine nicht triviale Algebrennorm definiert werden kann, ist isomorph zu den komplexen Zahlen. Dabei genügt es bereits anzunehmen, dass jedes

von Null verschiedene Element der Algebra eine Linksinverse besitzt; die entsprechende Formulierung per Rechtsinverse ist ebenfalls hinreichend. Der Isomorphismus $A \rightarrow \mathbb{C}$ wird dabei implementiert durch $a \mapsto \lambda$, wobei jeweils λ die eine komplexe Zahl im Spektrum $\sigma(a)$ von a ist. Ebenfalls auf Mazur zurückgehend gilt (siehe BONSALL und DUNCAN [30](1973), Seite 73): Ist A eine normierte Divisionsalgebra über $\mathbb{R}$, so existiert ein Isomorphismus φ von A auf $\mathbb{R}$, $\mathbb{C}_{\mathbb{R}}$ oder $\mathbb{H}$; durch $\|x\|_{\varphi} := |\varphi(x)|$, $x \in A$, ist dann eine zu der gegebenen Algebranorm von A äquivalente Algebranorm auf A erklärt; des Weiteren gilt $\|xy\|_{\varphi} = \|x\|_{\varphi}\|y\|_{\varphi}$ für alle $x, y \in A$ und somit $\|x\|_{\varphi} = \rho(x)$ (Spektralradius, Definition 2.12.3) für alle $x \in A$. Als Folgerung der erstgenannten Aussage von Mazur hat man: Jede normierte Algebra A über $\mathbb{C}$ mit Eins, in der für alle $x, y \in A$ die Gleichung $\|xy\| = \|x\|\|y\|$ gilt, ist isometrisch isomorph zu $\mathbb{C}$. Für einen Beweis siehe zum Beispiel RICKART [246](1960), Seite 39, DORAN und BELFI [77], Seite 309, und LARSEN [196](1973), Seiten 38-40.

3.3.3 (BONSALL und DUNCAN [30](1973), Seite 157; PALMER [231](1994), Seite 677; RICKART [246](1960), Seite 45).

Sei A eine normierte Algebra über $\mathbb{K}$ mit einem schiefminimal idempotenten Element $p \in A$. Dies ist nach 2.1.52 beispielsweise der Fall, wenn ein minimales einseitiges Ideal $\mathfrak{m}$ von A mit $\mathfrak{m}^2 \neq \{0\}$ existiert. Nach einem Theorem von Mazur, siehe 3.3.2, ist die normierte Algebra pAp isomorph zu $\mathbb{R}$, $\mathbb{C}_{\mathbb{R}}$ oder $\mathbb{H}$. Im Fall von $\mathbb{K} = \mathbb{C}$ gilt, ebenfalls nach einem Theorem von Mazur, siehe wieder 3.3.2, die sich auf die Algebren beziehende Gleichung

$$pAp = \mathbb{C}p$$

und in diesem Fall existiert also ein $\ell \in A^{\#}$ mit

$$pxp = \ell(x)p \quad \text{für alle } x \in A;$$

wegen $|\ell(x)|\|p\| = \|\ell(x)p\| = \|pxp\| \leq \|p\|^2\|x\|$ ist dieses lineare Funktional ℓ auch stetig, $\ell \in A^*$.

Da für jedes $\mathbb{L} \in \{\mathbb{R}, \mathbb{C}, \mathbb{H}\}$ und jedes Idempotent $p \in A^{\times}$ $\mathbb{L}p$ ein Schiefkörper ist, hat man:

(a) In jeder normierten Algebra A über $\mathbb{R}$ ist ein idempotentes Element $p \in A^{\times}$ genau dann schiefminimal idempotent, wenn für ein $\mathbb{L} \in \{\mathbb{R}, \mathbb{C}_{\mathbb{R}}, \mathbb{H}\}$ die normierte Algebra pAp isomorph zu $\mathbb{L}p$ ist.

(b) In jeder normierten Algebra A über $\mathbb{C}$ ist ein idempotentes Element $p \in A^{\times}$ genau dann schiefminimal idempotent, wenn die normierte Algebra pAp isomorph zu $\mathbb{C}p$ ist.

3.3.4 (JACOBSON [152](1974), Seite 418). Sei X ein Vektorraum über $\mathbb{K}$ mit einer nicht ausgearteten quadratischen Form $Q: X \rightarrow \mathbb{K}$. Man sagt, Q *erlaube Komposition*, wenn man auf X eine bilineare Multiplikation $(x, y) \mapsto x \bullet y$ definieren kann, so dass gilt:

$$Q(x \bullet y) = Q(x)Q(y) \quad \text{für alle } x, y \in X.$$

Ist X mit diesem Produkt $\bullet$ eine nichtassoziative Algebra über $\mathbb{K}$ ohne Eins, dann kann man durch Modifikation der Multiplikation X eine Eins verschaffen: Sei u ein Element aus X mit $Q(u) \neq 0$. Seien L_u beziehungsweise R_u die durch u bestimmten Links- und Rechtsmultiplikationen in X . Dann ist X mit der Multiplikation

$$xy := (R_u^{-1}x) \bullet (L_u^{-1}y) \quad \text{für alle } x, y \in X$$

eine nichtassoziative Algebra über $\mathbb{K}$ mit u^2 als Eins.

3.3.5 Definition (JACOBSON [152](1974), Seite 419).

Eine *Kompositions-Algebra* ist eine nichtassoziative Algebra A über $\mathbb{K}$ mit einer Eins und einer nicht ausgearteten quadratischen Form $Q: A \rightarrow \mathbb{K}$, so dass für alle $x, y \in A$ die Gleichung $Q(xy) = Q(x)Q(y)$ gilt. Genauer schreibt man dann eine Kompositions-Algebra als (A, Q) .

3.3.6 Satz (JACOBSON [152](1974), Seite 422). *Sei A eine nichtassoziative Algebra über $\mathbb{K}$ mit Eins und einer nicht ausgearteten quadratischen Form $Q: A \rightarrow \mathbb{K}$. Dann sind äquivalent:*

- (a) (A, Q) ist eine Kompositions-Algebra.
- (b) A ist eine alternative Algebra über $\mathbb{K}$ mit einer Involution $*$ mit $x^*x = Q(x)e$ für alle $x \in A$.

3.3.7 Cayley-Dickson (MCCRIMMON [210](2004), Seite 64). Sei (A, q) eine endlich-dimensionale Kompositions-Algebra über $\mathbb{K}$. Dann existiert nach Satz 3.3.6 eine Involution $*$ mit $x^*x = q(x)e$ für alle $x \in A$. Nach dem sogenannten *Cayley-Dickson-Verdoppelungs-Prozess* kann man aus A eine nichtassoziative $*$ -Algebra konstruieren, deren Dimension das Doppelte der Dimension von A ist. Sei dazu $k \in \mathbb{K}^\times$. Dann ist die *Cayley-Dickson-Algebra* $CD(A, k)$ gleich dem Vektorraum A über $\mathbb{K}$ zusammen mit einer formalen Kopie des Vektorraumes A über $\mathbb{K}$, in Zeichen $A \oplus Am$, versehen mit der üblichen Vektorraumstruktur über $\mathbb{K}$

$$(u \oplus vm) + (x \oplus ym) := (u + x) \oplus (v + y)m \quad \text{und} \\ \lambda(x \oplus ym) := \lambda x \oplus \lambda ym \quad \text{für alle } u, v, x, y \in A, \lambda \in \mathbb{K}$$

und Produkt, Involution und einer nicht ausgearteten quadratischen Form $Q: A \rightarrow \mathbb{K}$ definiert durch:

$$(u \oplus vm)(x \oplus ym) := (ux + ky^*v) \oplus (yu + vx^*)m, \\ (x \oplus ym)^* := x^* \oplus -ym, \\ Q(x \oplus ym) := q(x) - kq(y) \quad \text{für alle } u, v, x, y \in A.$$

Insbesondere hat man somit

$$(e \oplus 0m)(x \oplus ym) = (x \oplus ym)(e \oplus 0m) = x \oplus ym \quad \text{für alle } x, y \in A$$

und des Weiteren

$$(0 \oplus em)(0 \oplus em) = (ke^*e) \oplus 0m = k(e \oplus 0m),$$

per entsprechenden Identifizierungen also $m^2 = ke$. Setzt man $A_0 := \mathbb{K}$ mit $*$ = Id_{A_0} und $q(x) = |x|^2$ für alle $x \in A_0$, so ist $A_1 := CD(A_0, k_1) = \mathbb{K} \oplus \mathbb{K}i_0$ mit $i_0^2 = k_1$, also $CD(\mathbb{R}, -1) = \mathbb{C}$. Weiteres iterieren liefert: $A_2 := CD(A_1, k_2) = A_1 \oplus A_1i_1$ mit $i_1^2 = k_2$, speziell $CD(\mathbb{C}, -1) = \mathbb{H}$. $A_3 := CD(A_2, k_3) = A_2 \oplus A_2i_2$ mit $i_2^2 = k_3$. $CD(\mathbb{H}, -1)$ heißt die *Algebra der Cayley'schen Oktonionen* und wird mit $\mathbb{O}$ bezeichnet. $\mathbb{O}$ ist eine nicht kommutative, nicht assoziative nichtassoziative Divisionsalgebra über $\mathbb{R}$.

Da $A_{n+1} := CD(A_n, k_{n+1})$, $n \geq 0$, $k_{n+1} \in \mathbb{K}^\times$ genau dann eine alternative Algebra über $\mathbb{K}$ ist, wenn A_n assoziativ ist, sind die A_n für $n \geq 4$ keine Kompositions-Algebren mehr.

Genau die $A_{n+1} := CD(A_n, k_{n+1})$ für $A_0 := \mathbb{K}$, für alle $k_n \in \mathbb{K}^\times$ und $n \in \{0, 1, 2\}$ stellen alle möglichen Kompositions-Algebren über $\mathbb{K}$ dar. (JACOBSON [152](1974), Seite 425 und MCCRIMMON, ebd.)

$\mathbb{R}$, $\mathbb{C}$, $\mathbb{H}$ und $\mathbb{O}$ sind die einzigen endlichdimensionalen, alternativen, reellen nichtassoziativen Divisionsalgebren (SCHAFER [256](1966), Seite 48).

3.4 Hilberträume

Hilberträume über $\mathbb{K}$ wurden in Definition 2.4.6 eingeführt.

3.4.1 Spur I (BOURBAKI [34](1974), II.4.3). Sei $n \in \mathbb{N}^\times$, X ein n -dimensionaler Vektorraum über $\mathbb{K}$ und $T \in L(X)$. Man betrachte das duale Paar $(X, X^\#)$. Dann lässt sich T auch schreiben als $\sum_{i=1}^n \langle \cdot, x_i^\# \rangle y_i$ mit $x_i^\# \in X^\#$, $y_i \in X$, $i = 1, \dots, n$. (Ist zum Beispiel $b_1, \dots, b_n$ eine Basis von X und ist $b_1^\#, \dots, b_n^\#$ die dazu duale Basis in $X^\#$, also $\langle b_i, b_j^\# \rangle = \delta_{ij}$ (Kronecker-Delta), $i, j = 1, \dots, n$, so wählt man $y_i := T(b_i)$ und $x_i^\# := b_i^\#$, $i = 1, \dots, n$.) Der Wert $\sum_{i=1}^n \langle y_i, x_i^\# \rangle$ heißt die *Spur* von T und ist unabhängig von der gewählten Darstellung von T . Dadurch ist eine lineare Abbildung $sp: L(X) \rightarrow \mathbb{K}$ definiert, die *Spur* oder auch *Spurabbildung* genannt wird.

Sei $B = \{b_1, \dots, b_n\}$ eine beliebige Basis von X . Sei $M = (m_{ij})$ die Matrix von T bezüglich B , das heißt, M ist definiert durch die Gleichungen $T(b_j) = \sum_{i=1}^n m_{ij} b_i$, $j = 1, \dots, n$. Dann ist die Spur von T gleich der Summe der Diagonalelemente von M , also gleich $\sum_{i=1}^n \langle T(b_i), b_i \rangle$, wenn man X als Hilbertraum auffasst, der mit seinem bezüglich der Basis B kanonisch definierten Skalarprodukt versehen ist. Somit ist die Spur von T gleich der Summe der in ihrer Vielfachheit aufgeführten Eigenwerte von $T_{(\mathbb{C})}$ (siehe 2.9.26).

Die Spurabbildung ist eine Spur im Sinne von Definition 3.2.24.

3.4.2. Sei H ein Hilbertraum über $\mathbb{K}$ und $\langle \cdot, \cdot \rangle$ sein Skalarprodukt. Betrachte das Sesquilinearsystem $(H, H; \langle \cdot, \cdot \rangle)$. Bezeichne gemäß 2.9.7 $\overline{H}$ den zu H konjugiert-komplexen Hilbertraum. Nach dem Fréchet-Riesz'schen Darstellungssatz existiert die Gradientenabbildung (siehe Definition 2.2.25) auf dem Dualraum H^* und ist eine semilineare Isometrie. Durch sie gilt: $H^* \cong \overline{H}$.

3.4.3 Definition. Seien H_1 und H_2 zwei Hilberträume über $\mathbb{K}$ und $\langle \cdot, \cdot \rangle_i$, $i = 1, 2$, ihre Skalarprodukte. Sei $T \in \mathcal{L}(H_1, H_2)$ und bezeichne $T' \in \mathcal{L}(H_2^*, H_1^*)$ die zu T duale Abbildung. Seien grad_i , $i = 1, 2$, die Gradientenabbildungen auf H_i^* , $i = 1, 2$. Dann ist die mit T^* bezeichnete *Hilbertraumadjungierte* von T definiert als das folgende Element aus $\mathcal{L}(H_2, H_1)$:

$$T^* := \text{grad}_1 \circ T' \circ \text{grad}_2^{-1}.$$

3.4.4 Bemerkung. Mit den Bezeichnungen der Definition 3.4.3 gilt:

$$T^*y = \text{grad}_1(x \mapsto \langle Tx, y \rangle_2) \equiv \text{grad}_1(\langle T \cdot, y \rangle_2) \quad \text{für alle } y \in H_2.$$

Daher ist

$$\langle Tx, y \rangle_2 = \langle x, T^*y \rangle_1 \quad \text{für alle } x \in H_1, y \in H_2$$

eine äquivalente Formulierung der Definition der Hilbertraumadjungierten T^* von T .

3.4.5 Orthonormalsysteme (WEIDMANN [288] (2000), Seiten 51, 52, 57; WERNER [291](2002), Kapitel V.4). Sei H ein Prähilbertraum über $\mathbb{K}$. Eine Teilmenge S von H heißt ein *Orthonormalsystem*, wenn für alle $x, y \in S$ die Gleichung $\langle x, y \rangle = \delta_{xy}$ (Kronecker-Delta) erfüllt ist; gilt dann auch noch zusätzlich die Gleichung $H = \text{cl span}(S)$, so heißt S eine *Orthonormalbasis* von H ; die Kardinalzahl einer Orthonormalbasis (die stets wohldefiniert ist), wird die *Hilbertraumdimension* von H genannt; wenn keine Verwechslung mit der in 2.1.39 erklärten Dimension von dimensionierten freien Linksmoduln zu befürchten ist, wird statt des Wortes Hilbertraumdimension auch einfach nur der Begriff Dimension verwendet.

Jedes Orthonormalsystem ist linear unabhängig.

Sei $S = \{x_\nu : \nu \in I\}$ ein Orthonormalsystem von H . Dann sind die folgenden vier Aussagen äquivalent:

(a) S ist eine Orthonormalbasis.

(b) $\forall x \in H : x = \sum_{\nu \in I} \langle x, x_\nu \rangle x_\nu$.

(c) $\forall x, y \in H : \langle x, y \rangle = \sum_{\nu \in I} \langle x, x_\nu \rangle \langle x_\nu, y \rangle$.

(d) $\forall x \in H : \|x\|^2 = \sum_{\nu \in I} |\langle x, x_\nu \rangle|^2$. (Parseval'sche Gleichung)

Ist S eine Orthonormalbasis, so gelten die beiden folgenden Aussagen:

(e) S ist bezüglich der durch Mengeninklusion definierten partiellen Ordnung ein maximales Orthonormalsystem.

(f) $(H, \text{span } S)$ ist ein in H trennendes Dualsystem.

Falls nun H ein Hilbertraum über $\mathbb{K}$ ist und S eine Teilmenge von H ist, die die Aussage (e) oder (f) erfüllt, so folgt auch umgekehrt, dass S eine Orthonormalbasis ist.

Keine Orthonormalbasis eines unendlichdimensionalen Hilbertraumes über $\mathbb{K}$ ist eine Basis, wie sie im Rahmen von freien Moduln in 2.1.39 definiert worden ist.

Zwei Hilberträume über den gleichen Körper $\mathbb{K}$ sind genau dann isometrisch isomorph, wenn sie die gleiche Hilbertraumdimension haben.

3.4.6 Satz von Hellinger-Toeplitz (WERNER [291](2002), Satz V.5.5). *Sei H ein Hilbertraum über $\mathbb{K}$ und $T \in L(H)$. Dann ist T genau dann stetig, wenn eine (notwendig lineare) Abbildung $R: H \rightarrow H$ existiert, so dass für alle $x, y \in H$ die Gleichung $\langle Tx, y \rangle = \langle x, Ry \rangle$ erfüllt ist.*

3.4.7. Für eine Verallgemeinerung des vorstehenden Satzes siehe EDWARDS [88] (1965), Abschnitt 8.11.

3.4.8 Bezeichnung. Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{K}$. Dann ist die Abbildung $*$: $\mathcal{L}(H) \rightarrow \mathcal{L}(H)$, die jedem $T \in \mathcal{L}(H)$ die Hilbertraumadjungierte $T^* \in \mathcal{L}(H)$ zuordnet, eine Konjugation auf $\mathcal{L}(H)$. Somit ist, wenn nicht ausdrücklich etwas anderes vereinbart ist, $\mathcal{L}(H)$ immer als $*$ -Banachalgebra über $\mathbb{K}$ aufzufassen. An dieser Stelle sei bereits auf Satz 3.5.26(c) hingewiesen.

3.4.9 (BONSALL und DUNCAN [29](1971), Seite 5; WEIDMANN [288](2000), Seite 106, Satz 2.5.1). Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{K}$. Es sei an Definition 2.11.15 erinnert. Dann gilt die Äquivalenz $(x, x^*) \in \Pi_H \Leftrightarrow (x \in S_H \text{ und } x^* = \langle \cdot, x \rangle)$. Folglich gilt — nebenbei sei an Satz 2.11.21(a) erinnert — für alle $T \in \mathcal{L}(H)$ die Gleichung

$$W(T) = V_{\text{spatial}}(T).$$

Setze $\gamma := \sup\{|\lambda| : \lambda \in W(T)\}$. Dann gilt:

(a) Ist $\mathbb{K} = \mathbb{C}$, so gilt $\gamma \leq \|T\| \leq 2\gamma$.

(b) Ist T selbstadjungiert, so gilt $\gamma = \|T\|$.

Siehe auch die Bemerkung 3.5.46 und für $\mathbb{K} = \mathbb{C}$ PALMER [232](2001), Seite 818.

3.4.10 Bemerkung und Definition (LI [198](2003), Seiten 8, 9). Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{R}$. Mit H , aufgefasst als Vektorraum über $\mathbb{R}$, ist der Vektorraum $H_{(\mathbb{C})}$ über $\mathbb{C}$, versehen mit dem Skalarprodukt

$$\langle (x, y), (v, w) \rangle := (\langle x, v \rangle + \langle y, w \rangle) + i(\langle y, v \rangle - \langle x, w \rangle)$$

für alle $(x, y), (v, w) \in H_{(\mathbb{C})}$, ein Hilbertraum über $\mathbb{C}$, der wieder mit $H_{(\mathbb{C})}$ bezeichnet wird. (Man beachte dazu, dass für alle $x, y \in H$ die Gleichung $\|(x, y)\|^2 = \|x\|^2 + \|y\|^2$ gilt.) Für Weiteres bezüglich $H_{(\mathbb{C})}$ siehe auch GOODEARL [117](1982), Seite 65.

Sei nun $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{C}$. Setze

$$\langle x, y \rangle_{\mathbb{R}} := \operatorname{Re} \langle x, y \rangle \quad \text{für alle } x, y \in H.$$

Dann ist $(H_{\mathbb{R}}, \langle \cdot, \cdot \rangle_{\mathbb{R}})$ ein Hilbertraum über $\mathbb{R}$, dessen Norm mit der von H übereinstimmt; und wenn nicht etwas anderes vereinbart ist, ist mit $H_{\mathbb{R}}$ dieser Hilbertraum über $\mathbb{R}$ gemeint.

3.4.11 Produkte von Hilberträumen (CONWAY [52](1990), Seite 24; MATHIEU [208](1998), Seite 219). Sei H_{ν} , $\nu \in I$, eine Familie von Hilberträumen, alle über den gleichen Körper $\mathbb{K}$. Auf dem Vektorraum $\bigoplus_{\nu \in I} H_{\nu}$ ist dann per

$$\langle x, y \rangle := \sum_{\nu \in I} \langle x_{\nu}, y_{\nu} \rangle \quad \text{für alle } x = (x_{\nu})_{\nu \in I}, \quad y = (y_{\nu})_{\nu \in I} \in \bigoplus_{\nu \in I} H_{\nu}$$

ein Skalarprodukt definiert. Die dadurch induzierte Norm ist gleich der 2-Norm auf $\bigoplus_{\nu \in I} H_{\nu}$.

Die Vervollständigung in dieser Norm, also $\ell_2(\bigoplus_{\nu \in I} H_{\nu})$, ist wieder ein Hilbertraum über $\mathbb{K}$; er heißt die *Hilbert'sche direkte* oder *ℓ_2 -direkte Summe* der H_{ν} , $\nu \in I$.

3.4.12 Definition. Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{K}$. Mit den Bezeichnungen aus Definition 2.11.15(f) lautet ein Konzept von Positivität auf $\mathcal{L}(H)$:

$$\text{Pos}(W; \mathcal{L}(H)) := \{T \in \mathcal{L}(H) : W_{\langle \cdot, \cdot \rangle}(T) \subseteq \mathbb{R}^+\}.$$

3.4.13 Bemerkung. Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{K}$ und $T \in \mathcal{L}(H)$. Dann sind äquivalent:

- (a) $T \in \text{Pos}(W; \mathcal{L}(H))$.
- (b) $\langle Tx, x \rangle \geq 0$ für alle $x \in H$.

Beweis. $\Rightarrow$: Sei $T \in \text{Pos}(W; \mathcal{L}(H))$. Dann gilt nach Definition, dass für alle $x \in H$ mit $\|x\| = 1$ die Ungleichung $\langle Tx, x \rangle \geq 0$ gilt. Für alle $x \in H^{\times}$ gilt somit $\langle Tx, x \rangle = \|x\|^2 \langle T(\frac{x}{\|x\|}), \frac{x}{\|x\|} \rangle \geq 0$. $\Leftarrow$: Gilt für ein $T \in \mathcal{L}(H)$ für alle $x \in H$ die Ungleichung $\langle Tx, x \rangle \geq 0$, so gilt sie insbesondere für alle $x \in H$ mit $\|x\| = 1$ und T ist in $\text{Pos}(W; \mathcal{L}(H))$. $\square$

3.4.14. Für einen Hilbertraum über $\mathbb{K}$ nennen PEDERSEN [234](1989), Seite 91, Nummer 3.2.8 und WEIDMANN [288](2000), Seite 140, ein $T \in \mathcal{L}(H)$ positiv, wenn T ein Element von $\mathcal{L}(H)_{(*)} \cap \text{Pos}(W; \mathcal{L}(H))$ ist. Die Menge der in diesem Sinne positiven Elemente von $\mathcal{L}(H)$ wird hier mit $\text{Pos}(*W; \mathcal{L}(H))$ bezeichnet. Man beachte, dass wegen der Bemerkung 3.4.13 und des Hinweises in 3.4.8 im Fall von $\mathbb{K} = \mathbb{C}$ die Mengen $\text{Pos}(W; \mathcal{L}(H))$ und $\text{Pos}(*W; \mathcal{L}(H))$ gleich sind.

3.4.15 Satz. Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{K}$. Dann gilt: $\text{Pos}(V; \mathcal{L}(H)) \subseteq \text{Pos}(W; \mathcal{L}(H))$.

Beweis. Folgt direkt aus Satz 2.11.21. $\square$

3.4.16 (WEIDMANN [288](2000)). Sei H ein Hilbertraum über $\mathbb{K}$. Für jedes $T \in \text{Pos}(*W; \mathcal{L}(H))$ existiert genau ein $S \in \text{Pos}(*W; \mathcal{L}(H))$ mit $T = S^2$ (ebd., Seite 303). Für alle $T \in \mathcal{L}(H)$ gilt $T^*T \in \text{Pos}(*W; \mathcal{L}(H))$ (ebd., Seite 168). Folglich hat man die Gleichung

$$\text{Pos}(*; \mathcal{L}(H)) = \text{Pos}(*W; \mathcal{L}(H)).$$

3.4.17 (WERNER [291](2002), Seite 212, Satz V.5.9). Sei H ein Hilbertraum über $\mathbb{K}$ und $T \in \mathcal{L}(H)$, $T \neq 0$, idempotent. Dann sind äquivalent: **(a)** T ist eine *Orthogonalprojektion* von H auf $\text{ran}(T)$, das heißt, $\forall x \in \ker(T) \forall y \in \text{ran}(T) : \langle x, y \rangle = 0$; **(b)** $\|T\| = 1$; **(c)** $T \in \mathcal{L}(H)_{(*)}$; **(d)** $T \in \mathcal{L}(H)_{\text{normal}}$; **(e)** $T \in \text{Pos}(W; \mathcal{L}(H))$; **(f)** $T \in \text{Pos}(*W; \mathcal{L}(H))$; **(g)** $\ker(T) = (\text{ran } T)^{\perp}$.

3.4.18. Sei H ein Hilbertraum über $\mathbb{K}$. Die Menge der Orthogonalprojektionen von $\mathcal{L}(H)$ wird mit $\text{OProj}(H)$ bezeichnet. Gemäß 3.4.8 und 3.4.17 sind die Projektionen $p \in \mathcal{L}(H)$ immer als Orthogonalprojektionen zu verstehen, das heißt, es gilt $\text{Proj}(\mathcal{L}(H)) = \text{OProj}(H)$. Nichtsdestotrotz wird zeitweilen der Deutlichkeit halber explizit das Wort Orthogonalprojektion verwendet.

3.4.19 Ordnungen von Orthogonalprojektionen (TAYLOR und LAY [275] (1980), Seite 348; MATHIEU [208] (1998), Seite 224). Sei H ein Hilbertraum über $\mathbb{K}$. Bezeichne $\leq_*$ bzw. $\leq_{*W}$ die durch den punktierten spitzen konvexen Kegel $\text{Pos}(*; \mathcal{L}(H))$ bzw. $\text{Pos}(*W; \mathcal{L}(H))$ induzierte partielle Ordnung auf $\mathcal{L}(H)$. (Es sei an 3.4.16 erinnert.) Bezeichne $\leq_{\text{UR}}$ die per (2.20) in 2.1.34 definierte partielle Ordnung auf der Menge der (wenn man möchte, abgeschlossenen) Unterräume von $\mathcal{L}(H)$. Seien $\leq_r$ und $\leq_\ell$ die Präordnungen und $\leq_{\ell r}$ die partielle Ordnung jeweils von 2.1.45 auf $\text{Idem}(\mathcal{L}(H))$. Seien $p, q \in \mathcal{L}(H)$ zwei Orthogonalprojektionen. Dann sind die folgenden sieben Aussagen äquivalent:

- | | |
|--|---------------------------|
| (a) $p \leq_* q$; | (b) $p \leq_{*W} q$; |
| (c) $\text{ran}(p) \leq_{\text{UR}} \text{ran}(q)$; | (d) $p \leq_{\ell r} q$; |
| (e) $p \leq_\ell q$; | (f) $p \leq_r q$; |
| (g) $\ px\ \leq \ qx\ $ für alle $x \in H$. | |

Falls nicht ausdrücklich etwas anderes vereinbart ist, meint $p \leq q$ eine beliebige der vorstehenden äquivalenten Aussagen und die Menge $\text{OProj}(H)$ ist stets mit dieser partiellen Ordnung $\leq$ versehen. $\text{OProj}(H)$ ist ein Verband; dabei ist das Supremum $p \vee q$ gleich der Orthogonalprojektion auf $\text{cl span}(\text{ran}(p) \cup \text{ran}(q))$ und das Infimum $p \wedge q$ gleich der Orthogonalprojektion auf $\text{ran}(p) \cap \text{ran}(q)$. 0 bzw. Id_H ist das kleinste bzw. größte Element von $\text{OProj}(H)$. Jede nicht leere Teilmenge von $\text{OProj}(H)$ besitzt ein Infimum und ein Supremum, zum Beweis siehe 3.7.34.

3.4.20 Träger (MATHIEU [208] (1998), Seite 327). Sei H ein Hilbertraum über $\mathbb{K}$ und $T \in \mathcal{L}(H)$. Eine Orthogonalprojektion auf einen abgeschlossenen Unterraum U von H werde mit p_U bezeichnet. Die Orthogonalprojektion

$$\text{pcran}(T) := p_{\text{cl ran}(T)}$$

auf den Abschluss des Bildes von T heißt die *Bildprojektion* (im Englischen *range projection*) von T . Die Orthogonalprojektion

$$\text{pker}(T) := p_{\text{ker}(T)}$$

auf den Kern von T heißt die *Kernprojektion* (im Englischen *null projection*) von T . Offensichtlich gilt

$$\text{pcran}(T) \circ T = T \quad \text{und} \quad T \circ \text{pker}(T) = 0. \quad (3.2)$$

Des Weiteren gilt

$$\text{pcran}(T) = \min \{p \in \text{OProj}(H) : pT = T\}, \quad (3.3)$$

$$\text{pcran}(T^*) = \min \{p \in \text{OProj}(H) : Tp = T\}, \quad (3.4)$$

$$\text{pker}(T) = \max \{p \in \text{OProj}(H) : Tp = 0\}. \quad (3.5)$$

Die Projektion

$$s_\ell(T) := \text{pcran}(T) \quad (3.6)$$

heißt auch der *Linksträger* (im Englischen *left support*) von T . Die Projektion

$$s_r(T) := \text{pcran}(T^*) \quad (3.7)$$

heißt der *Rechtsträger* (im Englischen *right support*) von T . Offensichtlich gilt

$$s_\ell(T) = s_r(T^*). \quad (3.8)$$

Falls T selbstadjungiert ist, heißt die Projektion

$$s(T) := s_\ell(T) = s_r(T) \quad (3.9)$$

der *Träger* (im Englischen *support*) von T . Es gilt

$$s(T^*T) = s_r(T) = \text{Id}_H - \text{pker}(T), \quad (3.10)$$

$$s(TT^*) = s_\ell(T), \quad (3.11)$$

$$\text{pker}(T^*T) = \text{pker}(T). \quad (3.12)$$

3.4.21 Partielle Isometrien (PEDERSEN [233](1979), Seite 23; STRĂȚILĂ und ZSIDÓ [270](1979), Seite 36; SUNDER [272](2000); WEIDMANN [288](2000), Seiten 116, 117; ZEIDLER [296] (1995), Seite 74).

Seien H_1 und H_2 Hilberträume über $\mathbb{K}$. Sei $U \in \mathcal{L}(H_1, H_2)$. Die folgenden Aussagen sind äquivalent

- (a) U ist eine Isometrie.
- (b) Für alle $x, y \in H_1$ gilt $\langle Ux, Uy \rangle = \langle x, y \rangle$.
- (c) $U^*U = \text{Id}_{H_1}$.
- (d) Falls $\{x_\nu : \nu \in I\}$ ein Orthonormalsystem in H_1 ist, ist auch $\{Ux_\nu : \nu \in I\}$ ein Orthonormalsystem in H_2 .
- (e) In H_1 existiert eine Orthonormalbasis $\{b_\nu : \nu \in I\}$, so dass $\{Ub_\nu : \nu \in I\}$ ein Orthonormalsystem in H_2 ist.

Beweis. (a) $\Rightarrow$ (b): Für $\mathbb{K} = \mathbb{R}$ siehe 2.2.41 mit $F = \langle \cdot, \cdot \rangle$. Für $\mathbb{K} = \mathbb{C}$ siehe 2.2.42 mit $B = \langle \cdot, \cdot \rangle$.

(b) $\Rightarrow$ (a): Klar.

(b) $\Rightarrow$ (c): Für alle $x \in H_1$ gilt $\langle U^*Ux, x \rangle = \langle x, x \rangle = \langle \text{Id}_{H_1}x, x \rangle$. Da $U^*U - \text{Id}_{H_1}$ selbstadjungiert ist, gilt mit 3.4.9 $\|U^*U - \text{Id}_{H_1}\| = \sup \{ \langle (U^*U - \text{Id}_{H_1})x, x \rangle : x \in S_{H_1} \} = 0$.

(c) $\Rightarrow$ (d): Sei $x_\nu, \nu \in I$, ein Orthonormalsystem in H_1 . Dann gilt $\langle Ux_\nu, Ux_\mu \rangle = \langle U^*Ux_\nu, x_\mu \rangle = \langle x_\nu, x_\mu \rangle = \delta_{\nu\mu}$ für alle $\nu, \mu \in I$.

(d) $\Rightarrow$ (e): Klar.

(e) $\Rightarrow$ (b): Sei $b_\nu, \nu \in I$, eine Orthonormalbasis gemäß (e). Für alle $x, y \in H_1$ gilt dann: $\langle Ux, Uy \rangle = \langle U(\sum_{\nu \in I} \langle x, b_\nu \rangle b_\nu), U(\sum_{\mu \in I} \langle y, b_\mu \rangle b_\mu) \rangle = \sum_{\nu \in I} \sum_{\mu \in I} \langle x, b_\nu \rangle \langle b_\mu, y \rangle \langle Ub_\nu, Ub_\mu \rangle = \sum_{\nu \in I} \langle x, b_\nu \rangle \langle b_\nu, y \rangle = \langle x, y \rangle$. $\square$

U heißt eine *partielle Isometrie*, wenn ein abgeschlossener Unterraum D von H_1 existiert, so dass sowohl $U \upharpoonright D$ eine Isometrie ist als auch $U \upharpoonright D^\perp = 0$ gilt. Das Bild $\text{ran}(U)$ von U ist dann abgeschlossen (WEIDMANN, ebd.). D heißt die *Anfangsmenge* (im Englischen *initial domain*) und $\text{ran}(U)$ heißt die *Endmenge* (im Englischen *final domain*) der partiellen Isometrie U .

Sei D ein abgeschlossener Unterraum von H_1 und F ein abgeschlossener Unterraum von H_2 . Bezeichne $p_D \in \text{OProj}(H_1)$ bzw. $p_F \in \text{OProj}(H_2)$ die Orthogonalprojektion auf D bzw. F . Bezeichne U^* die in 3.4.3 definierte Hilbertraumadjungierte von U . Dann sind die folgenden fünf Aussagen äquivalent:

- (a) U ist eine partielle Isometrie mit Anfangsmenge D und Endmenge F .
- (b) $\text{ran}(U) = F$ und für alle $x, y \in H_1$ gilt $\langle Ux, Uy \rangle = \langle p_D x, y \rangle$.
- (c) $U^*U = p_D$ und $UU^* = p_F$.
- (d) D und F weisen dieselbe Hilbertraumdimension auf.
- (e) U^* ist eine partielle Isometrie mit Anfangsmenge F und Endmenge D .

Ist U eine partielle Isometrie mit Anfangsmenge D und Endmenge F , so heißt p_D die *Anfangsprojektion* und p_F die *Endprojektion* von U .

Liegt bei den vorstehenden Äquivalenzen der Spezialfall vor, dass $D = H_1$ und $F = H_2$ gilt, so hat man die Äquivalenz der folgenden sieben Aussagen:

- (a) U ist ein isometrischer Isomorphismus von H_1 auf H_2 .
- (b) U ist eine surjektive Isometrie.
- (c) $U^*U = \text{Id}_{H_1}$ und $UU^* = \text{Id}_{H_2}$, das heißt im Fall von $H_1 = H_2$, dass U gemäß 3.4.8 unitär ist.
- (d) U ist bijektiv und es gilt $U^{-1} = U^*$.
- (e) U^* ist ein isometrischer Isomorphismus von H_2 auf H_1 .
- (f) U ist eine Isometrie und falls $\{b_\nu : \nu \in I\}$ eine Orthonormalbasis für H_1 ist, ist auch $\{Ub_\nu : \nu \in I\}$ eine Orthonormalbasis für H_2 .
- (g) U ist eine Isometrie und in H_1 existiert ein Orthonormalsystem $\{x_\nu : \nu \in I\}$, so dass $\{Ux_\nu : \nu \in I\}$ eine Orthonormalbasis für H_2 ist.

Beweis. (b) $\Rightarrow$ (f): Gelte (b) und sei $\{b_\nu : \nu \in I\}$ eine Orthonormalbasis für H_1 . Da U eine Isometrie ist, ist $\{Ub_\nu : \nu \in I\}$ ein Orthonormalsystem in H_2 . Sei $y \in H_2$. Dann existiert genau ein $x \in H_1$ mit $y = Ux$ und es gilt: $\|y\|^2 = \|Ux\|^2 = \|x\|^2 = \sum_{\nu \in I} |\langle x, b_\nu \rangle|^2 = \sum_{\nu \in I} |\langle y, Ub_\nu \rangle|^2$.

(f) $\Rightarrow$ (g): Klar.

(g) $\Rightarrow$ (c): Sei $y \in H_2$. Dann gilt $UU^*y = UU^*(\sum_{\nu \in I} \langle y, Ux_\nu \rangle Ux_\nu) = \sum_{\nu \in I} \langle y, Ux_\nu \rangle (UU^*)Ux_\nu = \sum_{\nu \in I} \langle y, Ux_\nu \rangle U(U^*U)x_\nu = \sum_{\nu \in I} \langle y, Ux_\nu \rangle Ux_\nu = y$. $\square$

Für zwei weitere äquivalente Aussagen sei an Lemma 2.11.29 erinnert. Isometrische Isomorphismen von H_1 auf H_2 werden *unitäre Operatoren* genannt; möchte man den Fall $\mathbb{K} = \mathbb{R}$ betonen, nennt man sie auch *orthogonale Operatoren*.

Als ein Beispiel eines zwar isometrischen, aber nicht surjektiven Operators hat man den sogenannten Rechtsshiftoperator auf dem Folgenraum ℓ_2 über $\mathbb{K}$, soll heißen $U: \ell_2 \rightarrow \ell_2, (a_1, a_2, \dots) \mapsto (0, a_1, a_2, \dots)$. Falls H_1 und H_2 isometrisch isomorph sind, heißen ein $T_1 \in \mathcal{L}(H_1)$ und ein $T_2 \in \mathcal{L}(H_2)$ *unitär äquivalent*, falls es einen unitären Operator $U \in \mathcal{L}(H_1, H_2)$ mit $T_1 = U^*T_2U$ gibt. Offensichtlich wird dadurch für jeden Hilbertraum H über $\mathbb{K}$ eine Äquivalenzrelation auf $\mathcal{L}(H)$ erklärt.

Bezüglich partieller Isometrien lassen sich auch Äquivalenzen formulieren, ohne dass explizit die Anfangs- und Endmengen erwähnt werden. Die folgenden Aussagen sind für ein $U \in \mathcal{L}(H_1, H_2)$ äquivalent:

- (a) U ist eine partielle Isometrie.
- (b) U^*U ist eine Projektion.
- (c) UU^* ist eine Projektion.
- (d) U^* ist eine partielle Isometrie.
- (e) $U \upharpoonright (\ker U)^\perp$ ist eine Isometrie.
- (f) $UU^*U = U$.

Beweis. (b) $\Rightarrow$ (a): Sei U^*U eine Projektion. Setze $D := (U^*U)(H_1)$. Dann ist D ein abgeschlossener Unterraum von H_1 und $p_D = U^*U$. Sei $x \in \text{ran}(p_D)$. Dann ist $\langle x, x \rangle = \langle U^*Ux, x \rangle = \langle Ux, Ux \rangle$, also $\|x\| = \|Ux\|$ und $U \upharpoonright D$ ist eine Isometrie. Sei $x \in D^\perp$, also $x \in (\text{ran}(p_D))^\perp = \ker(p_D^*) = \ker(p_D)$. Somit $\langle Ux, Ux \rangle = \langle U^*Ux, x \rangle = 0$, $\|Ux\| = 0$, $x \in \ker U$ und $U \upharpoonright D^\perp = 0$.

(b) $\Rightarrow$ (f): Klar.

(f) $\Rightarrow$ (b): Gelte (f). Dann hat man $(U^*U)(U^*U) = U^*(UU^*U) = U^*U$.

(a) $\Rightarrow$ (e): Gelte (a), das heißt, es existiert ein abgeschlossener Unterraum $D \subseteq H_1$ mit $U \upharpoonright D$ isometrisch und $U \upharpoonright D^\perp = 0$. Also $D^\perp \subseteq \ker U$, $(\ker U)^\perp \subseteq D^{\perp\perp}$. Da allgemein für jeden Unterraum U eines Hilbertraumes H die Gleichung $\text{cl}(U) = U^{\perp\perp}$ gilt, folgt hier $(\ker U)^\perp \subseteq D$ und $U \upharpoonright (\ker U)^\perp$ ist eine Isometrie.

(e) $\Rightarrow$ (a): Klar. $\square$

Gilt eine (und damit jede) der vorstehenden Äquivalenzen, so gilt für den Rechtsträger von U die Gleichung

$$s_r(U) = U^*U,$$

$s_r(U)$ ist also die Anfangsprojektion der partiellen Isometrie U , und für den Linksträger von U gilt die Gleichung

$$s_\ell(U) = UU^*,$$

$s_\ell(U)$ ist also die Endprojektion der partiellen Isometrie U .

3.4.22 Kompression (KADISON und RINGROSE [167](1983), Seite 120; DIXMIER [76](1969), Seite 17; MURPHY [217](2004), Seite 117; ALFSEN und SHULTZ [8](2001), Seite 168). Sei H ein Hilbertraum über $\mathbb{K}$.

(a) Ist A eine $*$ -Algebra über $\mathbb{K}$ und $a \in A$ ein selbstadjungiertes Element, so wird mit U_a die lineare Abbildung $A \rightarrow A, x \mapsto axa$ bezeichnet. (Siehe auch die Abbildung $Q: X \rightarrow L(X)$ in Definition 4.6.3.) Sei G ein abgeschlossener Unterraum von H . Sei $p \in \mathcal{L}(H)$ die Projektion von H auf G und $T \in \mathcal{L}(H)$. (Dann ist $U_p \in \mathcal{L}(\mathcal{L}(H))$ eine Projektion.) Dann wird die Abbildung, die zugleich eine Astriktion und eine Restriktion auf G von der Abbildung pT ist, also die Abbildung

$$T_p: G \rightarrow G, x \mapsto pT(x) \quad (= U_p(T)(x)),$$

die *Kompression* von T auf G genannt. Hierbei ist es üblich, die Astriktion unausgesprochen vorauszusetzen, so dass man zum Beispiel die Kompression per $T_p := pT \upharpoonright G$ definiert, dabei aber genauer eigentlich das durch $T_p := G \upharpoonright pT \upharpoonright G$ definierte Element aus $\mathcal{L}(G)$ meint. In diesem Sinne wird auch hier im Folgenden die Astriktion stillschweigend vorausgesetzt. Anstelle von T_p schreibt man auch T_G .

(b) Wegen $\langle pTpx, y \rangle = \langle x, pT^*py \rangle$ für alle $x, y \in G$ ist die Hilbertraumadjungierte von T_p gleich $pT^*|_G$, also gleich der Kompression der Hilbertraumadjungierten T^* von T : $(T^*)_p = (T_p)^*$. Somit ist die Abbildung

$$U_p(\mathcal{L}(H)) \rightarrow \mathcal{L}(G), \quad pTp \mapsto T_p$$

ein $*$ -Isomorphismus.

Beweis. Injektivität: Sei $S \in U_p(\mathcal{L}(H))$, also $S = pTp$ für ein $T \in \mathcal{L}(H)$. Falls $S|_G = 0$, dann wegen $S = pSp$ auch $S|_{G^\perp} = 0$, also $S = 0$. $\square$

(c) Ist A eine Teilmenge von $\mathcal{L}(H)$ und $p \in \mathcal{L}(H)$ eine Projektion, so setzt man:

$$A_p := \{T_p \in \mathcal{L}(pH) : T \in A\},$$

die sogenannte *Kompression* von A (unter p).

(d) Sei $A \subseteq \mathcal{L}(H)$ und $p \in \mathcal{L}(H)$ eine Projektion mit $p \in A''$. Dann gilt $A_p \subseteq ((A')_p)' \subseteq (A'')_p$ und $(A')_p \subseteq (A_p)'$.

Beweis. Sei $x \in A_p$, also $x = p\tilde{x}p|pH$ für ein $\tilde{x} \in A$. Sei $y \in (A')_p$, also $y = p\tilde{y}p|pH$ für ein $\tilde{y} \in A'$. Dann gilt: $p\tilde{x}pp\tilde{y}p = p\tilde{x}p\tilde{y}p = p\tilde{x}\tilde{y}p$ (da $p \in A''$) $= p\tilde{y}\tilde{x}p$ (da $\tilde{y} \in A'$) $= pp\tilde{y}\tilde{x}p = p\tilde{y}p\tilde{x}p$ (da $p \in A''$) $= p\tilde{y}pp\tilde{x}p$, also insbesondere $xy = yx$, das heißt, $A_p \subseteq ((A')_p)'$ und $(A')_p \subseteq (A_p)'$. Sei nun $x \in ((A')_p)'$, also $x = p\tilde{x}p|pH$ für ein $\tilde{x} \in \mathcal{L}(H)$, wobei man ohne Einschränkung $\tilde{x} = p\tilde{x}p$ annehmen darf. Sei des Weiteren $\tilde{z} \in A'$. Dann gilt auf H : $\tilde{x}\tilde{z} = p\tilde{x}p\tilde{z} = p\tilde{x}pp\tilde{z} = xp\tilde{z} = xpp\tilde{z} = xp\tilde{z}p$ (da $p \in A''$) $= p\tilde{z}pxp$ (da $x \in ((A')_p)'$ und $p\tilde{z}p|pH \in (A')_p$) $= \tilde{z}pxp$ (da $p \in A''$ und $\tilde{z} \in A'$) $= \tilde{z}p\tilde{x}p$ (da $xp = p\tilde{x}p$) $= \tilde{z}\tilde{x}$, also $\tilde{x} \in A''$, $x \in (A'')_p$, $((A')_p)' \subseteq (A'')_p$. $\square$

(e) Sei $A \subseteq \mathcal{L}(H)$ und $p \in \mathcal{L}(H)$ eine Projektion mit $p \in A'$. Dann gilt $(A_p)' \subseteq (A')_p$ und nach (d) gilt Gleichheit, wenn zusätzlich $p \in A''$ gilt. Falls A eine $*$ -Unteralgebra von $\mathcal{L}(H)$ ist, dann sind $U_p(A) \subseteq \mathcal{L}(H)$ und $A_p \subseteq \mathcal{L}(pH)$ $*$ -Algebren und die Abbildung

$$U_p(A) \rightarrow A_p, \quad pTp \mapsto T_p$$

ist ein $*$ -Isomorphismus.

Beweis. Sei $u \in (A_p)'$, also $u = p\tilde{u}p|pH$ für ein $\tilde{u} \in \mathcal{L}(H)$, wobei man ohne Einschränkung $\tilde{u} = p\tilde{u}p$ annehmen darf. Sei $v \in A$. Dann gilt auf H : $v\tilde{u} = vp\tilde{u}p = pv\tilde{u}p$ (da $p \in A'$) $= pvp\tilde{u}p = pvpup = upvp = p\tilde{u}vp = \tilde{u}v$ (da $p \in A'$), also $\tilde{u} \in A'$ und $u \in (A')_p$.

Die Aussage $(A')_p \subseteq (A_p)'$ folgt aus (d) und auch im vorliegenden Fall mit genau dem Beweis von (d). $\square$

(f) Sei $A \subseteq \mathcal{L}(H)$ und $p \in \mathcal{L}(H)$ eine Projektion mit $p \in A' \cap A''$. Durch anwenden von (d) und (e) auf A' erhält man: $A_p \subseteq ((A')_p)' = ((A'')_p)'' = (A_p)'' = (A'')_p$

Beweis. $(A')_p \subseteq ((A'')_p)'$, also $((A'')_p)'' \subseteq ((A')_p)' = (A'')_p \subseteq ((A'')_p)''$. $(A_p)'' = ((A')_p)' = (A'')_p$ $\square$

3.4.23 Definition (PALMER [232](2001), Seite 863). (Vergleiche Abschnitt 2.8.) Sei A eine $*$ -Algebra über $\mathbb{K}$. Eine $*$ -Darstellung (bzw. $*$ -Antidarstellung) von A auf einem Hilbertraum H über $\mathbb{K}$ ist ein $*$ -Homomorphismus (bzw. $*$ -Antihomomorphismus) $T: A \rightarrow \mathcal{L}(H)$. Dass hierbei direkt $\mathcal{L}(H)$ anstelle von $L(H)$ verwendet wird, also jede $*$ -Darstellung von vornherein normiert ist, ist eine unmittelbare Konsequenz des Satzes von Hellinger-Toeplitz, siehe 3.4.6.

Sei A eine $*$ -Unteralgebra von $\mathcal{L}(H)$, H ein Hilbertraum über $\mathbb{K}$. Ist keine Darstellung von A näher spezifiziert, dann heißt A *wesentlich*, wenn die auf A eingeschränkte identische Abbildung aus $\mathcal{L}(\mathcal{L}(H))$ eine wesentliche Darstellung von A auf H ist, mit anderen Worten, wenn für alle $x \in H^\times$ ein $T \in A$ existiert mit $Tx \neq 0$.

3.4.24 Bemerkung (SUNDER [271](1987), Seite 14). Sei A eine $*$ -Unteralgebra von $\mathcal{L}(H)$ für einen Hilbertraum über $\mathbb{C}$. Sei p die Projektion auf den abgeschlossenen Unterraum $\text{cl span}(AH)$. Dann gilt $T = pTp$ für alle $T \in A$, insbesondere also $p \in A'$ und $A_pH \subseteq pH$.

Bezeichne A_p die Menge von 3.4.22, hier also $A_p = \{T|pH \in \mathcal{L}(pH) : T \in A\}$. Dann ist A_p eine wesentliche $*$ -Unteralgebra von $\mathcal{L}(pH)$. Des Weiteren gilt: $A' = (A_p)' \oplus \mathcal{L}((\text{Id}_H - p)H)$ und $A'' = (A_p)'' \oplus \mathbb{C}(\text{Id}_H - p)$.

3.4.25 Satz. Sei T eine $*$ -Darstellung einer $*$ -Algebra über $\mathbb{C}$ auf einen Hilbertraum über $\mathbb{C}$. Dann ist T genau dann wesentlich, wenn die Menge $T_AH = \langle A, H \rangle_T = \{T_a x : a \in A, x \in H\}$ dicht in H ist.

Beweis. PALMER [232](2001), Seite 865: Beweis mit topologisch zyklischen Vektoren. $\square$

3.4.26. PALMER [232](2001), Seite 865, zeigt in einem Satz die Äquivalenz von vier Aussagen, von denen hier durch den Satz 3.4.25 zwei wiedergegeben sind. Bezüglich seines Beweises weist Palmer ausdrücklich darauf hin, dass $T_A H$ a priori kein Vektorraum zu sein braucht. Insbesondere kann $(T_A H)^\perp = \{0\}$ sein, ohne dass $T_A H$ dicht in H ist! Dass aber zumindest $cl(T_A H)$ ein Vektorraum ist, zeigt Satz 3.4.27, der eine Folgerung von Satz 3.4.25 ist. Einen einfachen Beweis von Satz 3.4.25 für den $\mathbb{K}$ -Fall und $\text{span}(T_A H)$ anstelle von $T_A H$ findet man bei TAKESAKI [274](2001), Seite 36.

3.4.27 Satz (PALMER [232](2001), Seite 866). *Sei T eine $*$ -Darstellung einer $*$ -Algebra A über $\mathbb{C}$ auf einem Hilbertraum H über $\mathbb{C}$. Dann ist H die orthogonale direkte Summe der abgeschlossenen T -invarianten Unterräume $H_T = cl(T_A H)$ und $H_0 = \bigcap \{ker(T_a) : a \in A\}$. Die Restriktion von T auf den Unterraum H_T (bzw. H_0) ist eine wesentliche (bzw. triviale) Sub- $*$ -Darstellung, die alle wesentlichen (bzw. trivialen) Sub- $*$ -Darstellungen enthält. Für die Projektion p auf H_T gilt: $pT_a = T_a p = T_a$ für alle $a \in A$.*

3.4.28 Bemerkung. Wegen Satz 3.4.27 genügt es, nur die wesentlichen $*$ -Darstellungen zu betrachten.

Die Formulierung in Satz 3.4.27 „ $T \upharpoonright H_0$ enthält alle trivialen Sub- $*$ -Darstellungen“ heißt, dass jede triviale Sub- $*$ -Darstellung die Form $T \upharpoonright G$ mit $G \subseteq H_0$ hat. Analog ist „ $T \upharpoonright H_T$ enthält alle wesentlichen Sub- $*$ -Darstellungen“ zu lesen als: Jede wesentliche Sub- $*$ -Darstellung von T hat die Form $T \upharpoonright G$ mit $G \subseteq H_T$. Siehe auch DIXMIER [76](1981), Seite 44.

3.4.29 (PALMER [232](2001), Seiten 979, 981, 986). Sei A eine $*$ -Algebra über $\mathbb{K}$. Der Durchschnitt der Kerne von allen $*$ -Darstellungen (diese jeweils auf Hilberträume über $\mathbb{K}$) heißt das *reduzierende Ideal* von A , in Zeichen A_{red} . (Man beachte, dass hierbei keine Irreduzibilität oder dergleichen gefordert wird.) A heißt *reduziert*, falls $A_{\text{red}} = \{0\}$ gilt; A heißt *$*$ -radikal*, falls $A_{\text{red}} = A$ gilt.

Sei A eine $*$ -Algebra über $\mathbb{C}$. Ist B eine $*$ -Unteralgebra von A , so gilt $B_{\text{red}} \subseteq A_{\text{red}} \cap B$. Des Weiteren gilt $\mathfrak{J}(A) \subseteq A_{\text{red}}$. Somit ist insbesondere jede reduzierte $*$ -Algebra über $\mathbb{C}$ und jede ihrer $*$ -Unteralgebren Jacobson-halbeinfach.

3.4.30 (DIEUDONNÉ [68](1987), Übung 12.15.5). Sei H ein Hilbertraum über $\mathbb{K}$ und $K \subseteq H$ konvex. Ist $x \in K$ ein Extrempunkt von K , so bilden die x enthaltenden offenen Kalotten von K eine Umgebungsbasis von x in K .

3.5 C*-Algebren

3.5.1 Definition (PALMER [232](2001), Seite 868).

- (a) Sei H ein Hilbertraum über $\mathbb{K}$. Eine abgeschlossene $*$ -Unteralgebra A von $\mathcal{L}(H)$ heißt eine *C^* -Algebra über $\mathbb{K}$ von Operatoren*. Genauer sagt man auch, dass A eine *C^* -Algebra über $\mathbb{K}$ von Operatoren auf H* sei.
- (b) Eine Banach- $*$ -Algebra über $\mathbb{R}$, welche isometrisch $*$ -isomorph zu einer C^* -Algebra über $\mathbb{R}$ von Operatoren ist, heißt eine *C^* -Algebra über $\mathbb{R}$* .
- (c) Eine $*$ -Algebra über $\mathbb{C}$, welche $*$ -isomorph zu einer C^* -Algebra über $\mathbb{C}$ von Operatoren ist, heißt eine *C^* -Algebra über $\mathbb{C}$* .

3.5.2 Bemerkung. Eine äquivalente Formulierung der Definition einer C^* -Algebra über $\mathbb{C}$ lautet: Eine $*$ -Algebra über $\mathbb{C}$, die eine treue $*$ -Darstellung als eine C^* -Algebra über $\mathbb{C}$ von Operatoren hat, heißt eine C^* -Algebra über $\mathbb{C}$.

3.5.3. Trivialerweise ist jede C^* -Algebra über $\mathbb{K}$ von Operatoren eine C^* -Algebra über $\mathbb{K}$.

3.5.4 Definition. Sei A eine $*$ -Algebra über $\mathbb{K}$.

(a) Erfüllt eine Norm $\|\cdot\|$ auf A die sogenannte *schwache C^* -Bedingung*

$$\|a^*a\| = \|a^*\|\|a\| \quad \text{für alle } a \in A,$$

dann heißt die Norm eine *schwache C^* -Norm*. Ist die Norm zugleich eine Algebranorm, dann heißt sie eine *schwache C^* -Algebranorm*.

(b) Erfüllt eine Norm $\|\cdot\|$ auf A die sogenannte *C^* -Bedingung*

$$\|a^*a\| = \|a\|^2 \quad \text{für alle } a \in A,$$

dann heißt die Norm eine *C^* -Norm*. Ist die Norm zugleich eine Algebranorm, dann heißt sie eine *C^* -Algebranorm*.

3.5.5 Bemerkung. Ist A eine normierte $*$ -Algebra über $\mathbb{K}$, deren Norm eine C^* -Norm ist, so ist nach Satz 3.2.26 ihre Involution eine Isometrie (sprich eine Konjunktion) und folglich ihre Norm auch eine schwache C^* -Algebranorm.

3.5.6 Bemerkung. Sei A eine $*$ -Algebra über $\mathbb{K}$, die eine Norm $\|\cdot\|$ trägt, die eine schwache C^* -Norm oder eine C^* -Norm ist. Für jede Projektion $p \in A^\times$ gilt dann $\|p\| = 1$. Man bemerke, dass diese Gleichung nach Bemerkung 3.2.12(b) insbesondere auch für eine eventuell vorhandene von Null verschiedene Eins von A gilt.

3.5.7 Satz (DORAN und BELFI [77](1986), Seite 166: SEBESTYÉN). *Jede C^* -Norm auf einer $*$ -Algebra über $\mathbb{C}$ ist eine Algebranorm.*

3.5.8 Satz. *Sei A eine $*$ -Algebra über $\mathbb{C}$. Dann sind äquivalent:*

- (a) A ist eine C^* -Algebra über $\mathbb{C}$.
- (b) Es gibt auf A eine vollständige schwache C^* -Norm.
- (b') Es gibt auf A eine vollständige schwache C^* -Algebranorm.
- (c) Es gibt auf A eine vollständige C^* -Norm.
- (c') Es gibt auf A eine vollständige C^* -Algebranorm.
- (d) Die Abbildung $\|\cdot\|: a \mapsto \sqrt{\rho(a^*a)}$ für alle $a \in A$ ist die einzige C^* -Norm auf A . Die Norm $\|\cdot\|$ ist vollständig.

Ist A kommutativ, dann sind zusätzlich dazu äquivalent (man beachte Bemerkung 2.4.17(e)):

- (e) Der Gelfand-Homomorphismus ist ein $*$ -Isomorphismus auf $C_0(\Gamma_A, \mathbb{C})$.
- (f) A ist $*$ -isomorph zu $C_0(\Omega, \mathbb{C})$ für einen lokal kompakten Hausdorffraum Ω .

Beweis. Diesen Satz mit (a), (c'), (d), (e) und (f) findet man mit Beweis bei PALMER [232](2001), Seite 947.

(c) und (c') sind wegen Satz 3.5.7 äquivalent. Siehe dazu auch DORAN und BELFI [77](1986), Seite 165, die einen Satz von ARAKI und ELLIOT zitieren, in dem die Äquivalenz von (c) und (c') auf direktem Wege gezeigt wird.

Die Äquivalenz von (b) und (c') wird bei DORAN und BELFI [77](1986) auf den Seiten 165 und 166 belegt.

(c') $\Rightarrow$ (b'): Klar nach Bemerkung 3.5.5.

(b') $\Rightarrow$ (c'): Siehe DORAN und BELFI [77](1986), Seite 45. Sie zeigen die nichttriviale Folgerung, dass (b') impliziert, dass die Involution stetig und schließlich eine Isometrie ist. Als alternativen Beweis zeigen sie auf Seite 198 die Implikation (b') $\Rightarrow$ (c') als Korollar des Satzes von Vidav-Palmer. $\square$

3.5.9 Bemerkung. Gemäß Satz 3.5.8 betrachte man bei Bedarf jede C^* -Algebra über $\mathbb{C}$ ausgestattet mit der Norm $\|a\| = \sqrt{\rho(a^*a)}$ für alle $a \in A$. Diese Norm wird als die *intrinsische* C^* -Norm der C^* -Algebra über $\mathbb{C}$ angesprochen.

3.5.10 Bemerkungen (PALMER [232](2001), Seiten 871 und 954). Sei A ist eine C^* -Algebra über $\mathbb{C}$, versehen mit ihrer intrinsischen C^* -Algebra-Norm. Jede treue $*$ -Darstellung ist eine Isometrie. Somit ist jede abgeschlossene $*$ -Unteralgebra über $\mathbb{C}$ von A auch eine C^* -Algebra über $\mathbb{C}$. Jede $*$ -Darstellung T von A ist kontraktiv und das Bild T_A ist eine C^* -Algebra über $\mathbb{C}$.

Sei nun B eine weitere C^* -Algebra über $\mathbb{C}$, ebenfalls versehen mit ihrer intrinsischen C^* -Algebra-Norm. Sei $T: A \rightarrow B$ ein $*$ -Homomorphismus. Dann ist T kontraktiv und $\ker(T)$ und $\text{ran}(T)$ sind beide abgeschlossen.

Siehe auch 4.7.22 für die allgemeinere analoge Aussage für JB^* -Tripel über $\mathbb{C}$.

Sind A und B zwei C^* -Algebren über $\mathbb{K}$ und ist $T: A \rightarrow B$ ein $*$ -Isomorphismus von A auf B , so ist T isometrisch (CHU, DANG, RUSSO und VENTURA [48](1993)).

3.5.11 Satz. Für eine normierte $*$ -Algebra A über $\mathbb{R}$ sind äquivalent:

- (a) A ist eine C^* -Algebra über $\mathbb{R}$.
- (b) A ist symmetrisch und die Norm ist eine vollständige schwache C^* -Norm.
- (c) A ist symmetrisch und die Norm ist eine vollständige C^* -Norm.
- (d) $\sigma(a^*a) \subseteq \mathbb{R}^+$ für alle $a \in A$ und die Norm ist eine vollständige C^* -Norm.
- (e) $-s^2$ ist quasi-invertierbar für alle $s \in A_{(*)}$.
- (f) $\|s\|^2 \leq \lambda \|s^2 + t^2\|$ für alle $s, t \in A_{(*)}$, für ein konstantes $\lambda \in \mathbb{R}$ und die Norm ist eine C^* -Norm.
- (g) $\|s\|^2 \leq \lambda \|s^2 + t^2\|$ für alle kommutierenden Paare von Elementen $s, t \in A_{(*)}$, für ein konstantes $\lambda \in \mathbb{R}$ und die Norm ist eine vollständige C^* -Norm.
- (h) $\|a\|^2 \leq \|a^*a + b^*b\|$ für alle $a, b \in A$.
- (i) Die Norm von A ist vollständig, $A_{(\mathbb{C})}$ kann so normiert werden, dass $A_{(\mathbb{C})}$ eine C^* -Algebra über $\mathbb{C}$ ist und A kanonisch isometrisch isomorph in $A_{(\mathbb{C})\mathbb{R}}$ eingebettet werden kann.
- (j) $\|a^*\| \|a\| \leq \|a^*a + b^*b\|$ für alle $a, b \in A$ und die Norm ist vollständig.
- (k) A ist hermitesch und die Norm ist eine vollständige schwache C^* -Norm.
- (l) Es existiert eine C^* -Algebra B über $\mathbb{C}$ und ein $*$ -Automorphismus $\bar{}$ auf B mit Periode 2, so dass $A = B_{(-)\mathbb{R}}$ gilt.

Beweis. (a) $\Rightarrow$ (c): GOODEARL [117](1982), Seite 68.

(c) $\Rightarrow$ (a): INGELSTAM [145](1964), Korollar 17.7 auf Seite 268.

(a) $\Rightarrow$ (h): INGELSTAM [145](1964), Seite 265.

(h) $\Rightarrow$ (a): INGELSTAM [145](1964), Korollar 18.8 auf Seite 269.

(c) $\Leftrightarrow$ (d): Satz 3.2.32. Siehe auch [227].

(a) $\Leftrightarrow$ (e) : PALMER [230](1970), Seite 196.

(a) $\Leftrightarrow$ (f) $\Leftrightarrow$ (g): KULKARNI und LIMAYE [190](1980).

(a) $\Leftrightarrow$ (i) : LI [198](2003), Seite 77.

(i) $\Leftrightarrow$ (j) : LI [198](2003), Seite 165.

(i) $\Leftrightarrow$ (k) : LI [198](2003), Seite 163; siehe dazu auch den Beweis von (i) $\Leftrightarrow$ (j).

(k) $\Rightarrow$ (b): Nach dem bisher bewiesenen gilt (k) $\Leftrightarrow$ (i) $\Leftrightarrow$ (a) $\Rightarrow$ (c), also ist A symmetrisch.

(b) $\Rightarrow$ (k): Satz 3.2.37.

(i) $\Rightarrow$ (l) : LI [198](2003), Seite 78, Satz 5.1.3: Man wähle $B := A_{(\mathbb{C})}$ gemäß (i). Als - wähle man die kartesische Involution (siehe 3.1.8) auf dem Vektorraum $A_{(\mathbb{C})}$.

(l) $\Rightarrow$ (i) : LI [198](2003), Seite 78. $\square$

3.5.12 Bemerkung (INGELSTAM [145](1964)). Man beachte, dass zum Beispiel $\mathbb{C}_{\mathbb{R}}$, versehen mit der Involution $x^* = x$ und dem Betrag als Norm zwar eine Banach-*-Algebra über $\mathbb{R}$ ist, aber wegen $1 + i^*i = 0$ nicht symmetrisch ist, also keine C*-Algebra über $\mathbb{R}$ ist.

3.5.13 Beispiele (GOODEARL [117](1982)).

(a) $C(X, \mathbb{K})$, versehen mit der Supremumsnorm und der durch $f \mapsto f^*$, $f^*(x) := \overline{f(x)}$, $x \in X$, definierten Involution, ist für jeden kompakten Hausdorff-Raum X eine C*-Algebra über $\mathbb{K}$.

(b) Ist A eine C*-Algebra über $\mathbb{C}$, dann ist $A_{\mathbb{R}}$ eine C*-Algebra über $\mathbb{R}$.

(c) Sind für ein $n \in \mathbb{N}^{\times}$ $A_1, \dots, A_n$ C*-Algebren über $\mathbb{K}$, so ist $A_1 \times \dots \times A_n$, versehen mit der Supremumsnorm, eine C*-Algebra über $\mathbb{K}$. Siehe auch 3.5.23.

(d) Die Hamilton'sche Algebra $\mathbb{H}$ der Quaternionen über $\mathbb{R}$ mit der Bildung der konjugierten Quaternionen als Involution und dem Betrag als Norm ist eine C*-Algebra über $\mathbb{R}$.

(e) Sei X ein kompakter Hausdorffraum und $Y \subseteq X$ ein kompakter Unterraum. Dann ist $\{f \in C(X, \mathbb{C}) : f(Y) \subseteq \mathbb{R}\}$ eine C*-Algebra über $\mathbb{R}$.

(f) Jede kommutative C*-Algebra über $\mathbb{R}$ ist sowohl von der Form

$$\left\{ f \in C(X, \mathbb{C}) : f(\sigma(t)) = \overline{f(t)} \quad \text{für alle } t \in X \right\}$$

mit X als einen kompakten Hausdorffraum als auch von der Form

$$\left\{ f \in C_0(X, \mathbb{C}) : f(\sigma(t)) = \overline{f(t)} \quad \text{für alle } t \in X \right\}$$

mit X als einen lokal kompakten Hausdorffraum, wobei jeweils $\sigma: X \rightarrow X$ ein Homöomorphismus mit $\sigma^2 = \text{Id}_X$ ist. Siehe hierzu auch LI [198] (2003), Seite 83.

(g) Sei $n \in \mathbb{N}^{\times}$, dann ist $M_{\mathbb{K}}(n)$ eine C*-Algebra über $\mathbb{K}$ und $M_{\mathbb{H}}(n)$ eine C*-Algebra über $\mathbb{R}$. Dabei ist die Involution jeweils per $(a_{ij}) \mapsto (\overline{a_{ji}})$ definiert und die Norm erhält man durch Identifizierung von $M_{\mathbb{R}}(n)$ (bzw. $M_{\mathbb{C}}(n)$, $M_{\mathbb{H}}(n)$) mit $\mathcal{L}(\mathbb{R}^n)$ (bzw. $\mathcal{L}(\mathbb{C}^n)$, $\mathcal{L}(\mathbb{R}^{4n})$).

3.5.14 Beispiel. Nach dem Satz von Schauder (siehe etwa WERNER [291] (2002)) ist ein $T \in \mathcal{L}(X, Y)$, X, Y Banachräume, genau dann kompakt, wenn die duale Abbildung T^* kompakt ist. Da außerdem nach WERNER [291](2002), Seite 66, Satz II.3.2, $K(X, Y)$ ein

abgeschlossener Unterraum von $\mathcal{L}(X, Y)$ ist, folgt aus der Definition 3.4.3 der Hilbertraum-Adjungierten, dass $K(H)$ für jeden Hilbertraum H über $\mathbb{K}$ eine C^* -Algebra über $\mathbb{K}$ ist. Für jeden separablen Hilbertraum H über $\mathbb{C}$ besitzt $K(H)$ eine schiefminimal idempotente Projektion, siehe DAVIDSON [60](1996), Seite 36, Lemma I.10.1.

3.5.15 Bemerkung (LI [198](2003), Seite 78). Sei A eine C^* -Algebra über $\mathbb{R}$. Dann ist die C^* -Algebra $A_{(\mathbb{C})}$ eine Komplexifizierung nach LI. (Beachte 3.2.16.)

3.5.16 Bemerkung (INGELSTAM [145](1964), Theorem 17.6, Seite 267). Sei A eine Banach- $*$ -Algebra über $\mathbb{R}$. Dann sind äquivalent:

- (a) A ist $*$ -isomorph zu einer C^* -Algebra über $\mathbb{R}$.
- (b) A ist symmetrisch und die Norm erfüllt

$$\|x\|^2 \leq \beta \|x^*x\| \quad \text{für alle } x \in X \text{ und für ein konstantes } \beta \in \mathbb{R}.$$

3.5.17 Gelfand-Homomorphismen (LI [198](2003), Seiten 28, 41 und 78; GOODEARL [117](1982), Seite 95). Sei A eine kommutative Banachalgebra über $\mathbb{R}$. Als Involution $-$ auf $\mathbb{C}_{\mathbb{R}}$ sei die komplexe Konjugation gewählt. Man versetze den Vektorraum $\mathcal{L}(A, \mathbb{C}_{\mathbb{R}})$ mit der Vektorraum-Involution $\bar{\cdot}: \ell \mapsto \bar{\ell}$, wobei $\bar{\ell}(x) := \overline{\ell(x)}$ für alle $x \in A$.

Da mit $\ell \in \Gamma_A$ auch $\bar{\ell} \in \Gamma_A$ gilt, ist auf dem Vektorraum $C_0(\Gamma_A, \mathbb{C}_{\mathbb{R}})$ die Abbildung

$$\diamond: f \mapsto - \circ f \circ -$$

eine Vektorraum-Involution.

Der von $C_0(\Gamma_A, \mathbb{C}_{\mathbb{R}})$ bezüglich dieser Involution selbstadjungierte Teil

$$\left(C_0(\Gamma_A, \mathbb{C}_{\mathbb{R}})\right)_{(\diamond)}$$

ist eine Algebra. Die Abbildung $x \mapsto \hat{x}$ von Γ_A nach dieser Algebra ist ein Algebra-Homomorphismus, der hier im Folgenden als die Gelfand-Abbildung bezeichnet wird.

Auf dieser Algebra $\left(C_0(\Gamma_A, \mathbb{C}_{\mathbb{R}})\right)_{(\diamond)}$ ist die Abbildung $*$: $f \mapsto - \circ f$ eine Algebra-Involution. Ist nun A eine Banach- $*$ -Algebra über $\mathbb{R}$, so ist die Gelfand-Abbildung ein $*$ -Algebra-Homomorphismus. Falls A sogar eine C^* -Algebra über $\mathbb{R}$ ist, so ist die Gelfand-Abbildung ein $*$ -Algebra-Isomorphismus.

3.5.18 Satz. Sei A eine C^* -Algebra über $\mathbb{K}$ mit Eins. Dann ist ein Funktional $\ell \in A^\#$ genau dann ein Zustand (Definition 3.2.24(c)), wenn es ein unitaler Zustand (Definition 2.11.1) ist. Mit anderen Worten ist demnach die Menge der Zustände von A gleich $\Phi_A(e)$.

Beweis. Für den „genau dann“-Teil wird $\ell(e) = 1$ durch den Satz 3.2.29 geliefert. Das für den „wenn“-Teil entscheidende $\ell \geq 0$ liefert GOODEARL [117](1982): $\mathbb{K} = \mathbb{R}$: Seite 104; $\mathbb{K} = \mathbb{C}$: Seite 50. $\square$

3.5.19 Definition. Sei A eine C^* -Algebra über $\mathbb{K}$. Als eine naheliegende Verallgemeinerung von 3.4.21 definiert man: Ein Element $x \in A$ heißt eine *partielle Isometrie*, wenn x^*x eine Projektion ist. Ist $x \in A$ eine partielle Isometrie, so heißt x^*x die *Anfangsprojektion* und xx^* die *Endprojektion* von x . Im Fall, dass A eine Eins hat, heißt ein $x \in A$ eine *Isometrie*, wenn $x^*x = e$ gilt; ein $x \in A$ heißt eine *Co-Isometrie*, wenn xx^* eine Isometrie ist (also wenn $xx^* = e$ gilt).

3.5.20 Bemerkung. Sei A eine C^* -Algebra über $\mathbb{K}$. Als eine unmittelbare Folgerung der Definition einer C^* -Algebra folgt mit 3.4.21: Ein $x \in A$ ist genau dann eine partielle Isometrie, wenn $xx^*x = x$ gilt.

3.5.21. Sei A eine C^* -Algebra über $\mathbb{K}$. In der Gelfand-Naimark-Segal-Konstruktion ist jedem Zustand $\ell \in S_{A^*}$ ein Prähilbertraum $A/\mathfrak{N}_\ell$ über $\mathbb{K}$ zugeordnet, wobei $\mathfrak{N}_\ell := \{x \in A : \ell(x^*x) = 0\}$ und $\langle x + \mathfrak{N}_\ell, y + \mathfrak{N}_\ell \rangle := \ell(y^*x)$ für alle $x, y \in A$ — es sei an 3.2.23 erinnert. Bezeichnet H_ℓ den jeweiligen Hilbertraum über $\mathbb{K}$, den man mit der Vervollständigung von $A/\mathfrak{N}_\ell$ erhält, so kann man die Hilbert'sche direkte Summe H_L (siehe 3.4.11) von den Hilberträumen H_ℓ , $\ell \in L$, für eine Teilmenge L der Menge S aller Zustände von A betrachten. Für jedes solche nicht leere L erhält man eine $*$ -Darstellung (T_L, H_L) von A auf dem Hilbertraum H_L , die bei $L = S$ *universelle Darstellung* heißt; siehe zum Beispiel KADISON und RINGROSE [167](1983), Seite 281.

Jede $*$ -Darstellung von A ist genau dann algebraisch irreduzibel, wenn sie topologisch irreduzibel ist (MURPHY [217](2004), Seite 152 und LI [198] (2003), Seite 97); es ist daher üblich, im Rahmen von C^* -Algebren über $\mathbb{K}$ schlichtweg von irreduziblen $*$ -Darstellungen zu reden.

Der Gelfand-Naimark-Segal-Konstruktion entsprechend, heißen zwei Elemente $x, y \in A$ *orthogonal* (zueinander), falls $xy^* = y^*x = 0$ gilt. (Siehe diesbezüglich auch 4.7.3.) Mit anderen Worten: Zwei Elemente aus A sind genau dann orthogonal, wenn sie es bezüglich des Annihilatorsystems $(A, A; B; A)$, $B : (x, y) \mapsto y^*x$, sind. Man bemerke, dass somit als Spezialfall gilt: Zwei partielle Isometrien $u, v \in A$ sind genau dann orthogonal, wenn $(uu^*)(vv^*) = (u^*u)(v^*v) = 0$ gilt.

3.5.22. Für jeden nicht kompakten Hausdorffraum X , ist $C_0(X, \mathbb{K})$ mit $f^* := \bar{f}$ als Involution ([257]) eine C^* -Algebra über $\mathbb{K}$ ohne Eins. (DORAN und BELFI [77](1986), Seite 5.)

Nach einer auf YOOD (1956) zurückgehenden Idee kann man jede C^* -Algebra über $\mathbb{K}$ ohne Eins durch Anknüpfen des Skalarbereiches $\mathbb{K}$ zu einer Eins verhelfen.

Gemäß Bemerkung 3.5.6 ist jede C^* -Algebra über $\mathbb{K}$, die im Fall von $\mathbb{K} = \mathbb{C}$ mit ihrer intrinsischen C^* -Norm versehen sei, die eine Eins hat, unital.

Hat A keine Eins, dann ist zwar A^1 mit der Norm $(a, \lambda) \mapsto \|a\| + |\lambda|$, $a + \lambda \in A^1$ (siehe Bemerkung 3.2.20) eine $*$ -Banachalgebra über $\mathbb{K}$, aber im Allgemeinen ist diese Norm keine C^* -Norm.

Bei RICKART [246](1960), Seite 186, findet man mit Beweis folgende Ausführung: Bezeichnet L_a die durch $a \in A$ bestimmt Links-Multiplikation in A , dann ist $\varphi : a \mapsto L_a$ ein isometrischer Isomorphismus von A nach $\varphi(A) \subseteq \mathcal{L}(A)$. (Denn: $\|L_a\| = \sup\{\|L_ab\| : b \in B_A\} \leq \|a\|$ und $\|a\|^2 = \|a^*a\| = \|aa^*\| = \|L_aa^*\| \leq \|L_a\|\|a^*\| = \|L_a\|\|a\|$, also auch $\|a\| \leq \|L_a\|$.) Sei B die Unter algebra von $\mathcal{L}(A)$, die von $\varphi(A)$ und Id_A erzeugt wird. Dann besteht B aus allen Operatoren der Form $L_a + \lambda \text{Id}_A$, $a \in A$, $\lambda \in \mathbb{K}$. Setze $(L_a + \lambda \text{Id}_A)^* := L_{a^*} + \bar{\lambda} \text{Id}_A$. Dann ist $L_a + \lambda \text{Id}_A \mapsto L_{a^*} + \bar{\lambda} \text{Id}_A$ eine Konjugation auf B , die Operatornorm in B ist eine C^* -Algebranorm und $(a, \lambda) \mapsto L_a + \lambda \text{Id}_A$ ist ein isometrischer $*$ -Isomorphismus von A^1 nach B .

Somit ist hier

$$\|(a, \lambda)\| := \sup \{\|ab + \lambda b\| : b \in B_A\} \quad \text{für alle } a \in A, \lambda \in \mathbb{K} \quad (3.13)$$

eine C^* -Algebranorm auf A^1 , die die Norm von A fortsetzt. (PALMER [232](2001), Seite 812, für $\mathbb{K} = \mathbb{C}$ und INGELSTAM [145](1964), Seite 264, für $\mathbb{K} = \mathbb{R}$.)

Für den Fall $\mathbb{K} = \mathbb{C}$ hat man also, dass A^1 eine C^* -Algebra über $\mathbb{C}$ ist. Für den Fall $\mathbb{K} = \mathbb{R}$ folgt mit Satz 3.2.31, dass A^1 , versehen mit der Norm (3.13), eine C^* -Algebra über $\mathbb{R}$ ist.

Man beachte, dass bei dieser Konstruktion die intrinsische Norm von A als eine C^* -Norm angenommen wird. So gab es in den frühen 1960er Jahren noch keine solche Konstruktion für Banach- $*$ -Algebren über $\mathbb{C}$ ohne Eins, die ausgestattet sind mit einer schwachen C^* -Norm. Die dafür entsprechende Konstruktion gab 1967 B. J. Vowden an; siehe zum Beispiel DORAN und BELFI [77](1986), Seite 40, Satz 14.1.

3.5.23 (CONWAY [53](1999), Seite 6). Sei A_ν , $\nu \in I$, eine Familie von C^* -Algebren über $\mathbb{C}$. Dann sind auch $\ell_\infty(\bigoplus_{\nu \in I} A_\nu)$ und $c_0(\bigoplus_{\nu \in I} A_\nu)$ C^* -Algebren über $\mathbb{C}$, wobei die Operationen kanonisch koordinatenweise definiert sind.

3.5.24 Satz (DORAN und BELFI [77](1986), Seite 157). Sei A eine Banach- * -Algebra über $\mathbb{C}$. Dann sind äquivalent:

- (a) $\|a^*a\| = \|a\|^2$ für alle $a \in A$.
- (b) $\|a^*a\| = \|a^*\|\|a\|$ für alle $a \in A$.
- (c) $\|n^*n\| = \|n^*\|\|n\|$ für alle $n \in A_{\text{normal}}$.
- (d) $\rho(a^*a) = \|a\|^2$ für alle $a \in A$.
- (e) $\rho(a^*a) \geq \|a\|^2$ für alle $a \in A$.

Wenn A unital ist, dann sind zu (a) bis (f) zusätzlich äquivalent:

- (g) $\|u\| = 1$ für alle $u \in A_{\text{unitär}}$.
- (h) $\|\exp(is)\| = 1$ für alle $s \in A_{(*)}$.

3.5.25 Satz (PALMER [232](2001), Seite 1181, Theorem 11.2.5). Sei $(A, \|\cdot\|)$ eine Banach- * -Algebra über $\mathbb{C}$ mit Eins. Dann sind äquivalent:

- (a) A ist eine C^* -Algebra über $\mathbb{C}$ und $\|\cdot\|$ ist ihre intrinsische C^* -Norm.
- (b) $\|\cdot\|$ ist eine C^* -Norm.
- (c) $A_{\text{unitär}} \subseteq B_A$.
- (d) $\exp(iA_{(*)}) \subseteq B_A$. (Siehe auch 3.2.19)
- (e) $A_{(*)} \subseteq A_{\text{herm}}$.

3.5.26 Satz. Sei A eine C^* -Algebra über $\mathbb{C}$. Dann gilt:

- (a) A ist symmetrisch.

Wenn A eine Eins hat, dann gilt zusätzlich:

- (b) $\text{cl co exp}(iA_{(*)}) = \text{cl co } A_{\text{unitär}} = B_A$.
- (c) $A_{(*)} = A_{\text{herm}}$.
- (d) Für die intrinsische C^* -Norm $\|\cdot\|$ gilt:

$$\begin{aligned} \|a\| &= \sup\{\omega(a^*a) : e \dashv \omega\}^{1/2} && \text{für alle } a \in A, \\ \|s\| &= \sup\{|\omega(s)| : e \dashv \omega\} && \text{für alle } s \in A_{(*)}. \end{aligned}$$

Beweis. (a) Siehe RICKART [246](1960), Seite 243. (b) Siehe PALMER [232] (2001), Seite 1177, und UPMEIER [279](1985), Seite 256, Theorem 15.24. (d) Siehe PALMER [232](2001), Seiten 947 und 951. $\square$

3.5.27 Satz (LI [198](2003), Seite 157). Sei A eine C^* -Algebra über $\mathbb{R}$ mit Eins. Dann gilt $\text{cl}\left(\text{co}(\cos(A_{(*)}) \cdot \exp(A_{(-*)}))\right) = B_A$.

3.5.28. Damit eine unitale Banach- $*$ -Algebra eine C^* -Algebra über $\mathbb{C}$ sein kann, muss sie genügend hermitesche Elemente besitzen: Der Satz von Vidav-Palmer besagt, dass eine unitale Banachalgebra A über $\mathbb{C}$ genau dann eine Involution zulasse, bezüglich derer sie eine C^* -Algebra über $\mathbb{C}$ sei, wenn sie gleich der komplexen linearen Hülle von A_{herm} sei und dies wiederum sei genau dann der Fall, wenn mit der üblichen Identifikation gälte: $A = A_{\text{herm},(\mathbb{C})}$, A also eine Vidav-Algebra wäre. Die besagte Involution ist dann gerade die kartesische Involution von $A_{\text{herm},(\mathbb{C})}$ (siehe Bemerkung 3.1.8 und 3.2.16). Genauer gilt:

3.5.29 Satz (PALMER [232](2001), Seite 951, Theorem 9.5.9). *Sei $(A, \|\cdot\|)$ ein Banachraum über $\mathbb{C}$. Wähle ein Element aus S_A und nenne es 1. Ein Produkt auf A , das $(A, \|\cdot\|)$ zu einer Banachalgebra mit 1 als Eins macht, soll ein »gutes Produkt« heißen. Dann sind die folgenden drei Aussagen äquivalent:*

- (a) *$\text{span}(A_{\text{herm}})$ ist dicht in A und es kann auf A ein gutes Produkt definiert werden.*
- (b) *Ein Produkt und eine Involution können auf A definiert werden, so dass A eine C^* -Algebra über $\mathbb{C}$ mit Eins 1 ist und die Norm $\|\cdot\|$ auf ihr eine C^* -Algebranorm ist.*
- (c) *Auf A kann ein gutes Produkt definiert werden, und zu jedem solchen Produkt gehört genau eine Involution, welche A zu einer C^* -Algebra über $\mathbb{C}$ mit genau $\|\cdot\|$ als einziger C^* -Algebranorm macht. $A_{\mathbb{R}}$ hat die direkte Summenzerlegung $A_{\mathbb{R}} = A_{\text{herm}} \dot{+} i \cdot A_{\text{herm}}$ und die Involution ist definiert per $h + ik \mapsto h - ik$ für alle $h, k \in A_{\text{herm}}$.*

3.5.30 Bemerkung (PALMER [232](2001), Seite 951). Der Satz 3.5.29 zeigt, dass eine unitale Banachalgebra A über $\mathbb{C}$ genau dann isometrisch isomorph zu einer C^* -Algebra über $\mathbb{C}$ ist, wenn die Gestalt der Einheitskugel B_A von A die gleiche ist, wie die von der abgeschlossenen Einheitskugel einer C^* -Algebra über $\mathbb{C}$ mit ihrer intrinsischen C^* -Algebranorm. Man muss nichts über die multiplikative Struktur wissen, außer, welches Element die Eins ist. Die besagte Gestalt ist vollkommen bestimmt von einer infinitesimalen Umgebung der Eins auf der Einheitssphäre S_A von A , da 2.11.27(b) zeigt, dass A_{herm} definiert werden kann als die Menge der Elemente, die hermitesch nach Vidav sind. Die Gestalt von B_A in einer C^* -Algebra über $\mathbb{C}$ bestimmt eindeutig die Involution. Siehe auch 3.5.54.

3.5.31 Bemerkung (DORAN und BELFI [77](1986), Seite 198). Hat man also einen Banachraum X über $\mathbb{C}$ und ist 1 ein Element aus S_X , dann macht im Fall, dass $X_{\mathbb{R}} = X_{\text{herm}} + i \cdot X_{\text{herm}}$ gilt, jedes gute Produkt X zwangsläufig zu einer C^* -Algebra über $\mathbb{C}$.

3.5.32 Bemerkung (BONSALL [28](1970), Seite 268). Interessant an dem Satz von Vidav-Palmer ist, dass, obwohl es eine Aussage über Operatoren auf Hilberträumen ist, es weder das Konzept von Hilberträumen noch Involutionen verwendet, sondern nur den numerischen Wertebereich benötigt.

3.5.33 Extremalpunkte (LI [197](1992), Seiten 94, 96; LI [198](2003), Seite 99; PALMER [232](2001), Seite 959, Theorem 9.5.16(a); MATHIEU [208](1998), Seite 280; KADISON und RINGROSE [167](1983), Seite 163).

Sei A eine C^* -Algebra über $\mathbb{K}$, die im Fall von $\mathbb{K} = \mathbb{C}$ mit ihrer intrinsischen C^* -Norm versehen sei. Dann ist ein $x \in B_A$ genau dann ein Extremalpunkt von B_A , wenn

$$(1 - x^*x)A(1 - xx^*) = \{0\} \quad (3.14)$$

gilt. Bezeichnet man für $\ell, r \in A$ mit $M_{\ell,r}$ das per

$$M_{\ell,r}(x) := \ell x r \quad \text{für alle } x \in A$$

definierte Element aus $\mathcal{L}(A)$, so ist also

$$\text{ex}(B_A) = \{x \in B_A : M_{1-x^*x, 1-xx^*} = 0\}.$$

Im Fall, dass A eine Eins hat, ist diese Eins ein Extrempunkt von B_A . Ist x ein Extrempunkt von B_A , so ist x eine partielle Isometrie. Für jeden Hilbertraum H über $\mathbb{K}$ ist ein $T \in \mathcal{L}(H)$ genau dann ein Extrempunkt von $B_{\mathcal{L}(H)}$, wenn T oder T^* eine Isometrie ist. (Siehe auch die Schlussbemerkung in 3.4.21 und MATHIEU [208](1998), Seite 293.) $\text{ex}(B_A)$ ist genau dann nicht leer, wenn A eine Eins besitzt, nämlich $e = x^*x + xx^* - x^*xx^*$, wobei $x \in A$ ein beliebiger Extrempunkt von B_A ist. Diese zuletzt genannte Äquivalenz gilt auch, wenn A eine *-Algebra über $\mathbb{C}$ ist, die mit einer C^* -Algebranorm versehen ist, und dann besitzt des Weiteren A genau dann eine Eins, wenn die Menge der Vertice von B_A nicht leer ist (siehe auch 2.11.31).

Für jede C^* -Algebra A über $\mathbb{C}$ mit Eins gilt

$$\begin{aligned} \{x \in S_A : A^* &= \text{span } \Phi_A(x)\} = \{x \in S_A : x \text{ ist ein Vertex von } B_A\} \\ &= A_{\text{unitär}} \subseteq \text{ex}(B_A); \end{aligned} \quad (3.15)$$

(siehe auch 4.7.3) und mit Satz 3.5.26(b) gilt dann also

$$B_A = \text{cl co ex}(B_A).$$

Obwohl im unendlichen Fall B_A nicht kompakt ist, gilt hier für B_A also eine an den Satz von Krein-Milman erinnernde Aussage.

Siehe auch AKEMANN und WEAVER [5](2002).

3.5.34 (PEDERSEN [233](1979), Seite 13). Sei A eine C^* -Algebra über $\mathbb{C}$, die mit ihrer intrinsischen C^* -Norm versehen sei. Seien p und q zwei idempotente Elemente aus $\text{Pos}(*\sigma; A)^\times$. Dann ist ein $x \in pAq$ genau dann ein Extrempunkt von B_{pAq} , wenn

$$(p - x^*x)A(q - xx^*) = \{0\}$$

gilt; also

$$\text{ex}(B_{pAq}) = \{x \in B_{pAq} : M_{p-x^*x, q-xx^*} = 0\}.$$

Insbesondere ist jede Projektion $p \in A^\times$ ein Extrempunkt von B_{pAp} .

3.5.35 (ALFSEN und SHULTZ [8](2001), Seite 75). Sei A eine C^* -Algebra über $\mathbb{C}$ mit Eins. Dann sind die Extrempunkte von $B_A \cap \text{Pos}(*\sigma; A)$ genau die von 0 verschiedenen Projektionen in A . Des Weiteren gilt:

$$\text{ex}(B_A \cap A_{(*)}) = \{s \in A : \text{Es existiert eine Projektion } p \in A \text{ mit } s = 2p - e\}.$$

In diesem Zusammenhang sei an 3.1.10 erinnert.

3.5.36 (MATHIEU [208](1998), Seite 292). Sei A eine C^* -Algebra über $\mathbb{C}$. Für jede Projektion $p \in A$ ist pAp eine C^* -Unteralgebra von A .

Sei $T \in A$ eine partielle Isometrie. Setze p als die Anfangsprojektion und q als die Endprojektion von T , also $p = T^*T$ und $q = TT^*$. Dann ist M_{T,T^*} (siehe 3.5.33) ein *-Isomorphismus von der C^* -Unteralgebra pAp auf die C^* -Unteralgebra qAq .

Falls A unital ist, sind die Projektionen genau die positiven, partiellen Isometrien von A (AKEMANN und WEAVER [5](2002)). Vergleiche 4.7.3.

3.5.37 Bemerkung (BONSALL und DUNCAN [31](1973)). Sei A eine unitale Algebra über $\mathbb{C}$. Dann gilt, beziehungsweise definiert man:

(a) (ebd., Seite 100) Der Dualraum A^* ist gleich der komplexen linearen Hülle von $\Phi_A(e)$.

(b) (ebd., Seite 103) Man definiert den sogenannten reellen linearen span von $\Phi_A(e)$ als die Menge

$$(A^*)_H := \{\lambda f - \mu g : f, g \in \Phi_A(e), \lambda, \mu \in \mathbb{R}^+\}.$$

Bezüglich der Motivation dieser Definition siehe MOORE [215](1971).

(c) Für den Dualraum A^* gilt: $A^* = (A^*)_H + i \cdot (A^*)_H$.

(d) Die Summe in (c) ist genau dann direkt, wenn A eine Involution zulässt, bezüglich der A eine C^* -Algebra über $\mathbb{C}$ ist, das heißt, genau dann, wenn die hermiteschen Elemente in A die Punkte des Dualraumes A^* trennen. (MOORE [215](1971), Seite 103 und DORAN und BELFI [77](1986), Seite 181.)

(e) Der Satz in (d) besagt anschaulich, dass, damit eine Banach- $*$ -Algebra über $\mathbb{C}$ eine C^* -Algebra über $\mathbb{C}$ sein kann, sie nicht zu viele Funktionale in der Menge von (b) liegen haben darf. ($(A^*)_H \cap i \cdot (A^*)_H$ darf nur die 0 enthalten.)

(f) Ist A eine C^* -Algebra über $\mathbb{C}$ und beachtet man, dass dann der Dualraum A^* gemäß Bemerkung 3.1.6 eine kanonische Involution trägt, dann gilt: $(A^*)_H = (A^*)_{(*)}$.

3.5.38. Die vorstehende Bemerkung 3.5.37 findet man bei PALMER [232](2001), Seite 953f, in einer ähnlichen Form wie Satz 3.5.29.

3.5.39 Satz. Sei A eine C^* -Algebra über $\mathbb{K}$. Sei $\ell \in S_{A^*}$ mit $\ell \geq 0$. Dann ist ℓ im Fall von $\mathbb{K} = \mathbb{C}$ selbstadjungiert. Im Fall $\mathbb{K} = \mathbb{R}$ ist zum Beispiel das Funktional $\ell: \mathbb{C}_{\mathbb{R}} \rightarrow \mathbb{R}$, $a + bi \mapsto a + b$ nicht selbstadjungiert.

Beweis. Für die Behauptung im komplexen Fall siehe GOODEARL [117](1982), Seite 49. Die Behauptung für den reellen Fall findet sich bei GOODEARL [117] (1982), Seite 102. $\square$

3.5.40 Satz. Sei A eine C^* -Algebra über $\mathbb{R}$. Dann ist $\text{Pos}(*\sigma; A)$ ein punktierte spitzer konvexer Kegel. Ist $x \in \text{Pos}(*\sigma; A)$, dann existiert genau ein $y \in \text{Pos}(*\sigma; A)$ mit $x = y^2$. Ist $x \in A$, dann gilt genau dann $x \in \text{Pos}(*\sigma; A)$, wenn ein $y \in A$ mit $x = y^*y$ existiert.

Beweis. LI [198](2003), Seite 83. $\square$

3.5.41 Satz. Sei A eine C^* -Algebra über $\mathbb{K}$. Dann gilt:

$$\text{Pos}(*; A) = \text{Pos}(*\sigma; A).$$

Beweis. Für den Fall $\mathbb{K} = \mathbb{R}$ folgt dies aus Satz 3.5.40, wobei man für die Inklusion „ $\subseteq$ “ lediglich noch die Bemerkung 2.2.13 zu beachten braucht. Für den Fall $\mathbb{K} = \mathbb{C}$ ist die Aussage standard, siehe zum Beispiel PALMER [232](2001), Seite 870, Theorem 9.2.16(b). $\square$

3.5.42. Als eine unmittelbare Folgerung des Satzes 3.5.41 hat man: Jede treue $*$ -Darstellung einer C^* -Algebra über $\mathbb{K}$ respektiert die durch den punktierten konvexen Kegel $\text{Pos}(*\sigma; A)$ induzierte partielle Ordnung. Dies ist im Hinblick auf 3.7.36 von Interesse.

3.5.43 Satz. Sei A eine C^* -Algebra über $\mathbb{C}$ und $a \in A$. Sei T eine treue $*$ -Darstellung von A auf einem Hilbertraum H über $\mathbb{C}$ (siehe Bemerkung 3.5.2). Dann sind äquivalent:

(a) $a \in \text{Pos}(*\sigma; A)$. (Definition 3.2.13)

(b) $a \in \text{Pos}(*; A)$. (Definition 3.2.24)

- (c) $a \in A_{(*)}$ und $\omega(a) \geq 0$ für alle Zustände ω von A .
- (d) $a = b^*b$ für ein $b \in A$. („ a besitzt eine Quadratwurzel“)
- (e) $a = b^2$ für ein $b \in \text{Pos}(*; A)$.
- (f) $a = b^2$ für genau ein $b \in \text{Pos}(*; A)$.
- (g) $T_a \in \text{Pos}(*; \mathcal{L}(H))$.
- (h) $T_a \in \text{Pos}(W; \mathcal{L}(H))$. (Definition 3.4.12)

Hat A eine Eins, dann sind zusätzlich zu (a) bis (h) äquivalent:

- (i) $a \in \text{Pos}(V; A)$.
- (j) $a \in A_{(*)}$ und $\|\lambda - a\| \leq \lambda$ für alle $\lambda \geq \|a\|$.
- (k) $a \in A_{(*)}$ und $\|\lambda - a\| \leq \lambda$ für ein $\lambda \geq \|a\|$.
- (l) $a \in \text{Pos}(V\sigma; A)$.

Beweis. (f) Siehe zum Beispiel MURPHY [217](2004) oder DIXMIER [75](1977).

(i) $\Leftrightarrow$ (l): Satz 2.12.12.

(i), (k): AKEMANN und WEAVER [5](2002) □

3.5.44 (KADISON und RINGROSE [167](1983), Seite 241, Satz 4.1.5). Sei A eine C^* -Algebra über $\mathbb{C}$ mit einer Eins e , B eine C^* -Unteralgebra von A mit $e \in B$. Dann gilt für alle $x \in B$ die Gleichung $\sigma_B(x) = \sigma_A(x)$. Folglich hat man $\text{Pos}(*\sigma; B) = B \cap \text{Pos}(*\sigma; A)$.

3.5.45. Sei A eine C^* -Algebra über $\mathbb{K}$ und $a \in A$. Dann ist $a^*a \in \text{Pos}(*; A)$ und nach den Sätzen 3.5.41, 3.5.40 und 3.5.43 existiert genau ein $b \in \text{Pos}(**\sigma; A)$ mit $a^*a = b^2$; dieses b heißt der *Absolutbetrag* von a , wird mit $|a|$ bezeichnet und man schreibt formal $|a| = \sqrt{a^*a}$.

3.5.46 Bemerkung. Sei $A \subseteq \mathcal{L}(H)$ eine C^* -Algebra von Operatoren über $\mathbb{K}$ auf einem Hilbertraum H , $T \in \text{Pos}(**\sigma; A)$ und $x \in H$ mit $\langle Tx, x \rangle = 0$. Dann ist $Tx = 0$.

Beweis. Nach den Sätzen 3.5.40 und 3.5.43(e) existiert ein $S \in \text{Pos}(**\sigma; A)$ mit $T = SS$. Also $0 = \langle Tx, x \rangle = \langle Sx, Sx \rangle = \|Sx\|^2$, $Sx = 0$, $Tx = SSx = 0$. □

3.5.47 Ordnungsbeschränkte Teilmengen (KADISON und RINGROSE [167] (1986), Seite 249; MATHIEU [208](1998), Seite 268; WEIDMANN [288](2000), Seite 108).

Sei A eine C^* -Algebra über $\mathbb{C}$ mit Eins. Dann gilt für alle Elemente x des geordneten Vektorraumes $(A_{(*)}, \text{Pos}(**\sigma; A))$:

- (a) $-\|x\|e \leq x \leq \|x\|e$ und
- (b) $\|x\| = \inf\{t \in \mathbb{R}^+ : -te \leq x \leq te\}$.

Folglich gilt dann wegen $-y \leq x \leq y \Rightarrow -\|y\|e \leq -y \leq x \leq y \leq \|y\|e$ für alle $x, y \in A_{(*)}$ die Implikation

- (c) $-y \leq x \leq y \Rightarrow \|x\| \leq \|y\|$.

Da in jedem prägeordneten Vektorraum für alle seine Elemente x die Äquivalenzen $0 \leq x \Leftrightarrow 0 \leq 2x \Leftrightarrow -x \leq x$ gelten, gilt insbesondere $0 \leq x \leq y \Rightarrow -y \leq 0 \leq x \leq y \Rightarrow \|x\| \leq \|y\|$. Aus $x \in [0, y] = \{z \in A_{(*)} : 0 \leq z \leq y\}$ folgt also $\|x\| \in [0, \|y\|]$ für alle $x, y \in A_{(*)}$. Ist andererseits eine Menge $M \subseteq A_{(*)}$ Norm-beschränkt, etwa durch zwei reelle Zahlen k und ℓ , so dass für alle $x \in M$ die Aussage $k \leq \|x\| \leq \ell$ gilt, so gilt nach (a) $-\ell e \leq x \leq \ell e$.

Nennt man eine Menge $M \subseteq A_{(*)}$ *ordnungsbeschränkt*, wenn es zwei Elemente $x, y \in A_{(*)}$ gibt, so dass für alle $z \in M$ die Aussage $x \leq z \leq y$ gilt, so kann man also sagen: Die Norm-beschränkten Teilmengen von $A_{(*)}$ stimmen mit den ordnungsbeschränkten Teilmengen von $A_{(*)}$ überein.

Man bemerke, dass wegen 3.4.9(b) alle hier gemachten Aussagen insbesondere für $A = \mathcal{L}(H)$, H ein beliebiger Hilbertraum über $\mathbb{K}$, gelten.

3.5.48 Kommutativität. Sei A eine C^* -Algebra über $\mathbb{C}$. Dann sind die folgenden drei Aussagen äquivalent: (i) A ist kommutativ; (ii) $(A_{(*)}, \text{Pos}(*; A))$ ist ein Verband; (iii) Der Dualraum $(A_{(*)})^*$ ist ein Verband.

Beweis. Diese Aussage findet sich bei TAKESAKI [274](2001), Seite 25, Übung 1, mit dem Hinweis auf die folgenden zwei Quellen: FUKAMIYA, MISONOU und TAKEDA [104](1954) und SHERMAN [259](1951). $\square$

Ist H ein Hilbertraum über $\mathbb{C}$ mit Dimension echt größer als eins, so ist $\mathcal{L}(H)$ nicht kommutativ.

In BLACKADAR [26](2005), Seite 20, I.5.1.4, findet man das folgende Beispiel. Sei $H = \mathbb{C}^2$, also $\mathcal{L}(H) \simeq M_{\mathbb{C}}(2)$. Betrachte die beiden folgenden Elemente aus $\text{Pos}(*; M_{\mathbb{C}}(2))$: $p := \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}$, also die Projektion auf $\mathbb{C} \times \{0\}$, und $q := \frac{1}{2} \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix}$, also die Projektion auf die Diagonale $\Delta := \{(x, x) \in \mathbb{C}^2 : x \in \mathbb{C}\}$. Dann sind $e := \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$ und $p + q = \frac{1}{2} \begin{pmatrix} 3 & 1 \\ 1 & 1 \end{pmatrix}$ obere Schranken für p und q . Aber es gibt kein $r \in \text{Pos}(*; M_{\mathbb{C}}(2))$ mit $p \leq r$, $q \leq r$, $r \leq e$ und $r \leq p + q$.

KADISON [165](1951), Seite 507, Theorem 6, zeigt, dass für jeden Hilbertraum über $\mathbb{C}$ der reelle Vektorraum $\mathcal{L}(H)_{(*)}$ ein Antiverband ist. Somit ist für Hilberträume über $\mathbb{C}$ mit Dimension echt größer als eins $\mathcal{L}(H)_{(*)}$ nicht nur kein Verband, sondern sogar so weit wie nur möglich davon entfernt, ein Verband zu sein. Nichtsdestotrotz gilt für alle Hilberträume über $\mathbb{C}$ die in 3.5.47 aufgeführte Implikation $0 \leq x \leq y \Rightarrow \|x\| \leq \|y\|$ für alle $x, y \in \mathcal{L}(H)_{(*)}$, welche eine der wesentlichen Eigenschaften der sogenannten Banach-Verbände ist, siehe zum Beispiel FREMLIN [99](2004), Definition 354A.

Archbold zeigte 1972: Für A eine C^* -Algebra über $\mathbb{C}$ ist $A_{(*)}$ genau dann ein Antiverband, wenn $\{0\}$ ein Prim- $\mathbb{C}$ -Ring (siehe 2.1.43) ist.

3.5.49 Definition (PALMER [232](2001), Seite 990). Sei A eine $*$ -Algebra über $\mathbb{C}$. Dann heißt A

- *regulär*, falls jede $*$ -Unteralgebra von A Jacobson-halbeinfach ist;
- *sehr proper* (im Englischen *very proper*), falls für alle $x \in A$ die Implikation $-x^*x \in \text{Pos}(*; A) \Rightarrow x = 0$ gilt;
- *proper* (im Englischen *proper*), falls für alle $x \in A$ die Implikation $x^*x = 0 \Rightarrow x = 0$ gilt;
- *halbproper* (im Englischen *semiproper*), falls für alle $h \in A_{(*)}$ die Implikation $h^2 = 0 \Rightarrow h = 0$ gilt;

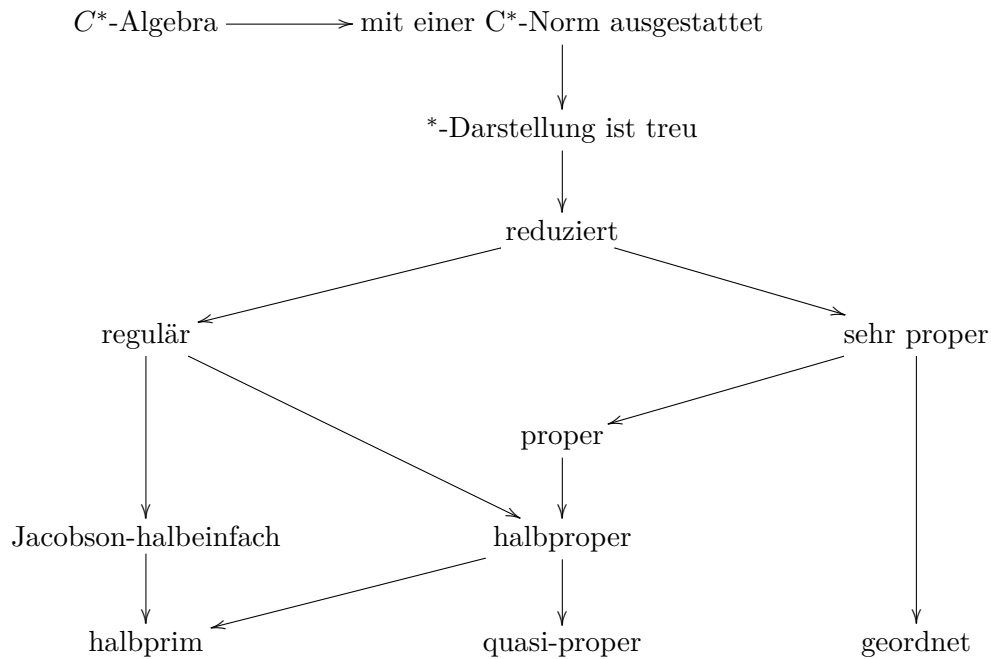
- *quasi-proper* (im Englischen *quasi-proper*), falls für alle $x \in A$ die Implikation $x^*x = 0 \Rightarrow xx^* = 0$ gilt;
- *geordnet* (im Englischen *ordered*), falls $\text{Pos}(*; A) \cap (-\text{Pos}(*; A)) = \{0\}$ gilt.

3.5.50. Sei A eine normierte $*$ -Algebra über $\mathbb{C}$. Dann sind äquivalent:

- A ist regulär.
- Für alle $h \in A_{(*)}$ gilt die Implikation $\rho(h) = 0 \Rightarrow h = 0$.
- Jede maximale, kommutative $*$ -Unteralgebra von A ist Jacobson-halbeinfach.

Zum Beweis beachte man Bemerkung 2.12.4(b) und siehe dann PALMER [232] (2001), Seite 991.

3.5.51 (PALMER [232](2001), Seiten 979 und 992). Im Rahmen von $*$ -Algebren sei an den in 3.4.29 eingeführten Begriff des reduzierenden Ideals erinnert. Für jede $*$ -Algebra über $\mathbb{C}$ gelten die folgenden Implikationen:



3.5.52 (LI [198](2003), Seite 79). Jede C^* -Algebra über $\mathbb{R}$ ist Jacobson-halbeinfach und somit insbesondere auch halbprim.

3.5.53 Approximative Eins (MURPHY [217](2004), Seite 77; LI [198] (2003), Seite 84; CONWAY [53](1999), Seite 18).

Ist A eine C^* -Algebra über $\mathbb{K}$, so ist es zur Vereinfachung von Berechnungen üblich, von einer approximativen Linkseins (a_ν) von A zusätzlich zu fordern, dass es ein aufsteigendes Netz in $\text{Pos}(*\sigma; A) \cap B_A$ ist; approximative Rechtseins und approximative Eins entsprechend. Diese Handhabung sei hiermit auch hier so vereinbart.

Durch Übergang zu den konjugierten Elementen erkennt man bei Betrachtung der Nullnetzbedingung, dass in C^* -Algebren über $\mathbb{K}$ sowohl jede approximative Links- als auch jede approximative Rechtseins eine approximative Eins ist.

Als einen Spezialfall eines Resultats von SEGAL aus dem Jahr 1947 hat man den Satz, dass jede C^* -Algebra A über $\mathbb{K}$ eine approximative Eins hat. Für einen Beweis siehe MURPHY [217](2004), Seite 78 und LI [198](2003), Seite 84. Zirka 1968 entdeckte DIXMIER, dass $\text{Pos}(*\sigma; A) \cap \text{int}(B_A)$ eine approximative Eins ist; sie wird die *kanonische*

approximative Eins genannt. Siehe auch FILLMORE [96](1996), Seite 16; KADISON und RINGROSE [167](1983), Lemma 4.2.11; PALMER [232](2001), Theorem 9.2.18; PEDERSEN [233](1979), Seite 11.

3.5.54 Eindeutigkeit der Involution. Vorab sei an 3.5.30 erinnert. Bezeichne $(A, *, \|\cdot\|)$ eine C^* -Algebra über $\mathbb{K}$, wobei $*$ ihre Involution und $\|\cdot\|$ ihre, im Fall von $\mathbb{K} = \mathbb{C}$ intrinsische, C^* -Norm sei.

Gemäß LI [198](2003), Seite 107, Theorem 5.6.5., gilt: Ist $(A, \|\cdot\|)$ eine Banachalgebra über $\mathbb{K}$ und sind $*$ und $\diamond$ zwei Involutionen auf A derart, dass $(A, *, \|\cdot\|)$ und $(A, \diamond, \|\cdot\|)$ zwei C^* -Algebren über $\mathbb{K}$ sind, so stimmen $*$ und $\diamond$ überein.

Für den Fall $\mathbb{K} = \mathbb{C}$ hat dies bereits RICKART [246](1960), Seite 247, Theorem 4.8.18., wie folgt festgestellt: Ist $(A, *, \|\cdot\|)$ eine C^* -Algebra über $\mathbb{C}$, so ist jede Involution $\diamond$ auf A , mit der die Norm $\|\cdot\|$ die C^* -Bedingung auf A erfüllt, identisch zur Involution $*$. Fordert man in dieser Formulierung die Involution $\diamond$ als symmetrisch, so gilt also die so modifizierte Aussage von Rickart auch für den Fall $\mathbb{K} = \mathbb{R}$.

3.5.55 Satz. *Jede C^* -Algebra über $\mathbb{C}$, deren zugrunde liegender Banachraum reflexiv ist, ist endlichdimensional.*

Beweis. Nach MEGGINSON [211](1998), Korollar 2.8.11 zum sogenannten Satz von Eberlein-Šmulian, ist jeder reflexive normierte Vektorraum schwach-Folgen-vollständig. SAKAI [252](1964), Seite 661, Satz 2, zeigt, dass jede schwach-Folgen-vollständige C^* -Algebra über $\mathbb{C}$ endlichdimensional ist. Siehe auch KADISON und RINGROSE [167](1986), Seite 772, Übung 10.5.17. $\square$

Es gilt aber noch mehr:

3.5.56 Satz (BECERRA GUERRERO, LÓPEZ PÉREZ und RODRÍGUEZ PALACIOS [17](2003), Seite 756, Theorem 2.5). *Jede C^* -Algebra über $\mathbb{C}$, deren abgeschlossene Einheitskugel eine relativ schwach offene, nicht leere Teilmenge (also zum Beispiel eine Scheibe) mit Durchmesser kleiner zwei enthält, ist endlichdimensional. Insbesondere ist also jede C^* -Algebra über $\mathbb{C}$, die die Radon-Nikodým-Eigenschaft aufweist, endlichdimensional.*

3.5.57 Atome in C^* -Algebren. Sei A eine C^* -Algebra über $\mathbb{K}$ und $p \in A^\times$ eine Projektion. p heißt ein *Atom* von A oder *atomar*, wenn die Gleichung $(pAp)_{(*)\mathbb{R}} = \mathbb{R}p$ gilt. A heißt *nicht atomar*, wenn A keine atomare Projektion enthält. Offensichtlich ist jede schiefminimal idempotente Projektion von A ein Atom von A .

3.6 Operator-Topologien

3.6.1 Bemerkung und Definition (WEIDMANN (2000), Seite 149; DUNFORD und SCHWARTZ [79](1963), Abschnitt XI.9). Seien H , H_1 und H_2 Hilberträume über $\mathbb{K}$ und sei $\langle \cdot, \cdot \rangle$ das Skalarprodukt von H_2 . Sei $1 \leq p < \infty$, dann betrachte man den Unterraum von $\mathcal{L}(H_1, H_2)$, der aus allen T besteht, für die $(\langle Tx_n, y_n \rangle)_{n \in \mathbb{N}} \in \ell_p$ gilt für alle orthonormalen Folgen $(x_n)_{n \in \mathbb{N}} \subseteq H_1$ und $(y_n)_{n \in \mathbb{N}} \subseteq H_2$. Dieser Unterraum ist ein Banachraum über $\mathbb{K}$ mit der sogenannten Schattenorm

$$\|\cdot\|_p: T \mapsto \sup \left\{ \left(\sum_{\nu} |\langle Tx_{\nu}, y_{\nu} \rangle|^p \right)^{1/p} \right\},$$

wobei das Supremum über alle Orthonormalsysteme $\{x_{\nu}\} \subseteq H_1$ und $\{y_{\nu}\} \subseteq H_2$ von H_1 beziehungsweise H_2 genommen wird, und heißt die *Schattenklasse* $C_p(H_1, H_2)$. Bei $H = H_1 = H_2$ schreibt man $C_p(H)$ für $C_p(H_1, H_2)$. $C_{\infty}(H)$ wird als $K(H)$ definiert.

SCHWARTZ und DUNFORD [79](1963), Seite 1100: Für $\mathbb{K} = \mathbb{C}$ gilt: Für alle $1 \leq p \leq \infty$ ist $C_p(H)$ eine Banachalgebra über $\mathbb{C}$, die im Fall von unendlichdimensionalem H keine Eins hat.

3.6.2 Spur II. Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{K}$. Für $1 \leq p \leq q \leq \infty$ gilt $C_p(H) \subseteq C_q(H)$, $\|\cdot\|_p \geq \|\cdot\|_q$; gilt zusätzlich $1 = \frac{1}{p} + \frac{1}{q}$, so gilt $\|ST\|_1 \leq \|S\|_p \|T\|_q$ für alle $S \in C_p(H)$, $T \in C_q(H)$.

Sei $\{b_\nu \in H : \nu \in I\}$, I eine Indexmenge, eine beliebige Orthonormalbasis für H . Dann ist die Abbildung

$$\text{spur}: \text{Pos}(*W; \mathcal{L}(H)) \rightarrow [0, \infty], \quad T \mapsto \sum_{\nu \in I} \langle Tb_\nu, b_\nu \rangle$$

unabhängig von der gewählten Orthonormalbasis und heißt *Spur* (im Englischen *trace*). (Bezüglich der Definition von Summen mit überabzählbarer Indexmenge siehe 2.4.1. Zum Beweis beachte Satz 3.4.16 und PALMER [232](2001), Seite 826, Satz 9.1.30.)

Ein $T \in \mathcal{L}(H)$ heißt *Hilbert-Schmidt-Operator*, wenn $\text{spur}(T^*T) < \infty$ ist, wobei T^* die Hilbertraumadjungierte ist. Die Menge der Hilbert-Schmidt-Operatoren auf H wird mit $HS(H)$ bezeichnet. Auf $HS(H)$ ist $(S, T) \mapsto \text{spur}(T^*S)$ ein Skalarprodukt mit dem $HS(H)$ ein Hilbertraum über $\mathbb{K}$ ist. Die von diesem Skalarprodukt induzierte Norm heißt *Hilbert-Schmidt-Norm* und wird durch $\|\cdot\|_{HS}$ symbolisiert. Es gilt $(HS(H), \|\cdot\|_{HS}) = (C_2(H), \|\cdot\|_2)$.

Die Elemente der Schattenklasse $C_1(H)$ heißen *Spurklasse-Operatoren* oder auch *nukleare Operatoren* auf H . Anstelle von $(C_1(H), \|\cdot\|_1)$ schreibt man $(N(H), \|\cdot\|)$ oder einfach nur $N(H)$; dabei ist zu beachten, dass im Allgemeinen weder $C_1(H)$ noch $C_2(H)$ in $\mathcal{L}(H)$, versehen mit der üblichen Operatornorm, abgeschlossen ist. $(C_1(H), \|\cdot\|_1)$ und $(C_2(H), \|\cdot\|_2)$ sind Banach-*-Algebren, die mit einer C^* -Norm ausgestattet sind (PALMER [232](2001), Seite 979).

Als Mengen gilt: $N(H) = HS(H)^2$. $N(H)$ ist ein *-Ideal in $\mathcal{L}(H)$. Wegen der eben erwähnten im Allgemeinen nicht vorliegenden Abgeschlossenheit in $\mathcal{L}(H)$ ist $N(H)$ aber im Allgemeinen keine C^* -Algebra. Es gilt $\|\cdot\|_1 = \text{spur}(|\cdot|)$. Des Weiteren kann man nach Definition von $C_1(H)$ in der Definition der Spur als Definitionsbereich auch ganz $C_1(H)$ einsetzen, wobei man dann natürlich als Wertebereich etwa $\mathbb{K}$ zu wählen hat. Dann ist *spur* ein lineares Funktional auf $C_1(H)$ und wird auch *Spurfunktional* genannt. REED und SIMON [245](1980) bemerken dazu, dass diese Vorgehensweise analog ist zu der Konstruktion des Lebesgue-Integrals, wo man $\int f d\mu$ zuerst für $f \geq 0$ definiert und dabei Werte in $[0, \infty]$ hat. Dann definiert man $L_1(\mu)$ als die Menge der f mit $\int |f| d\mu < \infty$. $L_1(\mu)$ ist ein Vektorraum und $f \mapsto \int f d\mu$ ist ein lineares Funktional.

Wegen der C^* -Bedingung gilt $\|\cdot\|_1 = \|\cdot\|$. Für $S \in N(H)$ und $T \in \mathcal{L}(H)$ gilt $\text{spur}(ST) = \text{spur}(TS)$.

Für Banachräume X und Y nennt man allgemeiner einen Operator $T \in \mathcal{L}(X, Y)$ *nuklear*, falls es Folgen $(x_n^*)_{n \in \mathbb{N}} \subseteq X^*$ und $(y_n)_{n \in \mathbb{N}} \subseteq Y$ mit $\sum_{n=1}^{\infty} \|x_n^*\| \cdot \|y_n\| < \infty$ gibt, so dass $Tx = \sum_{n=1}^{\infty} x_n^*(x)y_n$ für alle $x \in X$ gilt; siehe WERNER [291](2002), Seite 256, Definition VI.5.1.

Die beiden Abbildungen

$$\begin{aligned} f: \mathcal{L}(H) &\rightarrow N(H)^*, \quad S \mapsto (T \mapsto \text{spur}(ST)) && \text{und} \\ g: N(H) &\rightarrow K(H)^*, \quad R \mapsto (S \mapsto \text{spur}(RS)) \end{aligned}$$

sind isometrische Isomorphismen (WERNER [291](2002), Seite 273, Satz VI.6.4). Dies ist das nicht-kommutative Analogon zu den klassischen Resultaten $(c_0)^* \cong \ell_1$ und $(\ell_1)^* \cong \ell_\infty$ (SUNDER [271](1987)).

3.6.3 Schwache Operator-Topologie (PALMER [232](2001), Seite 886). Seien X und Y zwei topologische Vektorräume über $\mathbb{K}$. Die von der Produkttopologie von $(Y, \sigma(Y, Y^*))^X$ auf den Unterraum $\mathcal{L}(X, Y)$ induzierte Topologie heißt die **Topologie der punkweisen schwachen Konvergenz** oder die *schwache Operator-Topologie* auf $\mathcal{L}(X, Y)$ und wird mit *wo* bezeichnet.

Sei $\langle \cdot, \cdot \rangle$ die kanonische Bilinearform von Y . Setze $S := \{ \langle \cdot, x, y^* \rangle \in \mathcal{L}(X, Y)^\# : x \in X, y^* \in Y^* \}$ und $\mathcal{L}(X, Y)_\sim := \text{span}(S)$. Dann ist $\sigma(\mathcal{L}(X, Y), S) = \sigma(\mathcal{L}(X, Y), \mathcal{L}(X, Y)_\sim)$ (siehe Bemerkung 2.4.29(e)) die schwache Operator-Topologie von $\mathcal{L}(X, Y)$. Diese Topologie wird demgemäß von den mit $(x, y^*) \in X \times Y^*$ indizierten Halbnormen

$$T \mapsto |\langle Tx, y^* \rangle|$$

erzeugt.

Nebenbei sei bemerkt, dass für alle $T \in L(X, Y)$ die Gleichung $\|T\| = \sup\{|\langle Tx, y^* \rangle| : x \in B_X, y^* \in B_{Y^*}\}$ gilt.

Der Abschluss einer Teilmenge $A \subseteq \mathcal{L}(X, Y)$ in der schwachen Operator-Topologie wird mit $\text{cl}(wo; A)$ bezeichnet.

Sei H ein Hilbertraum über $\mathbb{K}$. Äquivalent zu der kanonischen Bilinearform von Y kann man im Fall von $X = Y = H$ für $\langle \cdot, \cdot \rangle$ das Skalarprodukt von H nehmen. Im Fall von $\mathbb{K} = \mathbb{C}$ genügt es, nur die Halbnormen mit $x = y$ zu verwenden; siehe dazu auch Bemerkung 2.2.42.

Für jedes $T \in \mathcal{L}(H)$ ist die Menge aller

$$U_F(T) := \{S \in \mathcal{L}(H) : |\langle (T - S)x, y \rangle| < 1, (x, y) \in F\}, \quad F \subseteq H \times H \text{ endlich,}$$

eine Umgebungsbasis von T in der schwachen Operator-Topologie.

Die schwache Operator-Topologie auf $\mathcal{L}(H)$ ist die Initialtopologie der Funktionenfamilie

$$\{\mathcal{L}(H) \rightarrow (H, \sigma(H, H^*)), T \mapsto Tx : x \in H\}$$

(DIXMIER [76](1981), Seite 35).

Da $\mathcal{L}(H)$ Hausdorff'sch ist, gilt nach den Bemerkungen 2.4.29(a) und (c) die Gleichung $\mathcal{L}(H)_\sim = (\mathcal{L}(H), wo)^*$, soll heißen, dass $\mathcal{L}(H)_\sim$ genau aus den *wo*-stetigen linearen Funktionalen auf $\mathcal{L}(H)$ besteht. Außerdem gilt per kanonischer Dualität: $(\mathcal{L}(H)_\sim, \|\cdot\|_{\mathcal{L}(H)^*})^* \cong \mathcal{L}(H)$ (PALMER [232](2001), Seite 886, Lemma 9.3.2(b)).

MURPHY [217](2004), Seite 136: *wo*-konvergente Folgen in $\mathcal{L}(H)$ sind beschränkt; dies gilt im Allgemeinen nicht für *wo*-konvergente Netze in $\mathcal{L}(H)$.

3.6.4 Schwach*-Operator-Topologie (HOLMES [137](1975), Seite 215). Seien X und Y normierte Vektorräume über $\mathbb{K}$. Die von der Funktionenfamilie

$$\{\mathcal{L}(X, Y^*) \rightarrow \mathbb{K}, T \mapsto \langle y, Tx \rangle : x \in X, y \in Y\}$$

auf $\mathcal{L}(X, Y^*)$ induzierte Topologie heißt die *schwach*-Operatortopologie*; sie wird mit *w*o* bezeichnet. Nebenbei sei bemerkt, dass für alle $T \in L(X, Y^*)$ die Gleichung $\|T\| = \sup\{|\langle y, Tx \rangle| : x \in B_X, y \in B_Y\}$ gilt.

Setze $S := \{ \langle y, \cdot, x \rangle \in \mathcal{L}(X, Y^*)^\# : x \in X, y \in Y \}$. Für vollständige X, Y gilt dann: $\mathcal{L}(X, Y^*) \cong (\text{span}(S), \|\cdot\|_{\mathcal{L}(X, Y^*)^*})^*$. Ist H ein Hilbertraum über $\mathbb{K}$ und setzt man $X := H$ und $Y := H^*$, so stimmt diese Topologie offensichtlich überein mit der schwachen Operatortopologie auf $\mathcal{L}(H)$.

3.6.5 Ultraschwache Operator-Topologie (TAKESAKI [274](2001), Seite 36; DIXMIER [76](1981), Seite 36). Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{K}$. Sei $T \in \mathcal{L}(H)$ und $\ell \in$

$K(H)^*$. Es gilt $\mathcal{L}(H) \cong K(H)^{**}$. Genauer ist mit den in Bemerkung 3.6.2 definierten Abbildungen f und g zu jedem $T \in \mathcal{L}(H)$ ein Funktional $k_T := (g^{*-1} \circ f)(T)$ aus $K(H)^{**}$ assoziiert und man kann $\langle T, \ell \rangle$ definieren als $\langle \ell, k_T \rangle$. Somit ist $K(H)^*$ als ein Untervektorraum von $\mathcal{L}(H)^*$ auffassbar und dieser Untervektorraum wird mit $\mathcal{L}(H)_*$ bezeichnet. Da $\mathcal{L}(H)_* \cong N(H)$, ist $\mathcal{L}(H)_*$ ein Prädual von $\mathcal{L}(H)$. Dementsprechend kann man auf $\mathcal{L}(H)$ die durch das Dualsystem $(\mathcal{L}(H), \mathcal{L}(H)_*)$ induzierte lokal konvexe Topologie $\sigma(\mathcal{L}(H), \mathcal{L}(H)_*)$ betrachten. Sie wird die σ -schwache Operator-Topologie oder ultraschwache Operator-Topologie auf $\mathcal{L}(H)$ genannt und mit *uwo* bezeichnet.

Da $\mathcal{L}(H) \cong N(H)^*$ ist, kann man gemäß der Bemerkung 2.4.31 das Dualsystem $(\mathcal{L}(H), \mathcal{L}(H)_*)$ auffassen als das Dualsystem $(N(H)^*, N(H))$. Für $\mathcal{L}(H)$, aufgefasst als $N(H)^*$, ist also die ultraschwache Operator-Topologie gleich der **schwach*-Topologie** $\sigma(N(H)^*, N(H))$. Die traditionelle Bezeichnung *ultraschwache Operator-Topologie* stammt noch aus der Zeit, wo die Theorie der Banachräume und der schwach*-Topologie noch nicht voll etabliert war (CONWAY [53](1999)).

Wegen $\langle T, \ell \rangle = \langle \ell, k_T \rangle = \langle \ell, (g^{*-1} \circ f)(T) \rangle = \langle \ell, (g^{-1*} \circ f)(T) \rangle = \langle g^{-1}(\ell), f(T) \rangle = \text{spur}(Tg^{-1}(\ell))$ induzieren die mit $S \in N(H)$ indizierten Funktionale

$$\ell_S: T \mapsto \text{spur}(ST)$$

aus $\mathcal{L}(H)^*$ die ultraschwache Operator-Topologie auf $\mathcal{L}(H)$.

TAKESAKI [274](2001), Seite 67, zeigt, dass genau die Elemente von $\mathcal{L}(H)_*$ als Funktionale der Form $T \mapsto \sum_{n \in \mathbb{N}} \langle Tx_n, y_n \rangle$ geschrieben werden können, wobei die $(x_n)_{n \in \mathbb{N}}$ und $(y_n)_{n \in \mathbb{N}}$ Folgen des Hilbertraumes H sind mit $\sum_{n \in \mathbb{N}} \|x_n\|^2 < \infty$ und $\sum_{n \in \mathbb{N}} \|y_n\|^2 < \infty$. Somit wird die ultraschwache Operator-Topologie von $\mathcal{L}(H)$ von den Halbnormen

$$T \mapsto \left| \sum_{n \in \mathbb{N}} \langle Tx_n, y_n \rangle \right| \quad \text{mit} \quad \sum_{n \in \mathbb{N}} \|x_n\|^2 < \infty \quad \text{und} \quad \sum_{n \in \mathbb{N}} \|y_n\|^2 < \infty$$

erzeugt.

Im Fall von $\mathbb{K} = \mathbb{C}$ genügt es, nur die Halbnormen mit $(x_n)_{n \in \mathbb{N}} = (y_n)_{n \in \mathbb{N}}$ zu verwenden; siehe dazu auch Bemerkung 2.2.42.

Für jedes $T \in \mathcal{L}(H)$ ist die Menge aller

$$\left\{ S \in \mathcal{L}(H) : \left| \sum_{\nu \in \mathbb{N}} \langle (T - S)x_{\mu\nu}, y_{\mu\nu} \rangle \right| < 1, \mu = 1, \dots, n \right\}$$

mit $(x_{\mu\nu})_{\nu \in \mathbb{N}}, (y_{\mu\nu})_{\nu \in \mathbb{N}} \subseteq H$, $\sum_{\nu \in \mathbb{N}} \|x_{\mu\nu}\|^2 + \|y_{\mu\nu}\|^2 < \infty$, $\mu = 1, \dots, n$, $n \in \mathbb{N}$, eine Umgebungsbasis von T in der ultraschwachen Operator-Topologie.

Der Abschluss einer Teilmenge $A \subseteq \mathcal{L}(H)$ in der ultraschwachen Operator-Topologie wird mit $\text{cl}(uwo; A)$ bezeichnet.

Die Hilbertraum-Involution ist ultraschwach Operator-stetig.

Mit den Bezeichnungen aus dem Abschnitt über die schwache Operator-Topologie gilt, dass $\mathcal{L}(H)_*$ der Abschluss von $\mathcal{L}(H)_\sim$ in $\mathcal{L}(H)^*$ mit seiner Norm-Topologie ist. Gemäß Satz 2.4.35 (b) stimmt also auf beschränkten Teilmengen von $\mathcal{L}(H)$ die schwache Operator-Topologie überein mit der ultraschwachen Operator-Topologie.

3.6.6 Starke Operator-Topologie. Seien X und Y zwei topologische Vektorräume über $\mathbb{K}$. Die von der Produkttopologie von Y^X auf den Unterraum $\mathcal{L}(X, Y)$ induzierte Topologie heißt die *Topologie der punktwweisen Konvergenz* oder auch die *starke Operator-Topologie* auf $\mathcal{L}(X, Y)$ und wird mit *so* bezeichnet. Man beachte, dass $\mathcal{L}(X, Y)$ in der starken Operator-Topologie im Allgemeinen nicht Folgen-abgeschlossen

ist. Nach dem Satz von Banach-Steinhaus ist $\mathcal{L}(X, Y)$ in der starken Operator-Topologie Folgen-abgeschlossen, wenn X und Y Banachräume über $\mathbb{K}$ sind. (MATHIEU [208](1998), Seite 111.)

Sind X und Y Banachräume über $\mathbb{K}$, dann ist für jedes $T \in \mathcal{L}(X, Y)$ die Menge aller

$$U_F(T) := \{S \in \mathcal{L}(X, Y) : \|(T - S)x\| < 1, x \in F\}, \quad F \subseteq X \text{ endlich,}$$

eine Umgebungsbasis von T in der starken Operator-Topologie.

Die starke Operator-Topologie wird von den folgenden mit $x \in X$ indizierten Halbnormen erzeugt:

$$p_x : T \mapsto \|Tx\|.$$

WERNER [291](2002), Seite 377, Beispiel VIII.2(c), zeigt: Sind X und Y normierte Räume über $\mathbb{K}$, dann gilt $(\mathcal{L}(X, Y); so)^* = \mathcal{L}(X, Y)_\sim$, das heißt, alle Elemente von $(\mathcal{L}(X, Y); so)^*$ sind von der Form

$$T \mapsto \sum_{\nu=1}^n \langle Tx_\nu, y_\nu^* \rangle$$

mit $x_\nu \in X$, $y_\nu^* \in Y^*$, $\nu = 1, \dots, n$, $n \in \mathbb{N}$.

Ist Y endlichdimensional, dann stimmt die starke Operator-Topologie mit der schwachen Operator-Topologie überein. Der Abschluss einer Teilmenge $A \subseteq \mathcal{L}(X, Y)$ in der starken Operator-Topologie wird mit $\text{cl}(so; A)$ bezeichnet.

Sei H ein Hilbertraum über $\mathbb{K}$. Für den Fall $X = Y = H$ gilt: Außer wenn H endlichdimensional ist, ist die Hilbertraum-Involution nicht stark Operator-stetig.

MATHIEU [208](1998), Seite 334: Da $\mathcal{L}(H)$ *so*-Folgen-abgeschlossen ist, sind *so*-konvergente Folgen nach dem Satz von Banach-Steinhaus beschränkt. Für Netze gilt das Analoge im Allgemeinen nicht: Zum Beispiel ist das Netz $T_U := n(1 - p_U)$, $U \in I$, wobei $I := \{V \subseteq \mathcal{L}(H) : V \text{ endlichdimensionaler Unterraum von } H\}$ kanonisch durch Inklusion gerichtet ist, p_U die Projektion auf U ist und $n = \dim(U)$ gilt, zwar gegen 0 *so*-konvergent, aber unbeschränkt.

3.6.7 Ultrastarke Operator-Topologie (DIXMIER [76](1981), Seite 35; TAKESAKI [274](2001), Seite 68). Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{K}$. Die von der Menge der mit $\ell \in \mathcal{L}(H)_*$, $\ell \geq 0$, indizierten Halbnormen

$$p_\ell : T \mapsto \ell(T^*T)^{1/2}$$

erzeugte lokal konvexe Topologie auf $\mathcal{L}(H)$ heißt die *σ -starke Operator-Topologie* oder *ultrastarke Operator-Topologie* oder *stärkste Operator-Topologie* und wird mit *uso* bezeichnet. Gemäß der Definition von $\mathcal{L}(H)_*$ sind im Fall von $\mathbb{K} = \mathbb{C}$ die Halbnormen p_ℓ , $\ell \in \mathcal{L}(H)_*$, $\ell \geq 0$, genau die Abbildungen

$$T \mapsto \left(\sum_{n \in \mathbb{N}} \|Tx_n\|^2 \right)^{1/2} \quad \text{mit } (x_n)_{n \in \mathbb{N}} \subseteq H, \sum_{n \in \mathbb{N}} \|x_n\|^2 < \infty.$$

Der Abschluss einer Teilmenge $A \subseteq \mathcal{L}(H)$ in der ultrastarken Operator-Topologie wird mit $\text{cl}(uso; A)$ bezeichnet.

Die Hilbertraum-Involution ist im Allgemeinen nicht ultrastark Operator-stetig.

3.6.8 Stark* Operator-Topologie. Sei H ein Hilbertraum über $\mathbb{K}$. Die *stark** Operator-Topologie *s*o* auf $\mathcal{L}(H)$ wird von der folgenden Menge von Halbnormen erzeugt:

$$\{T \mapsto \|Tx\|, T \mapsto \|T^*x\| : x \in H\},$$

wobei T^* die Hilbertraumadjungierte von T ist. BRATTELI und ROBINSON [37](1987), Seite 69, verwenden als Halbnormen die Abbildungen $T \mapsto \|Tx\| + \|T^*x\|$, $x \in H$.

Der Abschluss einer Teilmenge $A \subseteq \mathcal{L}(H)$ in der stark* Operator-Topologie wird mit $\text{cl}(s^*o; A)$ bezeichnet.

3.6.9 Ultrastark* Operator-Topologie. Sei H ein Hilbertraum über $\mathbb{K}$. Die von der Menge der mit $\ell \in \mathcal{L}(H)_*$, $\ell \geq 0$, indizierten Halbnormen

$$p_\ell: T \mapsto (\ell(T^*T) + \ell(TT^*))^{1/2}$$

erzeugte lokal konvexe Topologie auf $\mathcal{L}(H)$ heißt die σ -stark* Operator-Topologie oder *ultrastark* Operator-Topologie* und wird mit us^*o bezeichnet. Gemäß der Definition von $\mathcal{L}(H)_*$ sind im Fall von $\mathbb{K} = \mathbb{C}$ die Halbnormen p_ℓ , $\ell \in \mathcal{L}(H)_*$, $\ell \geq 0$, genau die Abbildungen $T \mapsto (\sum_{n \in \mathbb{N}} \|Tx_n\|^2 + \sum_{n \in \mathbb{N}} \|T^*x_n\|^2)^{1/2}$ mit $(x_n)_{n \in \mathbb{N}} \subseteq H$, $\sum_{n \in \mathbb{N}} \|x_n\|^2 < \infty$.

Diese Topologie ist die schwächste Topologie, die stärker als die ultrastarke Operator-Topologie ist und in der die Hilbertraum-Involution stetig ist.

3.6.10 Bemerkung (SUNDER [271], Seite 11). Es gilt folgendes Schema von Mengeninklusionen von (Operator-)Topologien:

$$\begin{array}{ccccccc} \text{Norm} & \supseteq & \text{Mackey} & \supseteq & \sigma\text{-stark}^* & \supseteq & \sigma\text{-stark} & \supseteq & \sigma\text{-schwach} \\ & & & & \text{IU} & & \text{IU} & & \text{IU} \\ & & & & \text{stark}^* & \supseteq & \text{stark} & \supseteq & \text{schwach} \end{array}$$

Mit den Topologieabkürzungen liest sich dieses Schema als:

$$\begin{array}{ccccccc} \|\cdot\| & \supseteq & \tau & \supseteq & us^*o & \supseteq & us o & \supseteq & uwo \equiv w^* \\ & & & & \text{IU} & & \text{IU} & & \text{IU} \\ & & & & s^*o & \supseteq & so & \supseteq & wo \end{array}$$

Dabei bezieht sich die Mackey-Topologie auf das Dualsystem $(\mathcal{L}(H), N(H))$, mit anderen Worten ist als Mackey-Topologie $\tau(\mathcal{L}(H), \mathcal{L}(H)_*)$ gemeint. Wenn H unendlich-dimensional ist, dann ist jede dieser Inklusionen echt. Da die schwache Topologie $\sigma(\mathcal{L}(H), \mathcal{L}(H)^*)$ stärker ist als die ultraschwache Operator-Topologie, sprich schwach*-Topologie $\sigma(\mathcal{L}(H), \mathcal{L}(H)_*)$, ist die schwache Topologie stärker als die schwache Operator-Topologie.

3.6.11. Auf der abgeschlossenen Einheitskugel $B_{\mathcal{L}(H)}$, H ein Hilbertraum über $\mathbb{K}$, gilt: $wo|_{B_{\mathcal{L}(H)}} = uwo|_{B_{\mathcal{L}(H)}}$, $so|_{B_{\mathcal{L}(H)}} = us o|_{B_{\mathcal{L}(H)}}$ und $s^*o|_{B_{\mathcal{L}(H)}} = us^*o|_{B_{\mathcal{L}(H)}}$; siehe dazu 3.6.5 und TAKESAKI [274](2001), Seite 69, Lemma 2.5 und LI [198](2003), Seite 62, Theorem 4.2.2.

3.6.12 Satz. Sei H ein Hilbertraum über $\mathbb{C}$. Dann gilt:

$$(\mathcal{L}(H), wo)^* = (\mathcal{L}(H), so)^* = (\mathcal{L}(H), s^*o)^*$$

$$\text{und } (\mathcal{L}(H), uwo)^* = (\mathcal{L}(H), us o)^*.$$

Beweis. PALMER [232](2001), Seite 886; DIXMIER [76](1981), Seite 39; □

3.6.13 Satz (LI [198] (2003), Seite 63). Sei H ein Hilbertraum über $\mathbb{R}$ und $\ell \in \mathcal{L}(H)^\#$. Dann sind die folgenden Aussagen äquivalent:

(a) ℓ ist uwo -stetig.

- (b) ℓ ist *uso-stetig*.
- (c) ℓ ist *us^{*}o-stetig*.
- (d) $\ell \upharpoonright \lambda \cdot B_{\mathcal{L}(H)}$ ist *wo-stetig* für alle $\lambda \in \mathbb{R}^+$.
- (e) $\ell \upharpoonright \lambda \cdot B_{\mathcal{L}(H)}$ ist *so-stetig* für alle $\lambda \in \mathbb{R}^+$.

3.6.14 Korollar. Sei H ein Hilbertraum über $\mathbb{K}$ und K eine konvexe Teilmenge von $\mathcal{L}(H)$. Dann gilt $cl(wo; K) = cl(so; K)$ und $cl(uwo; K) = cl(uso; K)$. Im Fall von $\mathbb{K} = \mathbb{C}$ gilt $cl(wo; K) = cl(so; K) = cl(s^*o; K)$ und $cl(uwo; K) = cl(uso; K) = cl(us^*o; K)$.

Beweis. Siehe Korollar 2.6.11. Quelle: CONWAY [53](1999), Seite 38; LI [198] (2003), Seite 62; BRATTELI und ROBINSON [37](1987), Seite 71. $\square$

3.6.15 Bemerkung. Sei H ein Hilbertraum über $\mathbb{K}$. Mit jeder der hier beschriebenen Operator-Topologie ist $\mathcal{L}(H)$ eine topologische Algebra.

Beweis. Siehe NAIMARK [221](1972), Seite 449; PALMER [232](2001), Seite 885. $\square$

3.6.16 Bemerkung. Sei H ein unendlichdimensionaler Hilbertraum über $\mathbb{K}$. Die Involution von $\mathcal{L}(H)$, also die Bildung der Hilbertraumadjungierten, ist *wo-*, *uwo-* und nach Konstruktion *s^{*}o-* und *us^{*}o-*stetig, aber weder *so-* noch *uso-*stetig. Somit ist $\mathcal{L}(H)_{(*)}$ in den *wo-*, *uwo-*, *s^{*}o-* und *us^{*}o-*Topologien abgeschlossen.

Beweis. Siehe DIXMIER [76](1969); HOLDGRÜN [135](1983), Seite 407; NAIMARK [221](1972), Seite 449; BRATTELI und ROBINSON [37](1987), Seite 69. $\square$

3.6.17 Satz (LI [198] (2003), Seite 69). Sei A eine $*$ -Unteralgebra von $\mathcal{L}(H)$, H ein Hilbertraum über $\mathbb{R}$. Dann gilt: $cl(wo; A) = cl(so; A) = cl(s^*o; A) = cl(uwo; A) = cl(uso; A) = cl(us^*o; A)$.

3.6.18 Satz. Sei A eine $*$ -Unteralgebra von $\mathcal{L}(H)$, H ein Hilbertraum über $\mathbb{C}$. Dann sind die folgenden Aussagen äquivalent:

- (a) A ist *wo-abgeschlossen*. (b) B_A ist *wo-abgeschlossen*.
- (c) A ist *so-abgeschlossen*. (d) B_A ist *so-abgeschlossen*.
- (e) A ist *uwo-abgeschlossen*. (f) B_A ist *uwo-abgeschlossen*.
- (g) A ist *uso-abgeschlossen*. (h) B_A ist *uso-abgeschlossen*.
- (i) A ist *s^{*}o-abgeschlossen*. (j) A ist *us^{*}o-abgeschlossen*.

Beweis. (a) bis (e): DIXMIER [76](1981), Seite 45.

(i) und (j) folgen aus (a) und (e) per Korollar 3.6.14 und Bemerkung 1.2.5. $\square$

3.7 W^* -Algebren

3.7.1 Definition (PALMER [232](2001), Seite 799; LI [198](2003), Seiten 63 und 123; SAKAI [253](1971), Seite 103).

- (a) Eine *von Neumann-Algebra* über $\mathbb{K}$ ist eine $*$ -Unteralgebra A von $\mathcal{L}(H)$ für einen Hilbertraum H über $\mathbb{K}$, für die $A = A''$ gilt. Genauer sagt man auch, dass A eine von Neumann-Algebra auf H über $\mathbb{K}$ sei.
- (b₁) Eine *W^* -Algebra* über $\mathbb{K}$ auf einem Hilbertraum H ist eine in der schwachen Operator-Topologie abgeschlossene $*$ -Unteralgebra von $\mathcal{L}(H)$, H ein Hilbertraum über $\mathbb{K}$.

- (b_2) Eine W^* -Algebra über $\mathbb{R}$ auf einem Hilbertraum H über $\mathbb{C}$ ist eine in der schwachen Operator-Topologie abgeschlossene $*$ -Unteralgebra von $\mathcal{L}(H)_{\mathbb{R}}$, H ein Hilbertraum über $\mathbb{C}$.
- (c_1) Eine W^* -Algebra über $\mathbb{R}$ ist eine Banach- $*$ -Algebra über $\mathbb{R}$, die isometrisch $*$ -isomorph zu einer von Neumann-Algebra über $\mathbb{R}$ ist.
- (c_2) Eine W^* -Algebra über $\mathbb{C}$ ist eine $*$ -Algebra über $\mathbb{C}$, die $*$ -isomorph zu einer von Neumann-Algebra über $\mathbb{C}$ ist.

3.7.2 Bemerkung. Eine äquivalente Formulierung der Definition einer W^* -Algebra über $\mathbb{C}$ lautet: Eine W^* -Algebra über $\mathbb{C}$ ist eine $*$ -Algebra über $\mathbb{C}$, die eine treue $*$ -Darstellung als eine von Neumann-Algebra über $\mathbb{C}$ besitzt.

Ist H ein Hilbertraum über $\mathbb{C}$, so nennt von Neumann eine Teilmenge M von $\mathcal{L}(H)$ einen *Ring*, falls M eine W^* -Algebra über $\mathbb{C}$ auf dem Hilbertraum H ist (VON NEUMANN [224](1929), II.1).

3.7.3 Bemerkung (USMANOV [281](2000), Seite 828). Eine andere Definition einer von Neumann-Algebra über $\mathbb{R}$ ist in dem Umfeld von STØRMER (Oslo, Norwegen), AYUPOV (Taschkent, Usbekistan) und anderen üblich; sie lautet: Sei H ein Hilbertraum über $\mathbb{C}$. Sei A eine $*$ -Unteralgebra über $\mathbb{R}$ von $\mathcal{L}(H)_{\mathbb{R}}$. Ist A als Teilmenge von $\mathcal{L}(H)$ in der schwachen Operator-Topologie abgeschlossen, $\text{Id}_H \in A$ und ist A “ziemlich reell“ in dem Sinne, dass $A \cap iA = \{0\}$ gilt, so heißt A eine *von Neumann-Algebra über $\mathbb{R}$ nach STØRMER-AYUPOV*.

Sei A eine von Neumann-Algebra über $\mathbb{R}$ nach STØRMER-AYUPOV und H der dazugehörige Hilbertraum über $\mathbb{C}$. Sei $F \subseteq \mathcal{L}(H)$ die kleinste von Neumann-Algebra über $\mathbb{C}$, die A als Menge umfasst. Dann gilt $F = A_{(\mathbb{C})}$. Siehe auch Satz 3.7.12.

3.7.4 Bemerkung (PALMER [232](2001), Seite 889). Sei H ein Hilbertraum über $\mathbb{K}$. Nach Satz 2.7.4 ist jede von Neumann-Algebra $A \subseteq \mathcal{L}(H)$ über $\mathbb{K}$ eine Norm-abgeschlossene $*$ -Unteralgebra A von $\mathcal{L}(H)$ und somit eine C^* -Algebra über $\mathbb{K}$ von Operatoren. Folglich ist jede von Neumann-Algebra über $\mathbb{K}$ und jede W^* -Algebra über $\mathbb{K}$ eine C^* -Algebra über $\mathbb{K}$. Jede W^* -Algebra über $\mathbb{K}$ auf H ist eine C^* -Algebra über $\mathbb{K}$.

(TOPPING [278](1971)). Sei H ein Hilbertraum über $\mathbb{K}$ und A eine Teilmenge von $\mathcal{L}(H)$. Dann ist nach Satz 2.7.4 und Bemerkung 3.6.15 die Menge A' abgeschlossen in jeder der in Abschnitt 3.6 definierten Operator-Topologie von $\mathcal{L}(H)$.

Insbesondere ist nach Satz 3.2.39 für jede $*$ -Teilmenge $A \subseteq \mathcal{L}(H)$ die Kommutante A' eine W^* -Algebra über $\mathbb{K}$ auf H und, wenn man auch den Satz 2.3.26(e) berücksichtigt, eine von Neumann-Algebra über $\mathbb{K}$; ist des Weiteren E eine von Neumann-Algebra auf H über $\mathbb{K}$ mit $A \subseteq E$, dann ist A'' wegen $A \subseteq A'' \subseteq E'' = E$ die kleinste von Neumann-Algebra über $\mathbb{K}$, die A umfasst.

Ist S eine Teilmenge von $\mathcal{L}(H)$ und bezeichnet S^* die Menge aller Hilbert-raumadjungierten T^* von T mit $T \in S$, so sagt man daher, dass $(S \cup S^*)''$ die von S erzeugte von Neumann-Algebra über $\mathbb{K}$ sei.

3.7.5 Von Neumann's Doppel-Kommutanten-Theorem. Sei H ein Hilbertraum über $\mathbb{K}$. Für jede $*$ -Unteralgebra A von $\mathcal{L}(H)$, die im Fall von $\mathbb{K} = \mathbb{R}$ Id_H enthält und im Fall $\mathbb{K} = \mathbb{C}$ wesentlich ist, gilt: A wo-abgeschlossen $\Leftrightarrow A = A''$.

Beweis. Für einen direkten Beweis siehe: SUNDER [271] (2001), Seite 888; LI [198](2003), Seite 67. $\square$

3.7.6 Von Neumann's Dichte-Theorem. Sei H ein Hilbertraum über $\mathbb{K}$. Für jede $*$ -Unteralgebra A von $\mathcal{L}(H)$, die im Fall von $\mathbb{K} = \mathbb{R}$ Id_H enthält und im Fall $\mathbb{K} = \mathbb{C}$ wesentlich ist, gilt $\text{cl}(wo; A) = A''$.

Beweis. Für den Fall $\mathbb{K} = \mathbb{C}$ siehe zum Beispiel: PALMER [232](2001), Seite 888; DIXMIER [76](1981), Seite 45; MURPHY [217](2004), Seite 115, Lemma 4.1.4. Für den Fall $\mathbb{K} = \mathbb{R}$ siehe Satz 2.7.6. $\square$

3.7.7. Der Satz 2.7.6 sagt, dass das von Neumann'sche Doppel-Kommutanten-Theorem und das von Neumann'sche Dichte-Theorem äquivalente Aussagen sind. Das ist speziell im Fall von $\mathbb{K} = \mathbb{R}$ von Interesse, da hier nur ein Beweis für das von Neumann'sche Doppel-Kommutanten-Theorem vorlag.

3.7.8. Für eine direkte Folgerung des von Neumann'schen Dichte-Theorems 3.7.6 betrachte man einen Hilbertraum H über $\mathbb{C}$. Sei A eine wesentliche $*$ -Unteralgebra von $\mathcal{L}(H)$. Dann ist nach dem Neumann'schen Dichte-Theorem A'' die kleinste wo -abgeschlossene $*$ -Unteralgebra von $\mathcal{L}(H)$, die A umfasst. Berücksichtigt man noch, dass nach Satz 2.3.26(i) $S' = (S \cup \{\text{Id}_H\})'$ gilt, so hat man folglich: Ist S eine Teilmenge von $\mathcal{L}(H)$, H ein Hilbertraum über $\mathbb{K}$, und bezeichnet S^* die Menge aller Hilbertraumadjungierten T^* von T mit $T \in S$, dann ist $(S \cup S^*)''$ die kleinste W^* -Algebra über $\mathbb{K}$ auf H , die S umfasst; man kann also von $(S \cup S^*)''$ als die von S erzeugte W^* -Algebra über $\mathbb{K}$ auf H sprechen.

3.7.9 Kaplansky's Dichte-Theorem (LI [198](2003), Seite 68, Theorem 4.4.1; LI [197](1992), Seite 35, Theorem 1.6.1). *Sei H ein Hilbertraum über $\mathbb{K}$, U und V zwei $*$ -Unteralgebren von $\mathcal{L}(H)$ mit $U \subseteq V$. Ist U wo -dicht in V , dann ist $B_U \tau(\mathcal{L}(H), \mathcal{L}(H)_*)$ -dicht in B_V .*

3.7.10. Ist also speziell W eine W^* -Algebra über $\mathbb{K}$ auf einem Hilbertraum H und $A \subseteq W$ eine in W wo -dichte $*$ -Unteralgebra von $\mathcal{L}(H)$ (also $\text{cl}(wo; A) = W$), so ist B_A in der Mackey-Topologie $\tau(\mathcal{L}(H), \mathcal{L}(H)_*)$ dicht in B_W , das heißt, für jedes $x \in B_W$ hat jede τ -Umgebung von x einen nicht leeren Durchschnitt mit B_A . Wegen 1.2.17(ℓ) und Bemerkung 3.6.10 ist also auch für jede Topologie $\tau \in \{us^*o, s^*o, uso, so, w^*, wo\}$ B_A τ -dicht in B_W . Für jedes $\tau \in \{us^*o, s^*o, uso, so, w^*, wo\}$ gilt nach 3.6.11 also

$$\text{cl}(\tau; A) \cap B_{\mathcal{L}(H)} = \text{cl}(\tau; B_A).$$

Im Fall von $\mathbb{K} = \mathbb{C}$ gilt

$$\text{cl}(so; A) \cap B_{\mathcal{L}(H)} \cap \text{Pos}(*\sigma; \mathcal{L}(H)) \subseteq \text{cl}(so; B_A \cap \text{Pos}(*\sigma; A)).$$

und

$$\text{cl}(so; A) \cap B_{\mathcal{L}(H)} \cap \mathcal{L}(H)_{(*)} \subseteq \text{cl}(so; B_A \cap A_{(*)}).$$

Folglich gilt:

3.7.11 (PALMER [232](2001), Seite 890). Sei H ein Hilbertraum über $\mathbb{C}$. Für jede wesentliche $*$ -Unteralgebra A von $\mathcal{L}(H)$ gilt: $B_{A''} = \text{cl}(\tau; B_A)$ mit $\tau \in \{wo, so, s^*o\}$. Siehe hierzu auch 3.7.44.

3.7.12 Satz (LI [198](2003), Seite 63). *Sei A eine $*$ -Unteralgebra A von $\mathcal{L}(H)$ für einen Hilbertraum H über $\mathbb{R}$. Dann ist A genau dann eine von Neumann-Algebra auf H über $\mathbb{R}$, wenn $A_{(\mathbb{C})}$ eine von Neumann-Algebra auf $H_{(\mathbb{C})}$ über $\mathbb{C}$ ist.*

3.7.13 Satz (PALMER [232](2001), Seite 889). *Sei H ein Hilbertraum über $\mathbb{K}$ und A eine Teilmenge von $\mathcal{L}(H)$. Dann sind die folgenden Aussagen äquivalent:*

- (a) A ist eine von Neumann-Algebra über $\mathbb{K}$.
- (b) Es gibt eine $*$ -Teilmenge $S \subseteq \mathcal{L}(H)$ mit $A = S'$.

(c) A ist eine W^* -Algebra über $\mathbb{K}$ auf H mit $\text{Id}_H \in A$.

(d) A ist eine $*$ -Algebra mit $\text{Id}_H \in A$ und wo -abgeschlossenem B_A .

Im Fall von $\mathbb{K} = \mathbb{C}$ ist zu (a) bis (d) zusätzlich äquivalent:

(d) A ist eine wesentliche W^* -Algebra über $\mathbb{C}$ auf H .

Beweis. (a) $\Rightarrow$ (b): Klar. (b) $\Rightarrow$ (c): Sätze 3.2.39, 2.3.26(h) und 2.7.4. (c) $\Rightarrow$ (a): Das ist der von Neumann'sche Doppel-Kommutanten-Satz 3.7.5. (c) $\Rightarrow$ (e): Jede $*$ -Unteralgebra von H die Id_H enthält ist wesentlich. (e) $\Rightarrow$ (a): Das ist das von Neumann'sche Dichte-Theorem 3.7.6. Bezüglich (a) $\Leftrightarrow$ (d) siehe LI [197](1992), Seite 36 und LI [198](2003), Seite 68. $\square$

3.7.14 Satz. Sei A eine W^* -Algebra über $\mathbb{K}$ auf einem Hilbertraum H . Dann hat A eine Eins.

Beweis. Für den Fall $\mathbb{K} = \mathbb{C}$ siehe MURPHY [217](2001), Seite 118, per approximativer Eins. Für den Fall $\mathbb{K} = \mathbb{R}$ hat man: KAPLANSKY [171](1968), Seite 130, bemerkt, dass jede W^* -Algebra über $\mathbb{R}$ auf einem Hilbertraum H ein in [171] auf Seite 3 definierter Baer-Ring ist. Jeder Baer-Ring gemäß [171] hat eine Eins. $\square$

Man beachte, dass die Eins in Satz 3.7.14 nicht gleich Id_H zu sein braucht.

3.7.15 Bemerkung. Sei A eine W^* -Algebra über $\mathbb{K}$ auf einem Hilbertraum H . Bezeichne p die Eins von A . Nach Satz 2.4.8 ist $A + \mathbb{K} \text{Id}_H$ wo -abgeschlossen in $\mathcal{L}(H)$, was man hier auch direkt zeigen kann: Dazu darf man gemäß den Sätzen 3.6.17 und 3.6.18 mit der so -Topologie argumentieren. Da im Fall $p = \text{Id}_H$ nichts zu zeigen ist, sei $p \neq \text{Id}_H$. Sei $b \in \text{cl}(so; A + \mathbb{K} \text{Id}_H)$. Sei (b_ν) ein Netz in $A + \mathbb{K} \text{Id}_H$, das gegen b so -konvergiert. Da $(\mathcal{L}(H); so)$ eine topologische Algebra ist, so -konvergiert das Netz $(b_\nu(\text{Id}_H - p))$ gegen $b(\text{Id}_H - p)$. Die b_ν lassen sich alle schreiben als $a_\nu + \lambda_\nu \text{Id}_H$, $a_\nu \in A$, $\lambda_\nu \in \mathbb{K}$. Wegen $b_\nu(\text{Id}_H - p) = (a_\nu + \lambda_\nu \text{Id}_H)(\text{Id}_H - p) = \lambda_\nu(\text{Id}_H - p)$ so -konvergiert das Netz $(\lambda_\nu(\text{Id}_H - p))$ gegen $b(\text{Id}_H - p)$. Das Netz $(\lambda_\nu(\text{Id}_H - p))$ liegt im eindimensionalen und daher so -abgeschlossenen Unterraum $\mathbb{K}(\text{Id}_H - p)$ und so -konvergiert daher gegen $\lambda(\text{Id}_H - p)$ für ein $\lambda \in \mathbb{K}$. Wegen $|\lambda_\nu - \lambda| \cdot \|(\text{Id}_H - p)x\| = \|(\lambda_\nu(\text{Id}_H - p) - \lambda(\text{Id}_H - p))x\| \rightarrow 0$ für alle $x \in X$, konvergiert das Netz (λ_ν) gegen λ . Somit so -konvergiert wegen $|\lambda_\nu - \lambda|\|x\| = \|(\lambda_\nu - \lambda)\text{Id}_H x\| \rightarrow 0$ für alle $x \in X$, das Netz $(\lambda_\nu \text{Id}_H)$ gegen $\lambda \cdot \text{Id}_H$ und damit das Netz $(a_\nu) = (b_\nu) - (\lambda_\nu \text{Id}_H)$ gegen $a := b - \lambda \text{Id}_H$. Da alle $a_\nu \in A$ und A so -abgeschlossen ist, gilt $a \in A$. Also $b = a + \lambda \text{Id}_H \in A + \mathbb{K} \text{Id}_H$ und $A + \mathbb{K} \text{Id}_H$ ist so - und damit auch wo -abgeschlossen in $\mathcal{L}(H)$.

Nach Satz 3.7.13(c) oder auch direkt nach dem von Neumann'schen Dichte-Theorem 3.7.6 gilt also $A + \mathbb{K} \text{Id}_H = (A + \mathbb{K} \text{Id}_H)''$ und $A + \mathbb{K} \text{Id}_H$ ist eine von Neumann-Algebra über $\mathbb{K}$. Nach Satz 2.3.26(i) ist $A' = (A + \mathbb{K} \text{Id}_H)'$. Somit hat man

$$A'' = A + \mathbb{K} \text{Id}_H \quad (3.16)$$

(also auch $A'' = A + \mathbb{K}(\text{Id}_H - p)$) und nach der Bemerkung 2.3.21 und den Gleichungen (2.7), (2.9) und (2.12) aus 2.1.33 ergibt sich die folgende Version des **von Neumann'schen Doppel-Kommutanten-Theorems für W^* -Algebren über $\mathbb{K}$ auf einem Hilbertraum:**

$$A = \{T \in A'' : Tp = T\} = A''p \quad (3.17)$$

$$= \{T \in A'' : pT = T\} = pA'' \quad (3.18)$$

$$= \{T \in A'' : pTp = T\} = pA''p = U_p(A''). \quad (3.19)$$

(Siehe auch 3.4.24.) Quellen dieser Bemerkung: TOPPING [278](1979), Seite 13; HOLDGRÜN [135](1983), Seite 415; HOLDGRÜN [136](1987), Seite 7; CONWAY [53](1999), Seite 243.

Als eine Folgerung des von Neumann'schen Doppel-Kommutanten-Theorems für W^* -Algebren über $\mathbb{K}$ auf einem Hilbertraum hat man:

3.7.16 Satz (MURPHY [217](2001), Seite 117). *Sei A eine W^* -Algebra über $\mathbb{K}$ auf einem Hilbertraum H und bezeichne p die Eins von A . Dann ist $p \in A'$, p ist eine Projektion, die Abbildung*

$$A \rightarrow A_p, \quad T \mapsto T_p \quad (\text{siehe 3.4.22})$$

ein isometrischer $$ -Isomorphismus und A_p ist eine von Neumann-Algebra auf pH über $\mathbb{K}$.*

Beweis. $p \text{ Eins} \Rightarrow pa = a = ap$ für alle $a \in A$, also $p \in A'$. Nach Bemerkung 3.2.12(b) ist p eine Projektion. Da p eine Eins von A ist, gilt $pTpx = Tx$ für alle $x \in H$, das heißt, $U_p(A) = A$. Daher ist einerseits nach 3.4.22 die Abbildung $T \mapsto T_p$ ein $*$ -Isomorphismus und andererseits $\|T_p\| = \sup \{\|T_px\| : x \in pH, \|x\| \leq 1\} = \sup \{\|pTpx\| : x \in H, \|x\| \leq 1\} = \|T\|$. Unmittelbar direkt aus der Gleichung 3.16 oder aus der Gleichung 3.19 liest man die Gleichung $A_p = (A'')_p$ ab, und da nach 3.4.22(f) die Gleichung $(A'')_p = (A_p)''$ gilt, folgt $A_p = (A_p)''$ und A_p ist eine von Neumann-Algebra. $\square$

3.7.17 Korollar. *Jede W^* -Algebra über $\mathbb{K}$ auf einem Hilbertraum H ist eine W^* -Algebra über $\mathbb{K}$.*

3.7.18. Ist A eine W^* -Algebra über $\mathbb{C}$ auf einem Hilbertraum H , die nicht wesentlich ist, dann kann man also einerseits gemäß Satz 3.7.16 den Hilbertraum H auf eH verkleinern (Kompression) oder andererseits gemäß Bemerkung 3.7.15 die Algebra A um das Element Id_H erweitern (Anknüpfen des Skalarbereiches), um sich auf den Fall einer von Neumann-Algebra zurückziehen zu können.

Als Folgerung von Satz 3.4.27 und von dem von Neumann'schen Dichte-Theorem 3.7.6 hat man allgemein:

3.7.19 Satz (PALMER [232](2001), Seite 890). *Sei A eine $*$ -Unteralgebra von $\mathcal{L}(H)$ für einen Hilbertraum H über $\mathbb{C}$. Sei p die Projektion auf den abgeschlossenen Unterraum $\text{cl}(AH)$ (also die Projektion von Satz 3.4.27). Dann ist $p \in \text{cl}(wo; A)$ und p ist Eins in $\text{cl}(wo; A)$. Des Weiteren ist $\{T \upharpoonright pH : T \in \text{cl}(wo; A)\}$ eine von Neumann-Algebra über $\mathbb{C}$ auf pH und A'' ist gleich $\mathbb{C}(\text{Id}_H - p) \oplus \text{cl}(wo; A)$.*

3.7.20 (LI [198](2003), Seite 67; LI [197](1992), Seite 17; CONWAY [53](1999), Seite 246). Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Sei $p \in \text{Proj}(A)$. Dann gilt:

- (a) A_p ist eine von Neumann-Algebra auf pH über $\mathbb{K}$ und es gilt $(A_p)' = (A')_p$.
- (b) Bezeichnet Z das Zentrum von A , so ist das Zentrum von A_p gleich Z_p .

3.7.21. Dass eine W^* -Algebra A über $\mathbb{K}$ auf einem Hilbertraum H ein Dualraum ist, ist klar: Man betrachte das Dualsystem $(\mathcal{L}(H), \mathcal{L}(H)_*)$. Nach Satz 2.4.36(a) gilt $(N(H)/A^\circ)^* \cong A^{\circ\circ}$ und da A schwach* abgeschlossen ist, folgt mit dem Bipolarensatz 2.5.6, dass A der Dualraum von $N(H)$ modulo dem Annihilator von A ist, soll heißen, A ist der Dualraum von $N(H)/A^\circ$.

Shōichirō Sakai formulierte 1956 eine abstrakte Charakterisierung der W^* -Algebren über $\mathbb{C}$, die 1996 von Isidro und Rodríguez Palacios auch auf den reellen Fall erweitert worden ist.

3.7.22 Satz. *Sei A eine C^* -Algebra über $\mathbb{K}$, die im Fall von $\mathbb{K} = \mathbb{C}$ mit ihrer intrinsischen C^* -Norm versehen sei. Dann sind äquivalent:*

- (a) *A ist eine W^* -Algebra über $\mathbb{K}$.*
- (b) *A hat ein Prädual.*

Beweis. Im Fall von $\mathbb{K} = \mathbb{C}$ siehe TAKESAKI [274](2001), Seite 133, Theorem 3.5. Siehe auch PEDERSEN [233](1979), Seite 67, Theorem 3.9.8, der (b) nach (a) mittels der monotonen Vollständigkeit — siehe hier Definition 3.7.32 — zeigt. Für den Fall $\mathbb{K} = \mathbb{R}$ hat man: CHU, DANG, RUSSO und VENTURA [48](1993), Seite 103, Korollar 2.5, zeigen, dass A genau dann eine W^* -Algebra über $\mathbb{R}$ ist, wenn A ein Prädual hat und A mit ihrer schwach*-Topologie eine topologische Algebra ist. ISIDRO und RODRÍGUEZ PALACIOS [148](1996), Seite 3410, Theorem 1.11, zeigen die Implikation von (b) nach (a). $\square$

3.7.23 (MATHIEU [208](1998)). Satz 3.7.22 zeigt, dass die *uwo*-stetigen *-Homomorphismen zwischen von Neumann-Algebren “die richtigen“ sind: Sei $T: A \rightarrow B$ ein unitaler *uwo*-stetiger *-Homomorphismus zwischen zwei von Neumann-Algebren A und B . Dann ist $T(A)$ eine von Neumann-Unteralgebra von B . (Siehe auch Satz 3.7.40.)

3.7.24 Bemerkung. Jede W^* -Algebra über $\mathbb{K}$ ist mit ihrer schwach*-Topologie eine topologische Algebra.

Beweis. Der Fall $\mathbb{K} = \mathbb{R}$ kam im Beweis von Satz 3.7.22 vor. Für den Fall $\mathbb{K} = \mathbb{C}$ siehe SAKAI [253](1971), Seite 18, Theorem 1.7.8. $\square$

3.7.25 Satz. *Sei A eine Banach*-Algebra über $\mathbb{R}$. Dann ist A genau dann eine W^* -Algebra über $\mathbb{R}$, wenn $A_{(\mathbb{C})}$ eine W^* -Algebra über $\mathbb{C}$ ist, so dass A kanonisch isometrisch isomorph in $A_{(\mathbb{C})\mathbb{R}}$ eingebettet ist. Ist A eine W^* -Algebra über $\mathbb{R}$, dann ist A in dieser W^* -Algebra $A_{(\mathbb{C})}$ schwach*-abgeschlossen.*

Sei $a \in A$, dann konvergiert ein Netz in A genau dann schwach gegen a , wenn es in der W^* -Algebra $A_{(\mathbb{C})}$ schwach* gegen a konvergiert.*

Beweis. LI [198](2003), Seite 123 und CHU, DANG, RUSSO und VENTURA [48] (1993), Seite 101. $\square$

3.7.26 Korollar. *Sei A eine C^* -Algebra über $\mathbb{R}$. Dann ist A genau dann eine W^* -Algebra über $\mathbb{R}$, wenn A eine treue Darstellung als eine W^* -Algebra über $\mathbb{R}$ auf einem Hilbertraum über $\mathbb{C}$ besitzt.*

Beweis. LI [198](2003), Seite 123 und CHU, DANG, RUSSO und VENTURA [48] (1993), Seite 103, Korollar 2.5. $\square$

3.7.27 Satz. *Jede W^* -Algebra über $\mathbb{K}$ besitzt eine Eins.*

Beweis. Für den Fall $\mathbb{K} = \mathbb{C}$ siehe TAKESAKI [274](2001), Seite 137. Für den Fall $\mathbb{K} = \mathbb{R}$ siehe LI [198](2003), Seite 123; man beachte auch ISIDRO und RODRÍGUEZ PALACIOS [148](1996), Seite 3409, Satz 1.9, die zum Beweis die Theorie der JB^* -Tripel verwenden. $\square$

3.7.28 Definition. Eine W^* -Algebra über $\mathbb{K}$ heißt ein *Faktor*, wenn sein Zentrum gleich $\mathbb{K} \cdot e$ ist.

3.7.29 Faktor I. (a) Das Zentrum Z einer von Neumann-Algebra A über $\mathbb{K}$ ist selbst auch eine von Neumann-Algebra über $\mathbb{K}$: $Z' = (A' \cap A)' \supseteq A'' \cup A' = A \cup A'$, also $Z'' \subseteq (A \cup A')' = A' \cap A'' = Z \subseteq Z''$ und somit $Z = Z''$.

(b) Ein Faktor ist also eine W^* -Algebra, die so nichtkommutativ wie nur möglich ist. (PEDERSEN [234](1989), Seite 174.)

(c) Sei H ein Hilbertraum über $\mathbb{R}$ und A eine von Neumann-Algebra auf H über $\mathbb{R}$. Es sei an Satz 3.7.12 erinnert. Bezeichne Z das Zentrum von A . Dann ist das Zentrum von $A_{(\mathbb{C})}$ gleich $Z_{(\mathbb{C})}$. Insbesondere ist A genau dann ein Faktor, wenn $A_{(\mathbb{C})}$ ein Faktor ist. (LI [198](2003), Seite 67.)

(d) Ist H ein Hilbertraum über $\mathbb{C}$, so ist $\mathcal{L}(H)$ ein Faktor. (KADISON und RINGROSE, Seite 182.)

(e) Für ein Beispiel eines Faktors, der nicht $*$ -isomorph zu $\mathcal{L}(H)$ für irgendeinen Hilbertraum H über $\mathbb{C}$ ist, siehe MURPHY [217](2004), Seite 182.

(f) Sei H ein Hilbertraum H über $\mathbb{K}$. Dann gilt $K(H)' = \mathbb{K} \cdot \text{Id}_H$. Da $K(H)$ gemäß Bemerkung 2.10.10 und Satz 3.7.22 nur genau für endlichdimensionales H eine W^* -Algebra über $\mathbb{K}$ ist, ist $K(H)$ auch genau nur für endlichdimensionales H ein Faktor.

(g) Kadison zeigte 1951, dass eine von Neumann-Algebra A über $\mathbb{C}$ genau dann ein Faktor ist, wenn $A_{(*)}$ ein Antiverband ist.

3.7.30 (PALMER [232](2001), Seite 895). Sei H ein Hilbertraum über $\mathbb{C}$. In 3.5.48 wurde erwähnt, dass $\mathcal{L}(H)_{(*)}$ ein Antiverband ist. 1946 zeigte JEAN-PIERRE VIGIER, dass gewisse Netze in $\mathcal{L}(H)$ ein Suprema haben; es gilt nämlich das folgende Monotonieprinzip:

3.7.31 Satz von Vigier. Sei H ein Hilbertraum über $\mathbb{K}$ und $(T_\nu)_{\nu \in I}$ ein monoton steigendes, nach oben ordnungsbeschränktes Netz in $\mathcal{L}(H)_{(*)}$. Dann existiert ein $T \in \mathcal{L}(H)_{(*)}$ mit:

$$(a) \quad T = \sup T_\nu.$$

$$(b) \quad T_\nu \xrightarrow{\tau} T \text{ für alle } \tau \in \{wo, so, w^*, uso, s^*o, us^*o\}.$$

Beweis. (Nach MATHIEU [208] (1998), Seite 335; CONWAY [53] (1999); STRĂTILĂ und ZSIDÓ [270](1979), Seite 39; SUNDER [271](1987), Seite 17.)

Sei $S \in \mathcal{L}(H)_{(*)}$ eine obere Schranke von $(T_\nu)_{\nu \in I}$. Ohne Einschränkung kann man $T_\nu \geq 0$ für alle $\nu \in I$ annehmen; sonst wähle man ein beliebiges $\mu \in I$, und der Beweis für das Netz $(T_\nu - T_\mu)_{\nu \geq \mu}$ beweist die Aussage für das Ausgangsnetz. Für alle $\nu \in I$ hat man also $0 \leq T_\nu \leq S$ und somit nach 3.5.47 auch $\|T_\nu\| \leq \|S\|$. Wegen der Äquivalenz $R \in \text{Pos}(*, \mathcal{L}(H)) \Leftrightarrow \langle Rx, x \rangle \geq 0 \forall x \in H$, ist für jedes $x \in H$ $(\langle T_\nu x, x \rangle)_{\nu \in I}$ ein monoton steigendes Netz in $\mathbb{R}$, das durch $\|S\| \|x\|^2$ beschränkt ist. Somit ist jedes solche Netz konvergent und man kann definieren:

$$m(x) := \sup_\nu \langle T_\nu x, x \rangle = \lim_\nu \langle T_\nu x, x \rangle \quad \text{für alle } x \in H$$

und die somit definierte Funktion $m: H \rightarrow \mathbb{R}$ ist positiv und im Fall von $\mathbb{K} = \mathbb{R}$ quadratisch. Per Polarisierung definiert man die gleichnamige Funktion $m: H \times H \rightarrow \mathbb{K}$, wobei im Fall von $\mathbb{K} = \mathbb{R}$ gemäß 2.2.40

$$m(x, y) := \frac{1}{4}(m(x+y) - m(x-y)) \quad \text{für alle } x, y \in H$$

und im Fall von $\mathbb{K} = \mathbb{C}$ gemäß 2.2.41

$$m(x, y) := \frac{1}{4} \sum_{k=0}^3 i^k m(x + i^k y) \quad \text{für alle } x, y \in H$$

gesetzt wird. $m: H \times H \rightarrow \mathbb{K}$ ist dann eine beschränkte, hermitesche, positive Sesquilinearform mit $\|m\| \leq \|S\|$. Für beliebiges, aber festes $x \in H$ ist die mit m_x bezeichnete Abbildung $y \mapsto m(x, y)$ eine stetige Linearform auf H mit Norm höchstens $\|m\|\|x\|$. Nach dem Darstellungssatz von Fréchet-Riesz, existiert daher genau ein $z_x \in H$, so dass $\langle y, z_x \rangle = m(x, y)$, also $\langle z_x, y \rangle = m(x, y)$ für alle $y \in H$ gilt. Setzt man $Tx := z_x$, so erhält man die Abbildung $T \in \mathcal{L}(H)$ mit $\|T\| = \|m\|$ und $\langle Tx, y \rangle = m(x, y)$ für alle $x, y \in H$. Dass die Abbildung m positiv ist, heißt gerade genau, dass $T \geq 0$ ist; daher ist $T \in \mathcal{L}(H)_{(*)}$. Da für alle $\nu \in I$, $x \in H$ $\langle Tx, x \rangle = m(x, x) \geq \langle T_\nu x, x \rangle$, also $\langle (T - T_\nu)x, x \rangle \geq 0$, ist T eine obere Schranke des Netzes $(T_\nu)_{\nu \in I}$. Ist andererseits $R \in \mathcal{L}(H)_{(*)}$ eine obere Schranke für $\{T_\nu : \nu \in I\}$, so gilt für jedes $x \in H$ $\langle Tx, x \rangle = m(x, x) = \lim_\nu \langle T_\nu x, x \rangle \leq \langle Rx, x \rangle$, also $T \leq R$, $T = \sup T_\nu$ in $\mathcal{L}(H)_{(*)}$ und (a) ist gezeigt. Nebenbei bemerkt man hier für den Fall $\mathbb{K} = \mathbb{C}$ $T_\nu \rightarrow T$ in der wo -Topologie. Da nach 3.6.11 auf beschränkten Teilmengen von $\mathcal{L}(H)$ die wo -Topologie mit der w^* -Topologie übereinstimmt, hat man somit für den Fall $\mathbb{K} = \mathbb{C}$ auch $T_\nu \rightarrow T$ in der w^* -Topologie. Dies lässt sich aber auch wie folgt direkt und allgemein gültig für den $\mathbb{K}$ -Fall zeigen:

Vorab aber erst noch eine kleine Zwischenrechnung: Für alle $\nu \in I$ gelten wegen $0 \leq T_\nu, -T_\nu \leq 0, T - T_\nu \leq T$ und $T_\nu \leq T, 0 \leq T - T_\nu$, die Ungleichungen $0 \leq T - T_\nu \leq T$, also $\|T - T_\nu\| \leq \|T\|$. Nun zurück zu der eigentlichen Rechnung. Seien $(x_n)_{n \in \mathbb{N}}$ und $(y_n)_{n \in \mathbb{N}}$ zwei Folgen in H mit $\sum_{n \in \mathbb{N}} \|x_n\|^2 < \infty$ und $\sum_{n \in \mathbb{N}} \|y_n\|^2 < \infty$. Sei $N \in \mathbb{N}$ und $\nu \in I$. Dann gilt: $\sum_{n \in \mathbb{N}} |\langle (T - T_\nu)x_n, y_n \rangle| \leq \sum_{n=0}^N |\langle (T - T_\nu)x_n, y_n \rangle| + \|T\| \sum_{n=N+1}^\infty \|x_n\| \|y_n\| =: R_{N,\nu}$. Für jedes $N \in \mathbb{N}$ erhält man somit ein Netz $(R_{N,\nu})_{\nu \in I}$ in $\mathbb{R}$, das gegen $R_N := \|T\| \sum_{n=N+1}^\infty \|x_n\| \|y_n\|$ konvergiert. Da $(R_N)_{N \in \mathbb{N}}$ eine Nullfolge ist, gilt $T_\nu \rightarrow T$ in der w^* -Topologie.

Aussage (b) des Satzes von Vigier wird in der Literatur oft nur für die so -Konvergenz formuliert. Diese zeigt sich beispielsweise wie folgt: Nach Satz 3.5.43 existiert zu jedem $Q \in \text{Pos}(*; \mathcal{L}(H))$ genau ein $R \in \text{Pos}(*; \mathcal{L}(H))$ mit $Q = R^2$ und im Folgenden sei $Q^{1/2}$ für solch ein R geschrieben. Für alle $\nu \in I$, $x \in H$ gilt: $\|(T - T_\nu)x\|^2 = \|(T - T_\nu)^{1/2}(T - T_\nu)^{1/2}x\|^2 \leq \|(T - T_\nu)^{1/2}\|^2 \|(T - T_\nu)^{1/2}x\|^2 = \|(T - T_\nu)^{1/2}\|^2 \|(T - T_\nu)^{1/2}\| \|(T - T_\nu)^{1/2}x\|$ (C^* -Bedingung) $= \|(T - T_\nu)\| \|(T - T_\nu)^{1/2}x\|^2 = \|(T - T_\nu)\| \langle (T - T_\nu)x, x \rangle \leq \|T\| \langle (T - T_\nu)x, x \rangle = \|T\| (m(x, x) - \langle T_\nu x, x \rangle) \rightarrow 0$, also $T_\nu \rightarrow T$ in der so -Topologie.

Auf beschränkten Teilmengen von $\mathcal{L}(H)$ stimmt gemäß 3.6.11 die so -Topologie mit der uso -Topologie überein, womit man auch $T_\nu \rightarrow T$ in der uso -Topologie vorliegen hat.

Da alle T_ν , $\nu \in I$ und T selbstadjungiert sind, folgt aus $T_\nu \rightarrow T$ in der uso -Topologie auch $T_\nu \rightarrow T$ in der us^*o -Topologie und damit wieder nach 3.6.11 auch $T_\nu \rightarrow T$ in der s^*o -Topologie $\square$

Dem Beweis entnimmt man folgende Anmerkung: Falls für alle Glieder T_ν des Netzes $(T_\nu)_{\nu \in I}$ die Ungleichung $T_\nu \geq 0$ gilt, so gilt für das Supremum T ebenfalls $T \geq 0$.

3.7.32 Definition. Eine C^* -Algebra A über $\mathbb{K}$ heißt *monoton vollständig*, wenn der geordnete Vektorraum $(A_{(*)}, \text{Pos}(*; A))$ monoton vollständig ist. (Siehe Definition 2.2.16.)

3.7.33 Korollar. Jede W^* -Algebra über $\mathbb{K}$ ist monoton vollständig.

3.7.34 (LI [197](1992), Seite 16; LI [198](2003), Seite 65). Als eine weitere Folgerung aus dem Satz von Vigier ergibt sich die folgende Diskussion.

Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Dann ist die Menge $\text{Proj}(A) = \text{OProj}(H) \cap A$ der Projektionen von A , versehen mit der mit $\leq$ bezeichneten von $\text{Pos}(*; A)$ induzierten partiellen Ordnung, ein Verband. Diese Aussage stimmt im Allgemeinen nicht mehr, wenn man anstelle der Menge $\text{Proj}(A)$ die Menge $A_{(*)}$ betrachtet (siehe 3.5.48) oder, wenn man A als eine C^* -Algebra über $\mathbb{K}$ von Operatoren auf H voraussetzt. Allgemeiner gilt: Sei p_ν , $\nu \in I$, eine Familie von Projektionen von

A . Dann existieren sowohl das Supremum $\sup_{\nu \in I} p_\nu$ als auch das Infimum $\inf_{\nu \in I} p_\nu$ und beide sind wiederum Projektionen von A . Genauer gilt für solch eine Familie und zur Beweisskizze folgendes: Für jedes endliche $F \subseteq I$ setze $p_F := \sup_{i \in F} p_i$. Für nicht leeres I ist (p_F) ein monoton steigendes Netz, welches durch Mengeninklusion auf I gerichtet ist. Da es in $\text{OProj}(H)$ nach oben durch Id_H beschränkt ist, existiert nach dem Satz von Vigier das Supremum in $\text{OProj}(H)$; ebenfalls nach diesem Satz so -konvergiert dieses Netz gegen dieses Supremum, welches wegen der so -Abgeschlossenheit von A ein Element von A ist. Dieses Supremum ist die Projektion von H auf den Abschluss von $\text{span } \bigcup_{\nu \in I} p_\nu H$. Analog hat man $\inf_{\nu \in I} p_\nu = so - \lim_{F \subseteq I \text{ endlich}} \inf_{i \in F} p_i =$ die Projektion von H auf $\bigcap_{\nu \in I} p_\nu H$. Die jeweilige Konvergenz gilt in jeder der im Satz von Vigier angegebenen Operatortopologie.

3.7.35. Sei H ein Hilbertraum über $\mathbb{K}$, M eine Teilmenge von $\mathcal{L}(H)_{(*)}$ und τ eine Topologie auf $\mathcal{L}(H)$. Setze:

$$M^{m,\tau} := \{T \in \mathcal{L}(H)_{(*)} : \text{Es gibt ein monoton steigendes Netz } (T_\nu)_{\nu \in I} \text{ mit} \\ T_\nu \in M \text{ für alle } \nu \in I, \text{ das } T \text{ als } \tau\text{-Grenzwert hat} \}$$

und

$$M_{m,\tau} := \{T \in \mathcal{L}(H)_{(*)} : \text{Es gibt ein monoton fallendes Netz } (T_\nu)_{\nu \in I} \text{ mit} \\ T_\nu \in M \text{ für alle } \nu \in I, \text{ das } T \text{ als } \tau\text{-Grenzwert hat} \}.$$

Es besteht folgender Zusammenhang mit den in 1.1.15 definierten Mengen M^m und M_m :

- (a) M monoton steigend vollständig $\Leftrightarrow M = M^m \Leftrightarrow M = M^{m,so}$.
- (b) M monoton fallend vollständig $\Leftrightarrow M = M_m \Leftrightarrow M = M_{m,so}$.

Falls M ein Untervektorraum von $\mathcal{L}(H)_{(*)}$ ist, dann gilt

$$(c) \quad M \text{ monoton vollständig} \Leftrightarrow M = M^m \Leftrightarrow M = M_m \Leftrightarrow M = M^{m,so} \Leftrightarrow M = M_{m,so}.$$

Beweis. Es genügt (a) zu zeigen, da (b) analog gezeigt wird. (c) folgt dann aus 2.2.15. Zum ersten „ $\Leftrightarrow$ “: Es ist nur „ $\Leftarrow$ “ zu zeigen. Sei dazu $(T_\nu)_{\nu \in I}$ ein nach oben ordnungsbeschränktes, monoton aufsteigendes Netz in M . Nach dem Satz von Vigier 3.7.31 existiert ein $T \in \mathcal{L}(H)_{(*)}$ mit $T = \sup T_\nu$. Also $T \in M^m$. Gilt $M = M^m$, dann $\sup T_\nu \in M$ und M ist monoton aufsteigend vollständig.

Zum zweiten „ $\Leftrightarrow$ “: „ $\Rightarrow$ “: Gelte $M = M^m$. Es ist $M^{m,so} \subseteq M$ zu zeigen. Sei dazu $T \in M^{m,so}$ beliebig. Es gibt also ein Netz $(T_\nu)_{\nu \in I}$ in M , das monoton steigend ist und T als so -Grenzwert hat.

Nun ist zu beachten, dass im Gegensatz zu Folgen im Allgemeinen Netze, die so -konvergent in $\mathcal{L}(H)$ sind, nicht beschränkt zu sein brauchen. Hier gilt aber folgendes: Da $T_\nu \xrightarrow{so} T \Leftrightarrow \forall x \in H : \|(T - T_\nu)x\| \rightarrow 0 \Leftrightarrow \forall x \in H : Tx = \lim_\nu T_\nu x \Rightarrow \forall x, y \in H : \langle Tx, y \rangle = \lim_\nu \langle T_\nu x, y \rangle$, also $\forall x \in H \quad \forall \nu \in I : \langle (T - T_\nu)x, x \rangle \geq 0$ und T ist eine obere Schranke von $(T_\nu)_{\nu \in I}$.

Nach dem Satz von Vigier 3.7.31 gilt $T = \sup T_\nu$ und damit $T \in M^m$. Wegen $M = M^m$ also $T \in M$.

„ $\Leftarrow$ “: Gelte $M = M^{m,so}$. Es ist $M^m \subseteq M$ zu zeigen. Sei $T \in M^m$. Es gibt also ein nach oben ordnungsbeschränktes, monoton steigendes Netz $(T_\nu)_{\nu \in I}$ in M mit T als Supremum. Nach dem Satz von Vigier 3.7.31 ist T der so -Grenzwert des Netzes $(T_\nu)_{\nu \in I}$. Also ist $T \in M^{m,so}$ und nach Voraussetzung $T \in M$. $\square$

3.7.36. Sei H ein Hilbertraum über $\mathbb{C}$. In 3.6.5 wurde die Menge $\mathcal{L}(H)_*$ definiert. Ist A eine W^* -Algebra über $\mathbb{C}$, so wird in den folgenden beiden Nummern 3.7.37 und 3.7.40 gezeigt, wie man in abstrakter Weise mit dem punktierten konvexen Kegel $\text{Pos}(*; A)$ ein dem $\mathcal{L}(H)_*$ entsprechendes A_* definieren kann, wobei im Fall von $A = \mathcal{L}(H)$ die in 3.6.5 definierte Menge $\mathcal{L}(H)_*$ übereinstimmt mit der in 3.7.37 definierten Menge A_* .

3.7.37 Definition (PALMER [232](2001), Seite 895). Sei A eine monoton vollständige C^* -Algebra über $\mathbb{C}$. Ein positives $\ell \in A^\#$ wird *normal* genannt, wenn für alle Suprema h von beschränkten, monoton steigenden Netzen $\{h_\nu\}_{\nu \in I} \subseteq (A_{(*)}; \leq_*)$ die Gleichung $\ell(h) = \sup \{\ell(h_\nu) : \nu \in I\}$ gilt. Die lineare Hülle in $A^\#$ aller normalen positiven linearen Funktionale von A wird mit A_* bezeichnet. Ein $\ell \in A^\#$ wird *normal* genannt, wenn es ein Element von A_* ist.

3.7.38. Für die Menge A_* aus Definition 3.7.37 gilt: $A_* \subseteq A^*$.

3.7.39 Satz. Sei A eine C^* -Algebra über $\mathbb{C}$ mit Eins und bezeichne A_* die Menge von Definition 3.7.37. A ist genau dann eine W^* -Algebra über $\mathbb{C}$, wenn der punktierte konvexe Kegel $A_{(*)}$ monoton vollständig ist und die Menge $A_* \cap S_{A^*}$ aller normalen Zustände von A die Punkte von A trennt, das heißt, für alle $x \in A^\times$ existiert ein $\ell \in A_*$ mit $\|\ell\| = 1$ und $\ell(x) \neq 0$.

Beweis. ALFSEN und SHULTZ [8](2001), Seite 108, Theorem 2.93, zitieren als Beweis PEDERSEN [233](1979), 2.4.4, KADISON und RINGROSE [167], Übung 7.6.38 und als Ursprung KADISON [166](1956). Siehe auch BRATTELI und ROBINSON [37](1987), Seite 78. (Es sei daran erinnert, dass nach Bemerkung 3.2.25(a) $A_{(*)}$ ein punktierter spitzer konvexer Kegel ist.) $\square$

3.7.40 Satz (PALMER [232](2001), Seite 895). Sei A eine W^* -Algebra über $\mathbb{C}$, versehen mit ihrer intrinsischen C^* -Norm. Sei A_* die Menge aus Definition 3.7.37. Dann gilt:

- (a) Sei X ein Prädual von A per $T: A \rightarrow X^*$. Dann ist $T^*(i_X(X))$ der abgeschlossene Unterraum A_* , das heißt, X und A_* sind stark eindeutige Präduale.
- (b) Sei $T: A \rightarrow \mathcal{L}(H)$ eine treue $*$ -Darstellung von A als eine von Neumann-Algebra über $\mathbb{C}$. Dann ist T_A abgeschlossen in der ultraschwachen Operator-Topologie von $\mathcal{L}(H)$ und die Funktionale $\ell \in A^\#$, die in der von der ultraschwachen Operator-Topologie von $\mathcal{L}(H)$ durch T induzierten Topologie auf A stetig sind, sind genau die $\ell \in A_* \subseteq A^*$.
- (c) Sei $T: A \rightarrow \mathcal{L}(H)$ eine treue $*$ -Darstellung von A . Dann ist T_A genau dann eine von Neumann-Algebra über $\mathbb{C}$, wenn T $\sigma(A, A_*)$ - w^* -stetig ist.

3.7.41 Bemerkung und Definition (PALMER [232](2001), Seite 896). Sei A eine W^* -Algebra über $\mathbb{C}$, versehen mit ihrer intrinsischen C^* -Norm. Die in Satz 3.7.40(b) auf A induzierte Topologie ist wohldefiniert und heißt die *ultraschwache Topologie* von A ; sie ist gerade die $\sigma(A, A_*)$ -Topologie, das heißt, die schwach*-Topologie von A .

3.7.42 Satz. Ein Prädual einer W^* -Algebra über $\mathbb{K}$ ist ein stark eindeutiges Prädual.

Beweis. Für den Fall $\mathbb{K} = \mathbb{C}$ ist das die Aussage von Satz 3.7.40; sie wird explizit bewiesen bei SAKAI [253](1971), Seite 30, Korollar 1.13.3. Für den Fall $\mathbb{K} = \mathbb{R}$ siehe ISIDRO und RODRÍGUEZ PALACIOS [148](1996), Seite 3409, Bemerkung 1.7. $\square$

3.7.43 Bemerkung. Das Prädual F einer W^* -Algebra A über $\mathbb{K}$ kann gemäß Lemma 2.10.3 — und wird es üblicherweise auch — als Teilmenge von A^* betrachtet werden, also als die Menge A_* der $\sigma(A, F)$ -stetigen Funktionale aus A^* .

3.7.44 (MATHIEU [208](1998), Seite 346 und LI [198](2003), Seiten 99-103). Als eine Anwendung des Kaplansky'schen Dichte-Theorems 3.7.9 hat man: Sei (T, H) eine universelle Darstellung einer C^* -Algebra A über $\mathbb{K}$. Dann gilt $T(A)'' \cong A^{**}$. In diesem Zusammenhang sei an Satz 3.5.55 erinnert.

3.7.45 Satz. *Sei A eine von Neumann-Algebra auf einem Hilbertraum H über $\mathbb{K}$.*

- (a) *Im Fall von $\mathbb{K} = \mathbb{R}$ gilt: $A_{(*)} = \text{cl span Proj}(A) = \text{span Pos}(*W; A) = \text{Pos}(*W; A) - \text{Pos}(*W; A)$ und $A = \text{cl span } A_{\text{unitär}}$, wobei der Fall $A \neq \text{span}(A_{\text{unitär}})$ vorliegen kann.*
- (b) *Im Fall von $\mathbb{K} = \mathbb{C}$ gilt: $A = \text{cl span Proj}(A) = \text{span Pos}(W; A) = \text{span } A_{\text{unitär}}$. Des Weiteren gilt $A_{\text{unitär}} = \exp(iA_{(*)})$. (Siehe auch 3.2.19.)*

Beweis. Für den Fall $\mathbb{K} = \mathbb{R}$ siehe LI [198](2003), Seite 65. Für den Fall $\mathbb{K} = \mathbb{C}$: Zum Beweis von $A = \text{cl span Proj}(A)$ siehe etwa CONWAY [53](1999), Seite 61, Satz 13.3.e oder ALFSEN und SHULTZ [8](2001), Seite 81, Lemma 2.37. Siehe auch die Bemerkung über unitäre Elemente in 3.2.19. Für, dass jedes unitäre Element die Form $\exp(T)$, T schiefselbstadjungiert, hat, siehe KADISON und RINGROSE [167](1986), Seite 728. $\square$

3.7.46 Bemerkungen. (a) SUNDER [271](1987), Seite 14, bemerkt: Gerade so wie $L_\infty(X, \mu)$ (als ein Norm-abgeschlossener Unterraum) von den Indikatorfunktionen erzeugt wird, wird jede von Neumann-Algebra über $\mathbb{C}$ (als ein Norm-abgeschlossener Unterraum) von seinen Projektionen erzeugt.

(b) Der Unterschied zwischen $\mathbb{K} = \mathbb{R}$ und $\mathbb{K} = \mathbb{C}$ kommt in Satz 3.7.45 daher, dass man für $\mathbb{K} = \mathbb{C}$ gemäß Bemerkung 3.1.5(e) in einem Vektorraum X , der mit einer Involution versehen ist, jedes $a \in X$ als Summe zweier selbstadjungierter Elemente schreiben kann, was im Allgemeinen beim Fall $\mathbb{K} = \mathbb{R}$ nicht zutrifft.

3.7.47 Satz (LI [198](2003), Seite 152). *Sei H ein Hilbertraum über $\mathbb{C}$. Sei A eine von Neumann-Algebra auf $H_{\mathbb{R}}$ über $\mathbb{R}$. Dann gilt $\text{cl}(wo; A_{\text{unitär}}) = B_A$*

3.7.48 Äquivalente Projektionen (MATHIEU [208](1998), Seiten 352-354). Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Seien $p, q \in A$ zwei Projektionen. Dann heißen p und q *äquivalent* (bezüglich A), in Zeichen $p \sim q$, wenn es eine partielle Isometrie $u \in A$ mit $u^*u = p$ und $uu^* = q$ gibt; man schreibt dann anstelle von $p \sim q$ auch genauer $p \sim_u q$ oder noch genauer $p \sim_u q (A)$ und klarer $p \sim_u q$ (bezgl. A) (im Englischen *rel. A*)).

$\sim$ ist eine Äquivalenzrelation auf $\text{Proj}(A)$ und wird auch *Murray-von Neumann-Äquivalenz* genannt. Man beachte, dass, damit zwei Projektionen einer von Neumann-Algebra äquivalent sind, nicht nur ihre Hilbertraumdimensionen ihrer Bildräume gleich sein müssen, sondern auch mindestens eine der die beiden Bildräume vermittelnden partiellen Isometrien ein Element der vorgegebenen von Neumann-Algebra sein muss.

Für alle $x \in A$ gilt

$$s_\ell(x) \sim s_r(x). \quad (3.20)$$

Es gilt die sogenannte *Formel von Kaplansky*:

$$(p \vee q) - p \sim q - (p \wedge q). \quad (3.21)$$

Falls $\|p - q\| < 1$ vorliegt, dann sind p und q sowohl unitär äquivalent (Definition 3.4.21) als auch äquivalent.

Sind p und q per einer partiellen Isometrie $u \in A$ äquivalent, so gilt zwar $p = u^*qu$, aber p und q müssen nicht unbedingt unitär äquivalent sein.

Durch

$$\begin{aligned} p \precsim q &: \Leftrightarrow \exists r \in \text{Proj}(A) : p \sim r \text{ und } r \leq q \\ &\Leftrightarrow \exists r \in \text{Proj}(A) : p \leq r \text{ und } r \sim q \end{aligned}$$

wird auf $\text{Proj}(A)$ eine Präordnung erklärt; da sie die sogenannte Cantor-Bernstein-Eigenschaft

$$(p \precsim q \text{ und } q \precsim p) \Rightarrow p \sim q$$

aufweist, ist durch sie auf der Menge $\text{Proj}(A)/\sim$ der durch $\sim$ induzierten Äquivalenzklassen von $\text{Proj}(A)$ sogar eine partielle Ordnung erklärt.

Beweis. Die Reflexivität ist klar. Zur Transitivität: Seien dazu $p, q, r \in \text{Proj}(A)$ mit $p \precsim q$ und $q \precsim r$. Es gibt also $a, b \in \text{Proj}(A)$ mit $p \sim a$, $a \leq q$, $q \sim b$, $b \leq r$. Also existieren partielle Isometrien $u, v \in A$ mit $p = u^*u$, $a = uu^*$, $q = v^*v$ und $b = vv^*$. Dann gilt $vav^*b = vav^*vv^* = vav^*$, also $vav^* \leq b$. vu ist eine partielle Isometrie, da $(vu)^*vu$ eine Projektion ist: $((vu)^*vu)^2 = (vu)^*vu(vu)^*vu = u^*v^*vu u^*v^*vu = u^*(v^*v)(uu^*)(v^*vu) = u^*qa(v^*vu) = u^*a(v^*vu) = (u^*uu^*)(v^*vu) = u^*v^*vu = (vu)^*vu$. Die Anfangsprojektion von vu ist $(vu)^*vu = u^*a(v^*v)u = u^*(aq)u = u^*au = u^*uu^*u = p$ und die Endprojektion von vu ist $vu(vu)^* = vu u^*v^* = vav^*$. Somit $p \sim_{vu} vav^* \leq b \leq r$, also $p \precsim r$.

Für die Antisymmetrie auf den Äquivalenzklassen siehe etwa MATHIEU [208] (1986), Seite 354 oder KADISON und RINGROSE [167] (1986), Seite 406. $\square$

Seien wie gehabt p und q zwei Projektionen von A . p heißt *kleiner als q relativ zu A* (im Englischen *weaker than*), falls $p \precsim q$ gilt. Gilt zwar $p \precsim q$, aber nicht $p \sim q$, so wird diese Situation mit $p \prec q$ bezeichnet und sagt, p sei *echt kleiner als q relativ zu A* (im Englischen *strictly weaker than*). Andere Sprechweisen für $p \precsim q$ sind zum Beispiel: p ist *subäquivalent* zu q , q *dominiert p relativ* (im Englischen *relatively dominates*) und p wird von q *dominiert* (im Englischen *is dominated by*).

3.7.49 (MATHIEU [208] (1998), Seite 356). Sei A eine von Neumann-Algebra über $\mathbb{K}$. Seien $p, q \in \text{Proj}(A)$ mit $p \sim q$ und sei $z \in A' \cap \text{Proj}(A)$. Dann gilt $pz \sim qz$.

Beweis. Es ist $p = u^*u$, $q = uu^*$ für eine partielle Isometrie $u \in A$. Es ist $pz = u^*uz = u^*z^*uz = (uz)^*(uz)$, womit uz eine partielle Isometrie von A ist. $qz = uu^*z = uzu^*z^* = (uz)(uz)^*$. $\square$

3.7.50 Definition. Sei H ein Hilbertraum über $\mathbb{K}$, A eine von Neumann-Algebra auf H über $\mathbb{K}$ und $T \in A$. Dann heißt die zentrale Projektion

$$\begin{aligned} c(T) &:= \min \{z \in A' \cap \text{Proj}(A) : Tz = T\} \\ &= \text{Id}_H - \max \{z \in A' \cap \text{Proj}(A) : Tz = 0\} \end{aligned}$$

der *zentrale Träger* (im Englischen *central support*, *central carrier* oder auch *central cover*) von T (bezüglich A). (Falls also $p \in A$ eine Projektion ist, gilt $c(p) = \min \{z \in A' \cap \text{Proj}(A) : pz \leq p\}$.)

3.7.51 Bemerkung. Sei H ein Hilbertraum über $\mathbb{K}$, A eine von Neumann-Algebra auf H über $\mathbb{K}$ und $T \in A$. Dann gilt

$$\text{ran } c(T) = \text{cl span}((AT)H).$$

Zum Beweis siehe etwa KADISON und RINGROSE [167] (1983), Seite 333, wobei für den dortigen Beweis noch angemerkt sei, dass nicht nur der Raum $\text{cl span}((AT)H)$ unter A

und A' invariant ist, sondern auch sein orthogonaler Komplementärraum, womit man dann hat, dass jedes $T \in A$ sowohl auf $\text{cl span}((AT)H)$ als auch auf $(\text{cl span}((AT)H))^\perp$ mit der Projektion auf $\text{cl span}((AT)H)$ kommutiert.

Des Weiteren gilt: $c(T) = c(T^*)$ (Beweisskizze: $T^* \circ \langle \cdot, h \rangle = \langle T^* \cdot, h \rangle = \langle \cdot, Th \rangle \leftrightarrow Th$), $c(T) = c(T^*T) = c(TT^*)$ (siehe CONWAY [53](1999), Seite 245) und $c(T) = c(\text{pker}(T)^\perp) = c(s_\ell(T)) = c(s_r(T))$.

3.7.52 (LI [198](2003), Seite 71). Sei A eine von Neumann-Algebra über $\mathbb{K}$. Seien $p, q \in \text{Proj}(A)$ mit $p \preceq q$. Dann gilt $c(p) \leq c(q)$. Insbesondere folgt aus $p \sim q$ die Gleichheit $c(p) = c(q)$ und hieraus insbesondere, dass aus $p \sim 0$ stets die Gleichung $p = 0$ folgt.

3.7.53 Bemerkung (MATHIEU [208](1998), Seite 356; CONWAY [53](1999), Seite 268). Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Seien $x, y \in A$. Dann sind die folgenden Aussagen äquivalent:

- (a) $c(x)c(y) = 0$.
- (b) $M_{x,y} = 0$. (Siehe 3.5.33)
- (c) $xAy = \{0\}$.
- (d) $\nexists p, q \in \text{Proj}(A)^\times, q \leq s_\ell(y), p \leq s_r(x) : p \sim q$.

Beweis. (a) $\Rightarrow$ (b): Für $T \in A$ bezeichne $\text{ann}(T)$ die zu T bezüglich des A zugrunde liegenden Ringes orthogonale Menge, das heißt,

$$\text{ann}(T) = \{x \in A : xT = Tx = 0\}.$$

Betrachte $A(1 - c(T))$. Sei $x \in A(1 - c(T))$. Dann ist $x = y - yc(T)$ für ein $y \in A$. Also $Tx = Ty - Tyc(T) = Ty - Tc(t)y = Ty - Ty = 0$ und $xT = yT - yc(T)T = yT - yT = 0$. Somit gilt $A(1 - c(T)) \subseteq \text{ann}(T)$. Sei nun $x \in \text{ann}(T)$. Es gilt dann auch $xc(T) = c(T)x = c(T)Tx = 0$, also $x = x - xc(T) = x(1 - c(T)) \in A(1 - c(T))$ und somit

$$\text{ann}(T) = A(1 - c(T)).$$

Gelte nun (a), also $c(x) \in \text{ann}(c(y))$. Dies ist äquivalent dazu, dass $c(x) = a(1 - c(c(y))) = a(1 - c(y))$ für ein $a \in A$. Also $c(x) \in \text{ann}(y)$, $c(x)y = 0$ und wegen $M_{x,y}T = xTy = xc(x)Ty = xTc(x)y$ somit $M_{x,y} = 0$.

(b) $\Rightarrow$ (d): Seien $p, q \in \text{Proj}(A)$ mit $q \leq s_\ell(y)$, $p \leq s_r(x)$, $p \sim_u q$ (bezgl. A). Zuerst wird die Implikation $M_{x,y} = 0 \Rightarrow M_{s_r(x), s_\ell(y)} = 0$ bewiesen: Sei dazu $z \in A$. Falls $xz = 0$ gilt, dann auch $x^*xz = 0$, also $zH \subseteq \ker(x^*x) = \text{pker}(x^*x)H$, $\text{pker}(x^*x)z = z$, $0 = (\text{Id}_H - \text{pker}(x^*x))z = s_r(x)z$, siehe die Gleichungen (3.10) in 3.4.20. Falls $zy = 0$ gilt, dann auch $y^*z^* = 0$, $z^*H \subseteq \ker(y^*) = \text{pker}(y^*)H$, $\text{pker}(y^*)z^* = z^*$, $0 = (\text{Id}_H - \text{pker}(y^*))z^* = s_r(y^*)z^* = s_\ell(y)z^*$, $zs_\ell(y) = 0$.

Aus $xy = 0$ für $a \in A$ folgt also (per $z := ay$) erst $s_r(x)ay = 0$ und (per $z := s_r(x)a$) dann $s_r(x)as_\ell(y) = 0$, womit die eingangs erwähnte Implikation gezeigt ist.

Gelte nun (b), dann folgt: $0 = s_r(x)u^*s_\ell(y) = s_r(x)u^*uu^*s_\ell(y) = s_r(x)u^*u u^*uu^*s_\ell(y) = s_r(x)pu^*qs_\ell(y) = pu^*q = u^*uu^*uu^* = u^*$, $u = 0$, $p = 0$, $q = 0$.

(d) $\Rightarrow$ (b): Per Kontraposition. Falls für ein $a \in A$ $xy \neq 0$ ist, muss $xa \text{ pcran}(y) \neq 0$ und damit auch $(\text{Id}_H - \text{pker}(x))a \text{ pcran}(y) \neq 0$. In anderen Zeichen also $s_r(x)as_\ell(y) \neq 0$. Setze $p := s_\ell(s_r(x)as_\ell(y))$ und $q := s_r(s_r(x)as_\ell(y))$. Da $\text{pcran}(s_r(x)as_\ell(y))H = s_r(x)as_\ell(y)H \subseteq s_r(x)H$, gilt $\text{pcran}(s_r(x)as_\ell(y))s_r(x) = \text{pcran}(s_r(x)as_\ell(y))$, $ps_r(x) = p$, $p \leq s_r(x)$; analog gilt wegen $\text{pcran}(s_\ell(y)a^*s_r(x))H \subseteq s_\ell(y)H$: $\text{pcran}(s_\ell(y)a^*s_r(x))s_\ell(y) = \text{pcran}(s_\ell(y)a^*s_r(x))$, $s_r(s_r(x)as_\ell(y))s_\ell(y) = s_r(s_r(x)as_\ell(y))$, $qs_\ell(y) = q$, $q \leq s_\ell(y)$. Und wegen Gleichung (3.20) in 3.7.48 gilt auch $p \sim q$.

(b) $\Rightarrow$ (a): Gelte (b), also $xAy = \{0\}$. Da $\text{ran } c(y) = \text{cl span}((Ay)H)$, also $xc(y) = 0$. Das heißt, $-xc(y) = x - x$, $x(\text{Id}_H - c(y)) = x$, $xc(y)^\perp = x$. Nach Definition von $c(x)$ somit $c(x) \leq c(y)^\perp$, das heißt, $c(x)c(y)^\perp = c(x)$, $c(x)(\text{Id}_H - c(y)) = c(x)$, $c(x)c(y) = 0$. $\square$

3.7.54 Faktor II. Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Dann gilt:

(a) A Faktor $\Rightarrow A' \cap \text{Proj}(A) = \{0, \text{Id}_H\}$.

(b) Im Fall von $\mathbb{K} = \mathbb{C}$ ist die hinreichende Bedingung in (a) auch notwendig, das heißt, A ist dann genau dann ein Faktor, wenn $A' \cap \text{Proj}(A) = \{0, \text{Id}_H\}$ gilt.

Beweis. (a) Sei A ein Faktor und $p \in A' \cap \text{Proj}(A)^\times$. Da $p \in A' \cap A$, gilt $p = \lambda \cdot \text{Id}_H$ für ein $\lambda \in \mathbb{K}$. Da $p \in \text{OProj}(H)$, gilt $\|p\| = 1$. Also $|\lambda| = 1$, $\lambda = \exp(ti)$ für ein $t \in \mathbb{R}$. Wegen $p^2 = p$ somit $\exp(2ti) = \exp(ti)$, $\exp(2ti)/\exp(ti) = 1$, $\exp(ti) = 1$, $\lambda = 1$ und $p = \text{Id}_H$.

(b) Da nach Bemerkung 3.7.29(a) $A' \cap A$ eine von Neumann-Algebra ist, gilt für $\mathbb{K} = \mathbb{C}$ nach Satz 3.7.45(b) die Gleichungskette $A' \cap A = \text{cl span Proj}(A' \cap A) = \text{cl span}(A' \cap \text{Proj}(A)) = \text{cl span} \{\text{Id}_H\} = \mathbb{C} \cdot \text{Id}_H$. $\square$

3.7.55 Vergleichssatz (MATHIEU [208](1998), Seite 357). Sei A eine von Neumann-Algebra über $\mathbb{K}$. Für alle $p, q \in \text{Proj}(A)$ existiert eine zentrale Projektion $z \in A' \cap \text{Proj}(A)$ mit $pz \preceq qz$ und $p(1-z) \preceq q(1-z)$.

Ist A ein Faktor, so tritt für alle $p, q \in \text{Proj}(A)$ genau einer der drei Fälle $p \sim q$, $p \not\sim q$ oder $q \not\sim p$ ein, das heißt, $\preceq$ ist auf $\text{Proj}(A)/\sim$ eine Totalordnung.

Im Fall $\mathbb{K} = \mathbb{C}$ ist auch umgekehrt A notwendig ein Faktor, wenn $\preceq$ auf $\text{Proj}(A)/\sim$ eine Totalordnung ist.

3.7.56 Klassifizierung I. Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Sei $p \in \text{Proj}(A)$. Die folgenden drei Aussagen sind offensichtlich äquivalent:

(e₁) $p \preceq p$ gilt nur mittels den kanonischen Ausdrücken $p \sim p$, $p \leq p$.

(e₂) $\forall q \in \text{Proj}(A) : ((p \sim q \text{ und } q \leq p) \Rightarrow p = q)$.

(e₃) $\nexists q \in \text{Proj}(A) : (p \sim q \text{ und } q < p)$.

Falls eine (und damit alle drei) der vorstehenden Aussagen zutrifft, heißt p *endlich* (im Englischen *finite*), andernfalls *unendlich* (im Englischen *infinite*).

Die folgenden beiden Aussagen sind äquivalent:

(ru₁) $\nexists q \in \text{Proj}(A)^\times$, q endlich : $q \leq p$.

(ru₂) $\forall q \in \text{Proj}(A)^\times : (q \leq p \Rightarrow q \text{ unendlich})$.

Falls eine (und damit beide) der vorstehenden Aussagen zutrifft, heißt p *rein unendlich* (im Englischen *purely infinite*).

Die folgenden beiden Aussagen sind offensichtlich äquivalent:

(eu₁) $\forall z \in A' \cap \text{Proj}(A)$, $pz \neq 0 : pz$ unendlich.

(eu₂) $\forall z \in A' \cap \text{Proj}(A) : (pz \text{ endlich} \Rightarrow pz = 0)$.

Falls eine (und damit beide) der vorstehenden Aussagen zutrifft, heißt p *echt unendlich* (im Englischen *properly infinite*).

Da wegen $(A_p)' = (A')_p$ (siehe 3.7.20(a)) für jedes $z \in A' \cap \text{Proj}(A)$ die Aussage $pzp \in (A_p)' \cap \text{Proj}(A_p)$ gilt und man umgekehrt ohne Einschränkung von einem $z \in A$ mit $pzp \in \text{Proj}(A_p)$ annehmen kann, dass $z \in \text{Proj}(A)$ ist (indem man mit der Schreibweise von 3.4.22 etwa $z_{1-p} = 0$ setzt), ist p genau dann echt unendlich, falls gilt:

(eu₃) $\nexists w \in A'_p \cap \text{Proj}(A_p)^\times : w$ in A_p endlich.

A heißt *endlich* (bzw. *unendlich*, bzw. *echt unendlich*, bzw. *rein unendlich*), falls Id_H endlich (bzw. unendlich, bzw. echt unendlich, bzw. rein unendlich) ist. Falls A rein unendlich ist, heißt A auch *vom Typ III*. Die folgenden vier Aussagen sind äquivalent:

(eu₄) A ist echt unendlich.

(eu₅) $\nexists z \in A' \cap \text{Proj}(A)^\times : z$ endlich.

(eu₆) Es gibt eine Folge $(p_n)_{n \in \mathbb{N}}$ von paarweise orthogonalen Projektionen von A mit $\sum_{n \in \mathbb{N}} p_n = \text{Id}_H$ und $p_n \sim \text{Id}_H \ \forall n \in \mathbb{N}$.

(eu₇) $\exists q \in \text{Proj}(A) : q \sim q^\perp \sim \text{Id}_H$.

Wegen $(A_p)_p = A_p$ ist p genau dann echt unendlich, wenn A_p echt unendlich ist. Falls p (bzw. A) echt unendlich ist, so ist offensichtlich p (bzw. A) unendlich.

p heißt *abelsch* (im Englischen *abelian*), falls A_p kommutativ ist. (Man verwechsle dies nicht mit dem Begriff zentral, der sich ja auf den Begriff der Kommutativität auf dem A zugrunde liegenden Ring bezieht.) In diesem Zusammenhang sei an 3.3.3 erinnert. Da für minimal idempotentes p jede Projektion in A_p gleich 0 oder p ist,

ist p im Fall von $\mathbb{K} = \mathbb{C}$ wegen Satz 3.7.45 genau dann
minimal idempotent, wenn p schiefminimal idempotent ist. (3.22)

(Zur Information sei hier die folgende Äquivalenz genannt: Für jede M -eingebettete C^* -Algebra A über $\mathbb{C}$ ist ein $q \in \text{Proj}(A)^\times$ genau dann minimal idempotent, wenn q schiefminimal idempotent ist, siehe HARMAND, WERNER und WERNER [126](1993), Seite 121.) Offensichtlich ist im Fall von $\mathbb{K} = \mathbb{C}$ jede schiefminimal idempotente Projektion abelsch. Wegen 3.7.20(b) ist in einem Faktor jede abelsche Projektion schiefminimal.

p heißt *treu* (im Englischen *faithful*), wenn $c(p) = \text{Id}_H$ gilt. Die folgenden drei Aussagen sind äquivalent:

(d₁) $\forall r \in \text{Proj}(A)^\times \ \exists q \in \text{Proj}(A)^\times, q$ abelsch : $q \leq r$.

(d₂) $\forall z \in A' \cap \text{Proj}(A)^\times \ \exists q \in \text{Proj}(A)^\times, q$ abelsch : $q \leq z$.

(d₃) $\exists q \in \text{Proj}(A) : q$ abelsch und treu.

Beweis. Siehe STRĂTILĂ und ZSIDÓ [270](1979), Seite 101. Oder auch PEDERSEN [233](1979), Seite 172, Satz 5.5.3, wobei man im Fall von $\mathbb{K} = \mathbb{R}$ anstelle seines Satzes 5.4.7 die hierige Bemerkung 3.7.53 verwende. □

Gilt eine (und damit alle drei) der vorstehenden Äquivalenzen, so heißt A *vom Typ I* oder auch *diskret* (im Englischen *discrete*). p heißt *diskret*, falls die Kompression A_p diskret ist.

Die folgenden beiden Aussagen sind äquivalent:

(he₁) $\nexists z \in A' \cap \text{Proj}(A)^\times : z$ rein unendlich.

(he₂) $\forall z \in A' \cap \text{Proj}(A)^\times \ \exists q \in \text{Proj}(A)^\times : q \leq z$.

Zumindest für den Fall $\mathbb{K} = \mathbb{C}$ ist die folgende Aussage äquivalent zu jeder der beiden vorstehenden Aussagen:

(he₃) $\exists q \in \text{Proj}(A) : q$ treu.

Falls die Aussage (he₁) oder (he₂) (und damit beide) gilt, heißt A *halbendlich* (im Englischen *semi-finite*).

3.7.57 (LI [197](1992), Seite 268; PEDERSEN [233](1979), Seite 166; LI [198] (2003), Seite 168).

Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Dann gibt es eine maximale, zentrale, endliche Projektion z_1 und eine dazu orthogonale, maximale, zentrale, rein unendliche Projektion z_3 . Somit ist $z_2 := \text{Id}_H - z_1 - z_3$ eine zentrale, halbendliche, echt unendliche Projektion und es existiert die Peirce-Zerlegung von A bezüglich $\{z_1, z_3\}$:

$$A = Az_1 \dot{+} Az_2 \dot{+} Az_3.$$

Diese Zerlegung ist im folgenden Sinne eindeutig: Ist $A = Ap_1 \dot{+} Ap_2 \dot{+} Ap_3$ eine Zerlegung von A mit Ap_1 endlich, Ap_3 vom Typ III und Ap_2 halbendlich und echt unendlich mit $p_1, p_2, p_3 \in A' \cap \text{Proj}(A)$, $p_1 + p_2 + p_3 = \text{Id}_H$, so ist $p_i = z_i$, $i = 1, 2, 3$.

3.7.58 Klassifizierung II. Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Sei $p \in \text{Proj}(A)$. p heißt *halbabelsch* (im Englischen *semi-abelian*), falls $pA_{(*)}p$ kommutativ ist. Offensichtlich ist p halbabelsch, falls p abelsch ist; im Fall von $\mathbb{K} = \mathbb{C}$ gilt auch die Umkehrung. Nach Lemma 3.2.5 ist jedes Atom (siehe Definition 3.5.57) von A halbabelsch. Jede halbabelsche Projektion ist endlich. Die folgenden fünf Aussagen sind äquivalent:

(ha₁) p ist halbabelsch.

(ha₂) $pA_{(*)}p \left(\stackrel{(3.1)}{=} (pAp)_{(*)} \right)$ ist gleich $(A' \cap A)_{(*)}p \left(= (A'_p \cap A_p)_{(*)} \right)$.

(ha₃) $\forall z \in A' \cap \text{Proj}(A)^\times, z \leq p \quad \nexists q_1, q_2 \in \text{Proj}(A)^\times, q_1 q_2 = 0, q_1 \sim q_2 : z = q_1 + q_2$.

(ha₄) $\nexists q \in \text{Proj}(A_p)^\times : q$ nicht zentral.

(ha₅) $\forall q \in \text{Proj}(A_p)^\times : q$ zentral.

A heißt *halbdiskret* (im Englischen *semi-discrete*), falls gilt:

(hd) $\forall z \in A' \cap \text{Proj}(A)^\times \quad \exists q \in \text{Proj}(A)^\times, q$ halbabelsch : $q \leq z$.

A heißt *halbkontinuierlich* (im Englischen *semi-continuous*), falls gilt:

(hk) $\nexists q \in \text{Proj}(A)^\times : q$ abelsch.

p heißt *halbkontinuierlich*, falls A_p halbkontinuierlich ist.

A heißt *kontinuierlich* (im Englischen *continuous*), falls gilt:

(k) $\nexists q \in \text{Proj}(A)^\times : q$ halbabelsch.

p heißt *kontinuierlich*, falls A_p kontinuierlich ist.

A heißt *vom Typ II*, falls A halbdiskret und kontinuierlich ist. Insbesondere weist Typ II keine Atome auf. Da jeder Typ III kontinuierlich ist (siehe LI [197](1992), Seite 306, Satz 6.8.1, und LI [198](2003), Seite 186, Satz 8.6.4(1)), weist auch Typ III keine Atome auf.

Offensichtlich fallen im Fall von $\mathbb{K} = \mathbb{C}$ jeweils die Begriffe halbdiskret und diskret, halbkontinuierlich und kontinuierlich zusammen.

3.7.59 Zerlegung im Fall von $\mathbb{K} = \mathbb{C}$ (STRÄTILÄ und ZSIDÓ [270](1979), Seite 100). Sei H ein Hilbertraum über $\mathbb{C}$ und A eine von Neumann-Algebra auf H über $\mathbb{C}$.

(a) Definiere die beiden folgenden zentralen Projektionen aus A :

$$p_{\text{he}} := \sup \{q \in A' \cap \text{Proj}(A) : A_q \text{ halbdiskret}\} \quad \text{und} \quad p_{\text{ru}} := p_{\text{he}}^\perp.$$

Dann hat man die Peirce-Zerlegung von A bezüglich p_{he} : $A = A_{\text{he}} \dot{+} A_{\text{ru}}$, wobei $A_{\text{he}} := Ap_{\text{he}}$ halbdiskret und $A_{\text{ru}} := Ap_{\text{ru}}$ rein unendlich (also vom Typ III) ist.

(b) Definiere die beiden folgenden zentralen Projektionen aus A :

$$p_{\text{e}} := \sup \{q \in A' \cap \text{Proj}(A) : q \text{ endlich}\} \quad \text{und} \quad p_{\text{eu}} := p_{\text{e}}^\perp.$$

Dann hat man die Peirce-Zerlegung von A bezüglich p_{e} : $A = A_{\text{e}} \dot{+} A_{\text{eu}}$, wobei $A_{\text{e}} := Ap_{\text{e}}$ endlich und $A_{\text{eu}} := Ap_{\text{eu}}$ echt unendlich ist.

(c) Definiere die beiden folgenden zentralen Projektionen aus A :

$$p_{\text{d}} := \sup \{q \in A' \cap \text{Proj}(A) : q \text{ diskret}\} \quad \text{und} \quad p_{\text{k}} := p_{\text{d}}^\perp.$$

Dann hat man die Peirce-Zerlegung von A bezüglich p_{d} : $A = A_{\text{d}} \dot{+} A_{\text{k}}$, wobei $A_{\text{d}} := Ap_{\text{d}}$ diskret (also vom Typ I) und $A_{\text{k}} := Ap_{\text{k}}$ vom Typ II oder vom Typ III ist. Man bemerke, dass, falls A eine schiefminimale Projektion enthält, diese in A_{d} enthalten sein muss.

(d) Setze:

$$\begin{aligned} p_1 &:= p_{\text{d}} p_{\text{e}} & (= p_{\text{d}} p_{\text{e}} p_{\text{he}}), \\ p_2 &:= p_{\text{d}} p_{\text{eu}} & (= p_{\text{d}} p_{\text{eu}} p_{\text{he}}), \\ p_3 &:= p_{\text{k}} p_{\text{e}} & (= p_{\text{k}} p_{\text{e}} p_{\text{he}}), \\ p_4 &:= p_{\text{k}} p_{\text{eu}} & (= p_{\text{k}} p_{\text{eu}} p_{\text{he}}), \\ p_5 &:= p_{\text{ru}} & (= p_{\text{ru}} p_{\text{eu}}), \\ p_{\text{I}} &:= p_{\text{d}}, \quad p_{\text{II}} := p_{\text{k}} p_{\text{he}}, \quad p_{\text{III}} := p_5. \end{aligned}$$

Offensichtlich sind die Projektionen $p_1, \dots, p_5$ paarweise orthogonal; ebenso sind p_I, p_{II}, p_{III} paarweise orthogonal.

- (e) A heißt vom Typ I_{fin} , falls A vom Typ I und endlich ist.
- A heißt vom Typ I_∞ , falls A vom Typ I und unendlich ist.
- A heißt vom Typ II_1 , falls A vom Typ II und endlich ist.
- A heißt vom Typ II_∞ , falls A vom Typ II und echt unendlich ist.
- (f) A lässt sich eindeutig zerlegen als

$$A = A_I \dot{+} A_{II} \dot{+} A_{III},$$

wobei $A_\gamma := Ap_\gamma$ vom Typ γ , $\gamma = I, II, III$, ist.

- (g) A lässt sich eindeutig zerlegen als

$$A = A_{I_{fin}} \dot{+} A_{I_\infty} \dot{+} A_{II_1} \dot{+} A_{II_\infty} \dot{+} A_{III},$$

wobei $A_{I_{fin}} := Ap_1$ vom Typ I_{fin} , $A_{I_\infty} := Ap_2$ vom Typ I_∞ , $A_{II_1} := Ap_3$ vom Typ II_1 , $A_{II_\infty} := Ap_4$ vom Typ II_∞ und $A_{III} := Ap_5$ vom Typ III ist.

- (h) Falls A ein Faktor ist, so ist A genau von einem der Typen $I_{fin}, I_\infty, II_1, II_\infty$ oder III.

(i) (TAKESAKI [274](2001), Seiten 155 und 298). Falls A kommutativ ist, so ist A vom Typ I. Jede Teilprojektion einer endlichen Projektion ist endlich. Erfüllt A die Eigenschaft

$$\forall q \in \text{Proj}(A)^\times \quad \exists r \in \text{Proj}(A)^\times, \text{ } r \text{ schiefminimal} : r \leq q,$$

so ist A vom Typ I. (TAKESAKI, ebd., nennt A , falls es die besagte Eigenschaft besitzt, *atomic* (im Englischen *atomic*).)

- (j) (CONNES [51](1994), Seite 454). In der Theorie der von Neumann-Algebren auf einem separablen Hilbertraum über $\mathbb{C}$ gilt:

1. Es existieren unendlichdimensionale Faktoren, die keine schiefminimalen Projektionen haben.

2. Ein Faktor mit einem separablen Prädual ist genau dann vom Typ I, wenn er über eine schiefminimale Projektion verfügt.

Siehe auch KADISON und RINGROSE [167](1986), Seite 424, und DIXMIER [76](1981), Seite 140, Korollar 3.

(k) In CARADUS, PFAFFENBERGER und YOOD [42](1974), Seite 121, findet man die folgende Definition: Sei A ein Faktor über $\mathbb{C}$. Dann heißt A vom Typ I, wenn A mindestens eine schiefminimal idempotente Projektion enthält; vom Typ II, wenn A keine schiefminimal idempotenten Projektionen besitzt und zwei Projektionen $p, q \in A$ existieren, so dass $\text{ran}(p)$ von keiner partiellen Isometrie $U \in A$ isometrisch auf $\text{ran}(q)$ abgebildet wird; und vom Typ III sonst.

3.7.60 Zerlegung im Fall von $\mathbb{K} = \mathbb{R}$ (LI [198](2003), Seite 170). Sei H ein Hilbertraum über $\mathbb{R}$ und A eine von Neumann-Algebra auf H über $\mathbb{R}$.

- (a) A heißt vom Typ H , wenn A halbdiskret und halbkontinuierlich ist. Zwecks einer einfacheren Schreibweise setze $\Gamma := \{ I, H, II, III \}$. A lässt sich eindeutig zerlegen als

$$A = A_I \dot{+} A_H \dot{+} A_{II} \dot{+} A_{III},$$

wobei $A_\gamma := Ap_\gamma$, $p_\gamma \in A' \cap \text{Proj}(A)$, $\gamma \in \Gamma$, mit $\sum_{\gamma \in \Gamma} p_\gamma = \text{Id}_H$ und A_γ vom Typ γ , $\gamma \in \Gamma$, ist; p_I ist wie in 3.7.59 definiert; p_{III} ist wie in 3.7.59 die maximale, zentrale, rein unendliche Projektion von A ; p_{II} ergibt sich in anderer Weise als in 3.7.59 definiert. Des Weiteren ist $\widehat{A}_I := A_I \dot{+} A_H$ halbdiskret, $\widehat{A}_{II} := A_H \dot{+} A_{II}$ halbbendlich und halbkontinuierlich,

$\widehat{A_{II}} \dot{+} A_{III}$ halbkontinuierlich und $A_{II} \dot{+} A_{III}$ kontinuierlich. Insbesondere weist $A_{II} \dot{+} A_{III}$ keine Atome auf.

(b) Falls A ein Faktor ist, so ist A genau von einem der Typen γ , $\gamma \in \Gamma$. Bei $\gamma = I$ ist A von der Form $\mathcal{L}(H_n)$, wobei H_n ein Hilbertraum über $\mathbb{R}$ der Dimension n für ein $n \in \mathbb{N}^\times \cup \{\infty\}$ ist; dabei ist der Faktor genau dann endlich, wenn $n \in \mathbb{N}^\times$ ist, und genau dann echt unendlich, wenn $n = \infty$ ist.

Kapitel 4

Holomorphie und Tripelsysteme

4.1 Analytische und differenzierbare Abbildungen

4.1.1 Potenzreihe (UPMEIER [279](1985), Seite 4). Seien X und Y Banachräume über $\mathbb{K}$. Sei $(p_n)_{n \in \mathbb{N}}$ eine Folge von stetigen n -homogenen Polynomen $p_n \in \mathcal{P}^n(X, Y)$, $n \in \mathbb{N}$. Dann heißt $\sum_{n=0}^{\infty} p_n$ eine *Potenzreihe* von X nach Y ; ihr *Konvergenzradius* ist definiert als die größte erweiterte reelle Zahl $R \leq \infty$, mit der für jedes reelle r mit $0 \leq r < R$ die Reihe $\sum_{n=0}^{\infty} p_n$ auf $r \cdot B_X$ gleichmäßig konvergiert. Die Potenzreihe $\sum_{n=0}^{\infty} p_n$ heißt konvergent, wenn $R > 0$ ist.

4.1.2 Analytische Abbildung (UPMEIER [279](1985), Seiten 6 ff). Seien X und Y Banachräume über $\mathbb{K}$ und sei $\Omega \subseteq X$ offen. Eine Abbildung $f: \Omega \rightarrow Y$ heißt *analytisch*, wenn für jedes $a \in \Omega$ eine konvergente Potenzreihe $\sum_{n=0}^{\infty} p_n$, $p_n \in \mathcal{P}^n(X, Y)$, $n \in \mathbb{N}$, existiert mit $f(x) = \sum_{n=0}^{\infty} p_n(x - a)$ für alle x einer Umgebung von a ; um den Grundkörper $\mathbb{K}$ zu betonen, verwendet man die Begriffe *reell-* beziehungsweise *komplex-analytisch*.

Ein *Gebiet* (im Englischen *domain*) eines Banachraumes Z über $\mathbb{K}$ ist eine zusammenhängende offene Teilmenge des Banachraumes Z .

Sei $\Omega \subseteq X$ ein Gebiet in X . Eine analytische Abbildung $T: \Omega \rightarrow Y$ heißt *bianalytisch*, wenn $\text{ran}(T)$ ein Gebiet in Y ist, T injektiv ist und $T^{-1}: \text{ran}(T) \rightarrow X$ analytisch ist; in diesem Fall heißen Ω und $T(\Omega)$ *bianalytisch äquivalent*. (UPMEIER [280](1987)], Seite 10.)

4.1.3 Beispiel (UPMEIER [279](1985), Seiten 27 und 28). Sei X eine unitale Banachalgebra über $\mathbb{K}$. Die Gruppe $G(X)$ aller invertierbaren Elemente von X ist offen und die Abbildung $x \mapsto x^{-1}$ auf $G(X)$ ist analytisch. Für jedes $x \in X$ konvergiert die Reihe $\exp(x) := e + \sum_{n=1}^{\infty} \frac{x^n}{n!}$ absolut. Für kommutierende $x, y \in X$ gilt $\exp(x + y) = \exp(x) \exp(y)$, womit insbesondere für jedes $x \in X$ $\exp(x)$ invertierbar ist: $\exp(x)^{-1} = \exp(-x)$. Die somit definierte Abbildung $\exp: X \rightarrow G(X)$ ist analytisch.

4.1.4 Gradient einer Abbildung (BOURGIN [36](1983), Seite 117). Seien X und Y normierte Vektorräume über $\mathbb{K}$ und sei $R \subseteq X$ beschränkt, $\Omega \subseteq X$ offen, $f: \Omega \rightarrow Y$ eine Abbildung und $a \in \Omega$. f heißt *R -subdifferenzierbar bei a* , wenn eine Abbildung $T: R \rightarrow Y$ existiert mit

$$\lim_{h \rightarrow 0, h > 0} \sup_{v \in R} \left\| \frac{1}{h} (f(a + hv) - f(a)) - T(v) \right\| = 0.$$

Falls solch ein T existiert, heißt es der *R -Subgradient* von f bei a , in Zeichen $\tilde{\nabla}_R f(a)$; und falls T — eventuell nach einer Fortsetzung von R auf X — als ein Element von $\mathcal{L}(X, Y)$ aufgefasst werden kann, heißt f *R -differenzierbar* bei a und T der *R -Gradient* von f bei a , in Zeichen $\nabla_R f(a)$.

Ist die Menge R einelementig, etwa $R = \{r\}$, so schreibt man $\nabla_r f(a)$ anstelle von $\nabla_R f(a)$; $\tilde{\nabla}_R f(a)$ entsprechend.

Existiert für alle $r \in X$ der $\{r\}$ -Subgradient von f bei a , so heißt f *subdifferenzierbar* bei a . Ist f bei jedem $a \in \Omega$ subdifferenzierbar, so heißt f subdifferenzierbar.

Existiert der B_X -Subgradient von f bei a , so heißt f *stark subdifferenzierbar* bei a (im Englischen *strongly subdifferentiable*). Ist f bei jedem $a \in \Omega$ stark subdifferenzierbar, so heißt f *stark subdifferenzierbar*.

4.1.5 Fréchet-Ableitung. Seien X und Y normierte Vektorräume über $\mathbb{K}$ und sei $\Omega \subseteq X$ offen und $f: \Omega \rightarrow Y$ eine Abbildung. Für das Folgende bemerke man, dass für endlichdimensionales X die beiden Mengen $L(X, Y)$ und $\mathcal{L}(X, Y)$ übereinstimmen (MEGGINSON [211](1998), Seite 29, Theorem 1.4.12). f heißt *Fréchet-differenzierbar* (oder auch einfach nur *differenzierbar*) (über $\mathbb{K}$) bei $a \in \Omega$, falls eine der folgenden äquivalenten Bedingungen erfüllt ist:

- (a) $\lim_{x \rightarrow a, x \neq a} \frac{\|f(x) - f(a) - T(x - a)\|}{\|x - a\|} = 0$ für ein $T \in \mathcal{L}(X, Y)$.
 (b) $\exists T \in \mathcal{L}(X, Y) \quad \forall \varepsilon > 0 \quad \exists \delta > 0 \quad \forall x \in \Omega, 0 < \|x - a\| < \delta :$

$$\frac{\|f(x) - f(a) - T(x - a)\|}{\|x - a\|} < \varepsilon.$$

(An dieser Stelle sei darauf aufmerksam gemacht, dass GILES [109](1982), Seite 141, in seiner Definition von Fréchet-Differenzierbarkeit von f nur $T \in L(X, Y)$ fordert! Für $\mathbb{K} = \mathbb{R}$ zeigt er, dass dann sein T genau dann stetig bei a ist, wenn f bei a stetig ist.)

- (c) $\lim_{x \rightarrow 0, x \neq 0} \frac{1}{\|x\|} (\|f(a + x) - f(a) - T(x)\|) = 0$ für ein $T \in \mathcal{L}(X, Y)$.

(d) f ist bei a stetig und $\lim_{x \rightarrow a, x \neq a} \frac{\|f(x) - f(a) - T(x - a)\|}{\|x - a\|} = 0$ für ein $T \in L(X, Y)$. In der Terminologie von Dieudonné [68] ist das die folgende Aussage: f ist bei a stetig und es existiert ein $T \in L(X, Y)$, so dass sich die beiden Abbildungen f und $x \mapsto f(a) + T(x - a)$ in a berühren.

(e) $\lim_{h \rightarrow 0, h \in \mathbb{K}^\times} \frac{1}{h} (f(a + hv) - f(a)) = T(v)$ (Normkonvergenz) für alle $v \in X$ und ein $T \in \mathcal{L}(X, Y)$, wobei die Konvergenz gleichmäßig bezüglich $v \in B_X$ (oder bei $X \neq \{0\}$ äquivalent: bezüglich $v \in S_X$) ist.

(f) Es existiert $T := \nabla_{B_X} f(a)$, also der B_X -Gradient (oder bei $X \neq \{0\}$ äquivalent: der S_X -Gradient) von f bei a .

- (g) $\exists T \in \mathcal{L}(X, Y) \quad \forall \varepsilon > 0 \quad \exists h_0 \in \mathbb{R}^\times \quad \forall h \in \mathbb{K}, 0 < |h| < h_0 \quad \forall v \in B_X :$

$$\left\| \frac{1}{h} (f(a + hv) - f(a)) - T(v) \right\| < \varepsilon.$$

Beweis. (f) $\Rightarrow$ (g): Klar im Fall von $\mathbb{K} = \mathbb{R}$. Liege also der Fall $\mathbb{K} = \mathbb{C}$ vor und gelte (f). Sei $\varepsilon > 0$. Dann existiert ein $h_0 > 0$, so dass für alle h mit $0 < h < h_0$ gilt: $\sup_{v \in B_X} \left\| \frac{1}{h} (f(a + hv) - f(a)) - T(v) \right\| < \varepsilon$. Sei nun $h \in \mathbb{C}$ mit $0 < |h| < h_0$. Setze $A := \sup_{v \in B_X} \left\| \frac{1}{h} (f(a + hv) - f(a)) - T(v) \right\|$. Dann ist $A < \varepsilon$ zu zeigen. Es ist $h = r \exp(i\theta)$ für ein $r \in \mathbb{R}$ und ein $\theta \in \mathbb{R}$. Es gilt $A = \sup_{v \in B_X} \left\| \frac{1}{h} (f(a + h \exp(-i\theta)v) - f(a)) - T(\exp(-i\theta)v) \right\| = \sup_{v \in B_X} \left\| \exp(-i\theta) \left(\frac{1}{\exp(-i\theta)h} (f(a + (\exp(-i\theta)h)v) - f(a)) - T(v) \right) \right\| < \varepsilon$. Man bemerke, dass hierbei entscheidend einging, dass B_X mit ihrer kreisförmigen Hülle übereinstimmt.

(g) $\Rightarrow$ (f): Klar.

Bezüglich (g) $\Rightarrow$ (e) bemerke man nur: Gelte (g) und sei $v \in X$ mit $\|v\| > 1$. Dann gilt $T(v) = \|v\| \cdot T(v/\|v\|) = \|v\| \cdot \lim_{h \rightarrow 0, h \in \mathbb{K}^\times} \frac{f(a+hv/\|v\|) - f(a)}{h} = \lim_{h \rightarrow 0, h \in \mathbb{K}^\times} \frac{f(a + \frac{h}{\|v\|}v) - f(a)}{\frac{h}{\|v\|}} = \lim_{h \rightarrow 0, h \in \mathbb{K}^\times} \frac{f(a + hv) - f(a)}{h}$. $\square$

Anstelle von T schreibt man $Df(a)$ oder auch $f'(a)$ und nennt $Df(a)$ die *Fréchet-Ableitung* (oder einfach auch nur die *Ableitung*) von f bei a .

Wenn f bei jedem Punkt von Ω differenzierbar ist, dann heißt f *Fréchet-differenzierbar* (oder auch einfach nur *differenzierbar*) und die (*Fréchet-*)*Ableitung* von f ist die Abbildung $Df: \Omega \rightarrow \mathcal{L}(X, Y), a \mapsto Df(a)$; im Fall von $\mathbb{K} = \mathbb{C}$ heißt f dann *holomorph*. Die Abbildung f heißt *zweimal differenzierbar* bei $a \in \Omega$, wenn Df bei $a \in \Omega$ differenzierbar ist. In diesem Fall ist die zweite Ableitung $D^2f \in \mathcal{L}(X, \mathcal{L}(X, Y)) \cong \mathcal{L}^2(X, Y)$ definiert per $D^2f(a) := (D(Df))(a)$ für alle $a \in X$. Unter D^0f soll f verstanden werden.

Sei U ein Unterraum von X . Falls f in $a \in \Omega$ differenzierbar ist, dann ist die Abbildung $f|_{\Omega \cap (a + U)}$ bei a differenzierbar und ihre Ableitung $Df(a)|_U$ bei a heißt die *partielle Fréchet-Ableitung* (oder einfach nur *partielle Ableitung*) von f bei a bezüglich U und wird mit $D_U f(a)$ bezeichnet. Für Weiteres zur partiellen Ableitung siehe PFLAUMANN und UNGER [237](1974), Seiten 71-79, und GERHARDT [107](2006), Seiten 73-94.

4.1.6 Gâteaux- und Richtungsableitung (JEGGLE [156](1979), Kapitel 1.2; PACHALE [229](1974), Kapitel 2.3; ZEIDLER [295](1986), Kapitel 4; NACHBIN [219](1981), 2. Teil).

Seien X und Y normierte Vektorräume über $\mathbb{K}$. Sei $\Omega \subseteq X$ offen, $f: \Omega \rightarrow Y$ eine Abbildung, $a \in \Omega$ und $v \in X$. Da a ein innerer Punkt von Ω ist, existiert ein $r > 0$, so dass man mit $W := r \cdot \text{int}(B_{\mathbb{K}})$ eine sternförmige, offene Umgebung in $\mathbb{K}$ von $0 \in \mathbb{K}$ vorliegen hat mit $a + tv \in \Omega$ für alle $t \in W$.

Man definiere die beiden Abbildungen

$$g_{a,v}: W \rightarrow Y, \quad t \mapsto f(a + tv)$$

und

$$q_{a,v}: \mathbb{R} \cap W \rightarrow Y_{\mathbb{R}}, \quad t \mapsto f(a + tv).$$

Die Abbildung f heißt *Gâteaux-differenzierbar* (über $\mathbb{K}$) *bei a in Richtung v* , falls die Abbildung $q_{a,v}$ bei $0 \in \mathbb{R}$ differenzierbar über $\mathbb{R}$ ist. Das in Y liegende Element

$$\delta f(a)(v) := Dq_{a,v}(0)(1)$$

heißt die *Gâteaux-Variation von f bei a bezüglich v* . Falls $\delta f(a)(v)$ für ein $v \neq 0$ existiert, so existiert für jedes $\lambda \in \mathbb{R}^\times$ auch $\delta f(a)(\lambda v)$, wobei dann $\lambda \cdot \delta f(a)(v) = \delta f(a)(\lambda v)$ gilt.

Die Abbildung f ist genau dann Gâteaux-differenzierbar bei a in Richtung v , wenn sie $\{-v, v\}$ -differenzierbar ist; in diesem Fall gilt dann

$$\delta f(a)(v) = \nabla_{S_{\mathbb{R}} \cdot v} f(a)(v).$$

Beweis. Sei f Gâteaux-differenzierbar bei a in Richtung v . Sei dementsprechend W eine sternförmige, offene Umgebung in $\mathbb{R}$ von 0 mit $a + tv \in \Omega$ für alle $t \in W$ und $q: W \rightarrow Y_{\mathbb{R}}, t \mapsto f(a + tv)$, differenzierbar bei 0 . Es existiert also ein $S \in \mathcal{L}(\mathbb{R}, Y_{\mathbb{R}})$ mit: $\forall \varepsilon > 0 \exists h_0 \in \mathbb{R}^\times \forall h \in \mathbb{R}, 0 < |h| < h_0 \forall \lambda \in \{-1, 1\} : \|\frac{1}{h}(q(h\lambda) - q(0)) - S(\lambda)\| < \varepsilon$, also $\|\frac{1}{h}(f(a + h\lambda v) - f(a)) - S(\lambda)\| < \varepsilon$. Definiere ein $T \in \mathcal{L}(X, Y)$ per $T(v) := S(1)$ und $Tx = 0$ für alle $x \notin \mathbb{K}v$. Damit gilt also $\sup_{w \in \{-v, v\}} \|\frac{1}{h}(f(a + hw) - f(a)) - T(w)\| < \varepsilon$, das heißt, $T = \nabla_{\{-v, v\}} f(a)$. Da $\delta f(a)(v) = Dq(0)(1) = S(1) = T(v)$ gilt, folgt die behauptete Gleichung.

Gilt umgekehrt $T = \nabla_{\{-v, v\}} f(a)$, so erhält man wieder das ebige $S \in \mathcal{L}(\mathbb{R}, Y_{\mathbb{R}})$ per $S(1) := T(v)$. $\square$

Existiert für jedes $v \in X$ die Gâteaux-Variation von f bei a bezüglich v und ist die dadurch definierte Abbildung

$$\delta f(a): X \rightarrow Y, \quad v \mapsto \delta f(a)(v)$$

linear und stetig, so wird diese Abbildung die *Gâteaux-Ableitung von f bei a* genannt, mit $Gf(a)$ bezeichnet und sagt in diesem Fall, dass f *Gâteaux-differenzierbar* (über $\mathbb{K}$) *bei a* sei. Die Abbildung $f: \Omega \rightarrow Y$ ist also genau dann Gâteaux-differenzierbar bei $a \in \Omega$, falls eine der folgenden äquivalenten Bedingungen erfüllt ist:

- (a) $\lim_{h \rightarrow 0, h \in \mathbb{R}^\times} \frac{1}{h} (f(a + hv) - f(a)) = T(v)$ (Normkonvergenz) für alle $v \in X$ und ein $T \in \mathcal{L}(X, Y)$.
- (b) $\exists T \in \mathcal{L}(X, Y) \quad \forall v \in X \quad \forall \varepsilon > 0 \quad \exists h_0 \in \mathbb{R}^\times \quad \forall h \in \mathbb{R}, 0 < |h| < h_0 :$

$$\left\| \frac{1}{h} (f(a + hv) - f(a)) - T(v) \right\| < \varepsilon.$$

Die Abbildung f heißt *komplex Gâteaux-differenzierbar* (über $\mathbb{K}$) *bei a in Richtung v* , falls die Abbildung $g_{a,v}$ bei $0 \in \mathbb{K}$ differenzierbar ist. Das dann in Y liegende Element

$$D_v f(a) := Dg_{a,v}(0)(1)$$

heißt dann die *Richtungsableitung* (über $\mathbb{K}$) *bei a in Richtung v* . Offensichtlich beschreiben im Fall von $\mathbb{K} = \mathbb{R}$ die Begriffe Gâteaux-differenzierbar bei a in Richtung v und komplex Gâteaux-differenzierbar bei a in Richtung v ein und denselben Sachverhalt.

Die Abbildung f ist genau dann komplex Gâteaux-differenzierbar bei a in Richtung v , wenn f $S_{\mathbb{K}} \cdot v$ -differenzierbar ist; in diesem Fall gilt dann

$$D_v f(a) = \nabla_{S_{\mathbb{K}} \cdot v} f(a)(v).$$

Beweis. Der Beweis geht ganz analog wie bei der Gâteaux-Differenzierbarkeit bei a in Richtung v . Sei f komplex Gâteaux-differenzierbar bei a in Richtung v . Sei dementsprechend $W := \rho \cdot \text{int}(B_{\mathbb{K}})$ für ein $\rho > 0$ derart, dass für alle $z \in W$ gilt: $a + zv \in \Omega$ und $g: W \rightarrow Y, z \mapsto f(a + zv)$, ist differenzierbar über $\mathbb{K}$ bei 0 . Es existiert also ein $S \in \mathcal{L}(\mathbb{K}, Y)$ mit: $\forall \varepsilon > 0 \quad \exists h_0 \in \mathbb{R}^\times \quad \forall h \in \mathbb{K}, 0 < |h| < h_0 \quad \forall \lambda \in S_{\mathbb{K}} :$ $\left\| \frac{1}{h} (g(h\lambda) - g(0)) - S(\lambda) \right\| < \varepsilon$, also $\left\| \frac{1}{h} (f(a + h\lambda v) - f(a)) - S(\lambda) \right\| < \varepsilon$. Definiere wieder ein $T \in \mathcal{L}(X, Y)$ per $T(v) := S(1)$ und $Tx = 0$ für alle $x \notin \mathbb{K}v$. Damit gilt also $\sup_{w \in S_{\mathbb{K}} \cdot v} \left\| \frac{1}{h} (f(a + hw) - f(a)) - T(w) \right\| < \varepsilon$, das heißt, $T = \nabla_{S_{\mathbb{K}} \cdot v} f(a)$. Da $D_v f(a) = Dg(0)(1) = S(1) = T(v)$ gilt, folgt die behauptete Gleichung. $\square$

Existiert für jedes $v \in X$ die Richtungsableitung von f bei a in Richtung v und ist die dadurch definierte Abbildung

$$X \rightarrow Y, \quad v \mapsto D_v f(a) \tag{4.1}$$

linear und stetig, so wird diese Abbildung die *komplexe Gâteaux-Ableitung von f bei a* genannt, mit $G_{\mathbb{K}} f(a)$ bezeichnet und sagt in diesem Fall, dass f *komplex Gâteaux-differenzierbar bei a* sei; liegt dabei speziell auch noch $\mathbb{K} = \mathbb{C}$ vor, so heißt f *Gâteaux-holomorph bei a* . (Es sei an dieser Stelle darauf hingewiesen, dass KAUP [180], Seite 5, hierbei lediglich die Existenz der Abbildung (4.1) fordert.)

Falls f Fréchet-differenzierbar bei a ist, gilt

$$Gf(a)(v) = D_v f(a) = Df(a)(v)$$

für alle $v \in X$.

4.1.7. Seien X und Y normierte Vektorräume über $\mathbb{K}$. Sei $\Omega \subseteq X$ offen, $f: \Omega \rightarrow Y$ eine Abbildung und $a \in \Omega$. Dann ist f genau dann Fréchet-differenzierbar bei a , wenn f bei a sowohl Gâteaux-differenzierbar als auch stark subdifferenzierbar ist.

Beweis. Dass f die behaupteten Eigenschaften hat, wenn es Fréchet-differenzierbar bei a ist, ist klar.

Zur Umkehrung: Bezeichne $\tau(v, h)$ den Differenzenquotienten $\frac{f(a+hv)-f(a)}{h}$ für beliebige $v \in X$, $h \in \mathbb{K}^\times$. Sei f bei a sowohl Gâteaux-differenzierbar als auch stark subdifferenzierbar. Aus der Gâteaux-Differenzierbarkeit von f bei a folgt, dass $\tau(v, h)$ für jedes $v \in X$ für $h \rightarrow 0$, $h \in \mathbb{K}^\times$, konvergiert und, dass die Abbildung $T: X \rightarrow Y$, $v \rightarrow \lim_{h \rightarrow 0, h > 0} \tau(v, h)$ linear und stetig ist. Dass f bei a stark subdifferenzierbar ist, heißt, dass es eine Abbildung $S: B_X \rightarrow Y$ gibt mit $\lim_{h \rightarrow 0, h > 0} \sup_{v \in B_X} \|\tau(v, h) - S(v)\| = 0$. Offensichtlich gilt $T \upharpoonright B_X = S$, das heißt, T ist der B_X -Gradient von f bei a . Nach 4.1.5(f) ist f bei a Fréchet-differenzierbar. $\square$

4.1.8 Asplundräume (PHELPS [238](1993)). Ein Banachraum X über $\mathbb{K}$ heißt ein *Asplundraum* (bzw. *schwacher Asplundraum*), wenn für jede offene, konvexe, nicht leere Teilmenge C von X jede stetige, konvexe Funktion $f: C \rightarrow \mathbb{R}$ bei jedem Punkt einer dichten G_δ -Teilmenge von C Fréchet-differenzierbar (bzw. Gâteaux-differenzierbar) ist.

Separable Banachräume über $\mathbb{K}$ sind schwache Asplundräume; für separable Banachräume über $\mathbb{R}$ wurde dies 1933 von Mazur gezeigt.

Der Folgenraum ℓ_∞ ist kein schwacher Asplundraum und der Folgenraum ℓ_1 ist kein Asplundraum.

Falls der Dualraum X^* eines Banachraumes X über $\mathbb{K}$ separabel ist, so ist X ein Asplundraum; für den Fall $\mathbb{K} = \mathbb{R}$ wurde dies 1968 von Asplund gezeigt.

Da ein reflexiver Banachraum X genau dann separabel ist, wenn sein Dualraum X^* separabel ist, folgt das Theorem: Ist X ein separabler, reflexiver Banachraum über $\mathbb{K}$, so ist X ein Asplundraum; für den Fall $\mathbb{K} = \mathbb{R}$ wurde dieses Theorem 1963 von J. Lindenstrauss gezeigt.

(ebd., Seite 82). Ein Banachraum über $\mathbb{K}$ ist genau dann ein Asplundraum, wenn sein Dualraum X^* die Radon-Nikodým-Eigenschaft hat.

(FRANCHETTI und PAYÁ [98](1993), Seite 68). G. Godefroy hat gezeigt, dass jeder Banachraum mit einer stark subdifferenzierbaren Norm ein Asplundraum ist.

In GILES [109](1982), Seite 211, und in BOURGIN [36](1983), Seite 161, Definition 5.7.3, findet man die Definition, wann ein Banachraum über $\mathbb{R}$, der ein Prädual besitzt, ein sogenannter *schwach*-Asplund-Raum* ist. Dieser Definition entsprechend, ist ein Banachraum über $\mathbb{K}$, der ein Prädual besitzt, genau dann ein schwach*-Asplund-Raum, wenn sein Prädual die Radon-Nikodým-Eigenschaft aufweist. Es sei an 2.10.12 erinnert, wonach das Prädual eines schwach*-Asplund-Raumes stets stark eindeutig ist.

4.1.9. (a) (ROSENTHAL [249](1984), Seite 134). Sei X ein Banachraum über $\mathbb{K}$ und $T \in \mathcal{L}(X)$. Dann gilt:

$$\sup \operatorname{Re} V_{\text{spatial}}(T) = \max \operatorname{Re} V(T) = \tilde{\nabla}_T \cdot \|(\operatorname{Id}_X) = \tilde{\nabla}_T \ln \|\exp(\cdot)\|(0),$$

wobei $\ln$ den auf $\mathbb{R}^{+\times}$ definierten üblichen natürlichen Logarithmus bezeichnet.

(b) (DUNFORD und SCHWARTZ [78](1958), Seite 447, Theorem 5). Sei X ein Banachraum über $\mathbb{K}$ und $x \in S_X$. Dann gilt für alle $y \in X$ die Gleichung

$$\max \operatorname{Re} V(x; y) = \tilde{\nabla}_y \cdot \|(x).$$

4.1.10. Seien X und Y normierte Vektorräume über $\mathbb{K}$ und sei U ein Unterraum von X . Ist $T \in \mathcal{L}(U, Y)$ eine Ableitung im Sinne einer der vorstehenden Nummern, so

wird im Fall von $X \cong \mathbb{K}$ üblicherweise T per der kanonischen surjektiven Isometrie $\varphi: \mathcal{L}(\mathbb{K}, Y) \rightarrow Y$, $R \mapsto R(1)$, als Element von Y betrachtet. Dass φ eine Isometrie ist, sieht man so: Da jedes $R \in \mathcal{L}(\mathbb{K}, Y)$ von der Form $t \mapsto ty$ für ein $y \in Y$ ist, gilt: $\|\varphi(R)\| = \|R(1)\| = \|y\| = \sup\{\|ty\| : |t| = 1\} = \|R\|$.

4.1.11 Glätte und Rauheit (MEGGINSON [211](1998), Seiten 485-487, 504; BECERRA GUERRERO und MARTÍN [18](2005)). Sei $(X, \|\cdot\|)$ ein normierter Vektorraum über $\mathbb{K}$ und $a \in S_X$. Es sei an Definition 2.11.1 erinnert. Die Norm $\|\cdot\|$ ist genau dann glatt bei a , wenn sie Gâteaux-differenzierbar bei a ist, und in diesem Fall gilt $G\|\cdot\|(a) = \operatorname{Re} x^*$, wobei x^* das eine Element von $\Phi_X(a)$ bezeichne. Insbesondere ist X also genau dann glatt, wenn die Norm von X Gâteaux-differenzierbar ist. Die Norm $\|\cdot\|$ ist genau dann Gâteaux-differenzierbar bei $x \neq 0$, wenn sie bei $x/\|x\|$ Gâteaux-differenzierbar ist. Analog gilt: Die Norm $\|\cdot\|$ ist genau dann Fréchet-differenzierbar bei $x \neq 0$, wenn sie bei $x/\|x\|$ Fréchet-differenzierbar ist.

Die Norm $\|\cdot\|$ heißt *Fréchet-glatt* (im Englischen *Fréchet smooth*) bei a , wenn sie bei a Fréchet-differenzierbar ist. X und seine Norm heißen *Fréchet-glatt*, wenn die Norm $\|\cdot\|$ bei allen $x \in S_X$ Fréchet-differenzierbar ist. X ist genau dann Fréchet-glatt, wenn X glatt ist und die Stützfunktion Φ_X stetig ist.

Die *Rauheit* (im Englischen *roughness*) von X bei a ist definiert als

$$\eta_X(a) := \limsup_{\|x\| \rightarrow 0} \frac{\|a+x\| + \|a-x\| - 2}{\|x\|}.$$

Synonym spricht man auch von der Rauheit der Norm von X bei a und dies auch entsprechend bei den beiden folgenden über die Rauheit definierten Begriffen. Nach DEVILLE, GODEFROY und ZIZLER [62](1993), Seite 3, ist X genau dann Fréchet-glatt bei a , wenn $\eta_X(a) = 0$ gilt.

Wegen der Dreiecksungleichung gilt offensichtlich immer $\eta_X(a) \leq 2$ ([62], Seite 8). Sei $\delta > 0$. X heißt δ -rau, falls für alle $x \in S_X$ die Ungleichung $\eta_X(x) \geq \delta$ gilt. X heißt *rau*, falls X für ein $\delta > 0$ δ -rau ist. X heißt *extrem rau* (im Englischen *extremely rough*), falls X 2-rau ist. Die kanonischen Normen von $C[0, 1]$ und ℓ_1 sind extrem rau. Das Prädual jeder unendlichdimensionalen von Neumann-Algebra über $\mathbb{C}$ ist extrem rau (BECERRA GUERRERO, LÓPEZ PÉREZ und RODRÍGUEZ PALACIOS [17](2003), Seite 758, Korollar 2.7). Nach DEVILLE, GODEFROY und ZIZLER [62](1993), Seite 8, Satz 1.11, ist im Fall, dass X ein Banachraum über $\mathbb{K}$ ist, X genau dann δ -rau, wenn jede nicht leere schwach*-Scheibe von B_{X^*} einen Durchmesser größer oder gleich δ hat.

Es sei an die Definition 2.11.11 eines stark exponierten Punktes erinnert. Nach DEVILLE, GODEFROY und ZIZLER [62](1993), Seite 3, Theorem 1.4(ii) gilt: Sei X ein Banachraum über $\mathbb{K}$, $a \in S_X$ und $x^* \in \Phi_X(a)$. Dann ist a mittels x^* genau dann ein stark exponierter Punkt von B_X , wenn die Norm von X^* bei x^* Fréchet-glatt ist.

4.1.12 Norm und Stützfunktion. (a) Sei X ein normierter Vektorraum über $\mathbb{K}$. Die Norm von X ist w -unterhalbstetig und die Norm von X^* ist w^* -unterhalbstetig.

(b) Sei X ein normierter Vektorraum über $\mathbb{K}$. Die Stützfunktion Φ_X ist $(n - w^*)$ -oberhalbstetig.

(c) Sei X ein normierter Vektorraum über $\mathbb{K}$ und $p \in S_X$. Die folgenden drei Aussagen sind äquivalent:

- (i) Die Norm von X ist stark subdifferenzierbar bei p .
- (ii) Φ_X ist $(n - \bar{n})$ -oberhalbstetig bei p .
- (iii) Für jedes $\varepsilon > 0$ enthält $\Phi_X(p) + \varepsilon \cdot B_{X^*}$ eine Scheibe von B_{X^*} .

(d) Sei X ein normierter Vektorraum über $\mathbb{K}$ und $p \in S_X$. Die folgenden drei Aussagen sind äquivalent:

- (i) Die Norm von X ist Fréchet-differenzierbar bei p .

(ii) Φ_X ist bei p sowohl $(n - \bar{n})$ -unterhalbstetig als auch $(n - \bar{n})$ -oberhalbstetig.

(iii) Φ_X ist $(n - n)$ -unterhalbstetig bei p .

(e) Sei X ein Banachraum über $\mathbb{K}$ und $p \in S_X$ und $\tau \in \{n, w, w^*\}$. Dann sind äquivalent:

(i) Φ_X ist $(n - \bar{\tau})$ -oberhalbstetig bei p .

(ii) Für jede Nullumgebung U von X^* enthält $\Phi_X(p) + U$ eine Scheibe von B_{X^*} .

(f) Sei X ein normierter Vektorraum über $\mathbb{K}$ und $p \in S_X$. Die folgenden vier Aussagen sind äquivalent:

(i) Die Norm von X ist Gâteaux-differenzierbar bei p .

(ii) Φ_X ist bei p sowohl $(n - \overline{w^*})$ -unterhalbstetig als auch $(n - \overline{w^*})$ -oberhalbstetig.

(iii) Φ_X ist bei p sowohl $(n - w^*)$ -unterhalbstetig als auch $(n - w^*)$ -oberhalbstetig.

(iv) Φ_X ist $(n - w^*)$ -unterhalbstetig bei p .

Beweis. Vorab sei darauf hingewiesen, dass die Aussagen (b) bis (f) in den gleich angegebenen Literaturstellen für $\mathbb{K} = \mathbb{R}$ formuliert sind. Wie man aber leicht sieht, sind die jeweils einzelnen Aussagen (i), (ii), usw. im Fall eines Raumes über $\mathbb{C}$ genau dann gültig, wenn sie für die reelle Strukturierung dieses Raumes gültig sind. Speziell für die Aussagen (c) und (e) sei diesbezüglich an 2.9.9 erinnert.

(a) B_X ist schwach abgeschlossen und B_{X^*} ist nach dem Satz von Alaoglu-Bourbaki schwach*-abgeschlossen. Nun siehe 1.2.9. (Man vergleiche auch mit MEGGINSON [211](1998), Seite 217, Theorem 2.5.21, und Seite 227, Theorem 2.6.14.) (b) CUDIA [55] (1964), Seite 298. (c) Siehe GREGORY [121](1980), Seite 19. Dabei beachte man, dass nach 2.11.4 $\partial\|\cdot\| \upharpoonright S_X = \Phi_X$ und $\text{ran } \partial\|\cdot\| \upharpoonright X^\times \subseteq S_{X^*}$ gilt. (d) (i) $\Leftrightarrow$ (ii): GREGORY [121](1980), Seite 18. (i) $\Leftrightarrow$ (iii): CUDIA [55](1964), Seite 304. (e) GILES, GREGORY und SIMS [108](1978), Seite 102. Man bemerke, dass man bei $\tau = w^*$ in (i) den Strich über dem τ weglassen kann, siehe 2.11.7. (f) (i) $\Leftrightarrow$ (ii): GREGORY [121](1980), Seite 19. (ii) $\Leftrightarrow$ (iii): 2.11.7. (i) $\Leftrightarrow$ (iv): CUDIA [55](1964), Seite 300. $\square$

4.1.13 (UPMEIER [279](1985), Seiten 7 und 8). Seien X und Y Banachräume über $\mathbb{K}$ und sei $\Omega \subseteq X$ offen und $n \in \mathbb{N}$. Jedes $p \in \mathcal{P}^n(X, Y)$ ist differenzierbar und die Ableitung Dp ist ein Element von $\mathcal{P}^{n-1}(X, \mathcal{L}(X, Y))$, da für alle $x, x_1 \in X$ die Gleichung $Dp(x)x_1 = n \tilde{p}(x_1, x, \dots, x)$ gilt. Des Weiteren gilt wegen $\tilde{Dp}(x_2, \dots, x_n)x_1 = n \tilde{p}(x_1, \dots, x_n)$ für alle $x_1, \dots, x_n \in X$ die Gleichung $\|\tilde{Dp}\| = n\|\tilde{p}\|$, womit die lineare Abbildung $\mathcal{P}^n(X, Y) \rightarrow \mathcal{P}^{n-1}(X, \mathcal{L}(X, Y))$, $p \mapsto Dp$, stetig ist.

Als Folgerung ergibt sich, dass jede analytische Abbildung $f: \Omega \rightarrow Y$ unendlich oft differenzierbar ist und die Ableitungen $f^{(m)}: \Omega \rightarrow \mathcal{L}^m(X, Y)$ für alle $m \in \mathbb{N}$ wiederum analytisch sind. Hat f bei $a \in \Omega$ die Potenzreihenentwicklung $\sum_{n=0}^{\infty} p_n$, so gilt $D^m f(a) = m! \tilde{p}_m$ für alle $m \in \mathbb{N}$.

4.1.14. Seien X und Y Banachräume über $\mathbb{C}$ und sei $\Omega \subseteq X$ offen. Wie in der Funktionentheorie gilt auch hier die Äquivalenz des Weierstrass-Konzeptes der komplex-analytischen Abbildungen und des Cauchy-Riemann-Konzeptes der holomorphen Abbildungen. Das heißt, eine Abbildung $f: \Omega \rightarrow Y$ ist genau dann komplex-analytisch, wenn sie holomorph ist (NACHBIN [218](1969), Seite 17).

4.2 Analytische Banach-Mannigfaltigkeiten

4.2.1 Analytische Banach-Mannigfaltigkeit (UPMEIER [279](1985)). Sei M ein topologischer Raum, der Hausdorff'sch ist. Eine *Karte* ist ein Tripel (Ω, p, X) , wobei Ω eine offene Teilmenge von M ist, X ein Banachraum über $\mathbb{K}$ und $p: \Omega \rightarrow X$ eine Abbildung ist, deren Bild $\text{ran}(p)$ eine offene Menge von X ist und die einen Homöomorphismus von Ω auf $\text{ran}(p)$ induziert. Ist $a \in \Omega$ mit $p(a) = 0$, so heißt (Ω, p, X) eine Karte *um* a . Die

Karten (Ω, p, X) und (Σ, q, Y) von M heißen *kompatibel*, wenn der Homöomorphismus $q \circ p^{-1}: p(\Omega \cap \Sigma) \rightarrow q(\Omega \cap \Sigma)$ bianalytisch ist. Ein *Atlas* von M ist eine Kollektion von paarweise kompatiblen Karten, die M überdecken. Ein bezüglich der Mengeninklusion maximaler Atlas versieht M mit der Struktur einer *Banach-Mannigfaltigkeit* (genauer *analytischen Banach-Mannigfaltigkeit*). Eine Banach-Mannigfaltigkeit über $\mathbb{K}$ ist also eine Hausdorff'sche Mannigfaltigkeit, die lokal über offene Teilmengen von Banachräumen über $\mathbb{K}$ per bianalytischen Kartenwechselabbildungen modelliert wird (KAUP [173](1977), Seite 43). Insbesondere ist zum Beispiel jede offene Teilmenge Ω eines Banachraumes X über $\mathbb{K}$ eine Banach-Mannigfaltigkeit, deren maximaler Atlas von $(\Omega, \text{Id}_X|_{\Omega}, X)$ generiert worden ist.

Eine Abbildung $g: M \rightarrow N$ zwischen Banach-Mannigfaltigkeiten heißt *analytisch*, wenn für jeden Punkt $m \in M$ eine Karte (Ω, p, X) von M und eine Karte (Σ, q, Y) von N existiert mit $m \in \Omega$, $g(\Omega) \subseteq \Sigma$ und ein kommutatives Diagramm

$$\begin{array}{ccc} \Omega & \xrightarrow{g} & \Sigma \\ p \downarrow & & \downarrow q \\ p(\Omega) & \xrightarrow{g_{\#}} & q(\Sigma) \end{array}$$

existiert, wobei $g_{\#}$ eine analytische Abbildung ist; $g_{\#}$ heißt die *lokale Darstellung* von g bezüglich p und q . Die Abbildung g heißt *bianalytisch*, wenn g bijektiv ist und sowohl g als auch g^{-1} analytisch sind. Die Gruppe aller bianalytischen Selbstabbildungen von M wird mit $\text{Aut}(M)$ bezeichnet und heißt die *analytische Automorphismengruppe* von M . Um besonders deutlich zu machen, dass sich die Automorphismen auf die Banach-Mannigfaltigkeits-Struktur beziehen, wird die Bezeichnung $\text{Aut}_{MF}(M)$ verwendet, wobei das im Index stehende „MF“ für das Wort „Mannigfaltigkeit“ steht.

Wie im endlichdimensionalen Fall werden Tangentialvektoren für Banach-Mannigfaltigkeiten mit dem Konzept des Keims einer analytischen Abbildung oder als Richtungsableitungen entlang einer glatten (im Englischen *smooth*) Kurve definiert; anders aber als im endlichdimensionalen Fall können diese Tangentialvektoren im Allgemeinen nicht mit dem Konzept der $\mathbb{K}$ -wertigen Derivationen definiert werden.

4.2.2 Keim (UPMEIER [279](1985), Seite 53). Sei M eine Banach-Mannigfaltigkeit über $\mathbb{K}$, Y ein Banachraum über $\mathbb{K}$ und $m \in M$. Betrachte Y -wertige analytische Abbildungen f , die auf einer vom jeweiligen f abhängigen offenen Umgebung Ω von m definiert sind. Zwei solcher Abbildungen heißen *äquivalent*, wenn sie auf einer offenen Umgebung von m übereinstimmen. Jede durch diese Relation definierte Äquivalenzklasse heißt ein *Keim der analytischen Y -wertigen Abbildungen bei m* . Die Menge $\mathcal{O}_m^Y$ all dieser Keime bei m ist in kanonischer Weise ein Vektorraum über $\mathbb{K}$. Bezeichnet $f_m \in \mathcal{O}_m^Y$ den Keim von f bei m , dann heißt

$$r_m: \mathcal{O}_m^Y \rightarrow Y, \quad f_m \mapsto f(m)$$

die *lineare Auswertungsabbildung* (bei m).

4.2.3 Tangentialraum (UPMEIER [279](1985), Seite 53). Sei M eine Banach-Mannigfaltigkeit über $\mathbb{K}$, Y ein Banachraum über $\mathbb{K}$ und (Ω, p, X) eine Karte von M . Für $m \in \Omega$ und $f \in \mathcal{O}_m^Y$ setze

$$\frac{\partial f}{\partial p}(m) := D(f \circ p^{-1})(pm) \in \mathcal{L}(X, Y).$$

Jeder Vektor $x \in X$ definiert eine lineare Abbildung

$$v := x \frac{\partial}{\partial p}: \mathcal{O}_m^Y \rightarrow Y, \quad f \mapsto \left(x \frac{\partial}{\partial p} \right) f := \frac{\partial f}{\partial p}(m)x.$$

Für $m \in M$ heißt die Menge $T_m(M) := \left\{ x \frac{\partial}{\partial p} : x \in X \right\}$, versehen mit der durch $x \mapsto x \frac{\partial}{\partial p}$ induzierten Banachraumstruktur, der *Tangentialraum von M bei m* ; seine Elemente heißen die *Tangentialvektoren von M bei m* .

4.2.4 Differential (UPMEIER [279](1985), Seite 55). Seien M und N zwei Banach-Mannigfaltigkeiten über $\mathbb{K}$ und sei $g: M \rightarrow N$ eine analytische Abbildung und $m \in M$. Dann existiert eine lineare Abbildung

$$g_m^*: \mathcal{O}_{g(m)}^Y \rightarrow \mathcal{O}_m^Y, \quad f \mapsto g_m^*(f) := (f \circ g)_m.$$

Für $v \in T_m(M)$ definiere man die lineare Abbildung

$$T_m(g)v: \mathcal{O}_{g(m)}^Y \rightarrow Y, \quad f \mapsto T_m(g)v(f) := v(g_m^*(f)).$$

Die lineare Abbildung

$$T_m(g): T_m(M) \rightarrow T_{g(m)}(N), \quad v \mapsto T_m(g)v$$

heißt das *Differential von g bei m* .

4.2.5 Tangentialbündel (UPMEIER [279](1985), Seite 58). Seien M und N zwei Banach-Mannigfaltigkeiten über $\mathbb{K}$. Die disjunkte Vereinigung $T(M) := \bigcup_{m \in M} \{m\} \times T_m(M)$ zusammen mit der kanonischen Projektion $\tau_M: T(M) \rightarrow M$, $(m, v) \mapsto m$, heißt das *Tangentialbündel von M* ; es ist ein sogenanntes Bündel von Banachräumen (siehe etwa LANG [195](1988)).

Sei $g: M \rightarrow N$ eine analytische Abbildung. Dann ist das *Differential* von g definiert als $T(g): T(M) \rightarrow T(N)$, $(m, v) \mapsto (g(m), T_m(g)v)$. Von g wird das folgende kommutative Diagramm induziert:

$$\begin{array}{ccc} T(M) & \xrightarrow{T(g)} & T(N) \\ \tau_M \downarrow & & \downarrow \tau_N \\ M & \xrightarrow{g} & N \end{array}$$

Ist $\Omega \subseteq M$ offen, so kann $T(\Omega)$ mit einer Teilmenge von $T(M)$ identifiziert werden per dem Differential $T(i): T(\Omega) \rightarrow T(M)$ der kanonischen Abbildung $i: \Omega \rightarrow M$. Für jede offene Teilmenge $A \subseteq X$ eines Banachraumes X kann das Tangentialbündel $T(A)$ mit $A \times X$ identifiziert werden per der Abbildung

$$k: T(A) \rightarrow A \times X, \quad \left(z, h \frac{\partial}{\partial z} \right) \mapsto (z, h).$$

Mittels der Identifikation per k kann man folglich für eine Karte (Ω, p, X) das Differential $T(p): T(\Omega) \rightarrow T(X)$ als die Abbildung $k \circ T(p)$ auffassen, also als

$$T(\Omega) \rightarrow p(\Omega) \times X, \quad \left(m, h \frac{\partial}{\partial p} \right) \mapsto (p(m), h) \quad (4.2)$$

wobei bereits auch $\tau_X \circ T(p)(T(\Omega)) = p \circ \tau_\Omega(T(\Omega)) = p(\Omega)$ berücksichtigt worden ist.

$T(M)$ trägt eine eindeutige Hausdorff-Topologie, so dass τ_M stetig ist und $T(M)$ einen bestimmten analytischen Atlas trägt (siehe UPMEIER, ebd.); mit diesem besagten Atlas ist $T(M)$ eine Banach-Mannigfaltigkeit über $\mathbb{K}$.

4.2.6 Vektorfeld (UPMEIER [279](1985), Seite 59). Sei M eine Banach-Mannigfaltigkeit über $\mathbb{K}$. Ein *analytisches Vektorfeld auf M* oder *infinitesimale Transformation* ist eine analytische Abbildung $X: M \rightarrow T(M)$ mit $\tau_M \circ X = \text{Id}_M$, das heißt, $X_m := X(m) \in$

$T_m(M)$ für alle $m \in M$. Der Vektorraum aller analytischen Vektorfelder auf M wird mit $\mathcal{T}(M)$ bezeichnet.

Sei $O \subseteq M$ offen, W ein Banachraum über $\mathbb{K}$, $f: O \rightarrow W$ eine analytische Abbildung und $X \in \mathcal{T}(M)$. Dann ist die Abbildung $Xf: O \rightarrow W$, $m \mapsto Xf(m) := X_m f_m$, mit $f_m \in \mathcal{O}_m^W$ der Keim von f bei m , analytisch.

Sei (Ω, p, Z) eine Karte von M . Für $m \in M$ gilt dann $X_m = h(m) \frac{\partial}{\partial p} \in T_m(M)$ mit $h(m) = T_m(p)X_m = X_m p_m = Xp(m) \in Z$, wobei im zweiten Gleichheitszeichen gemäß Auslegung (4.2) die Abbildung $T_m(p)$ als eine Abbildung $T_m(\Omega) \rightarrow Z$ aufgefasst wird.

Also kann das Vektorfeld X lokal, das heißt hier, beschränkt auf Ω , per der analytischen Abbildung $h = Xp: \Omega \rightarrow Z$ als $X = h \frac{\partial}{\partial p}$ geschrieben werden. Für eine beliebige analytische Abbildung $f: \Omega \rightarrow Z$, ist die analytische Abbildung $Xf: \Omega \rightarrow Z$ gegeben durch

$$Xf(m) = \left(h \frac{\partial}{\partial p} \right) f(m) = \frac{\partial f}{\partial p}(m) h(m).$$

4.2.7 Kommutator-Vektorfeld. Sei M eine Banach-Mannigfaltigkeit über $\mathbb{K}$. Zu jedem Paar $X, Y \in \mathcal{T}(M)$ existiert genau ein $Z \in \mathcal{T}(M)$, so dass für alle analytischen Abbildungen $f: Q \rightarrow W$ von offenen Teilmengen Q von M nach einem Banachraum W gilt: $Zf = Y(Xf) - X(Yf)$. Das Vektorfeld Z wird mit $[X, Y]$ bezeichnet und heißt der *Kommutator* von X und Y . $\mathcal{T}(M)$, versehen mit dem Kommutator als Produkt, ist eine Lie-Algebra und wird ebenfalls mit $\mathcal{T}(M)$ bezeichnet. (UPMEIER [279](1985), Seite 60.)

(KAUP [180], Seite 38, Beispiel 13.7(iii)). Dieses Kommutatorprodukt von Vektorfeldern ist in Abstimmung mit der Definition 2.1.29 des Kommutators für Algebren definiert: Ist X ein Banachraum über $\mathbb{C}$ und sind $S, T \in \mathcal{L}(X)$, so gilt für den Kommutator $[S, T] = ST - TS \in \mathcal{L}(X)$ die Gleichung

$$\left[S \frac{\partial}{\partial z}, T \frac{\partial}{\partial z} \right] = [S, T] \frac{\partial}{\partial z}.$$

Dies sei im Hinblick darauf erwähnt, dass manche Autoren das Kommutatorprodukt von Vektorfeldern mit umgekehrtem Vorzeichen definieren.

4.2.8 (UPMEIER [279](1985), Seite 61). Seien M und N zwei Banach-Mannigfaltigkeiten über $\mathbb{K}$. Sei $g: M \rightarrow N$ eine bianalytische Abbildung. Dann wird der Lie-Algebra-Isomorphismus $\mathcal{T}(M) \rightarrow \mathcal{T}(N)$, $X \mapsto T(g) \circ X \circ g^{-1}$ mit g_* bezeichnet. g_* ist also genau gerade so definiert, dass das folgende Diagramm kommutiert:

$$\begin{array}{ccc} T(M) & \xrightarrow{T(g)} & T(N) \\ \uparrow \mathcal{T}(M) \ni X & & \uparrow g_* X \in \mathcal{T}(N) \\ M & \xrightarrow{g} & N \end{array}$$

Ist $h: L \rightarrow M$ eine weitere bianalytische Abbildung, und ist die Komposition $g \circ h$ erklärt, so gilt

$$g_* \circ h_* = (g \circ h)_*. \quad (4.3)$$

Ist (Ω, p, X) eine Karte von M und ist (Σ, q, Y) eine Karte von N mit $g(\Omega) \subseteq \Sigma$, dann gilt

$$g_* \left(h \frac{\partial}{\partial p} \right)_{g(m)} = \left(\frac{\partial (q \circ g)}{\partial p}(m) h(m) \right) \frac{\partial}{\partial q} \Big|_{g(m)}$$

für alle $m \in \Omega$ und jeder analytischen Funktion $h: \Omega \rightarrow X$.

4.2.9 (UPMEIER [279](1985), Seite 62). Man betrachte folgenden Spezialfall: Sei Z ein Banachraum über $\mathbb{K}$ und $\Omega \subseteq Z$ offen. Dann haben die analytischen Vektorfelder auf Ω die Form $X = h \frac{\partial}{\partial z} = h(z) \frac{\partial}{\partial z}$ mit analytischem $h = X(\text{Id}_\Omega) : \Omega \rightarrow Z$; dabei wird üblicherweise des Öfteren $h(z)$ für die Funktion h geschrieben, und das z in $\frac{\partial}{\partial z}$ steht für die kanonische Kartenabbildung $\text{Id}_Z|_\Omega$; insbesondere steht $z \frac{\partial}{\partial z}$ für $\text{Id}_Z|_\Omega \frac{\partial}{\partial z}$. Ist $f : \Omega \rightarrow W$ eine analytische Abbildung, W ein Banachraum über $\mathbb{K}$, so gilt $\left(h \frac{\partial}{\partial z}\right) f(z) = Df(z)h(z)$. Das Kommutatorprodukt in $\mathcal{T}(\Omega)$ ist gegeben durch

$$\left[h(z) \frac{\partial}{\partial z}, k(z) \frac{\partial}{\partial z}\right] = \left(Dh(z)k(z) - Dk(z)h(z)\right) \frac{\partial}{\partial z}.$$

Ist $g : \Omega \rightarrow \Sigma$ eine bianalytische Abbildung von Ω auf eine offene Menge Σ eines Banachraumes Y über $\mathbb{K}$, dann ist der Isomorphismus $g_* : \mathcal{T}(\Omega) \rightarrow \mathcal{T}(\Sigma)$ von 4.2.8 gegeben durch

$$g_* \left(h(z) \frac{\partial}{\partial z} \right) = k(w) \frac{\partial}{\partial w},$$

mit $k(w) = Dg(g^{-1}(w)) \left(h(g^{-1}(w)) \right)$ für alle $w \in \Sigma$, das heißt, $k(h(z)) = Dg(z)(h(z))$ für alle $z \in \Omega$.

Für jedes $n \in \mathbb{Z}$, $n \geq -1$, bezeichnet

$$\mathcal{T}_n(Z) := \left\{ h \frac{\partial}{\partial z} : h \in \mathcal{P}^{n+1}(Z) \right\}$$

den Vektorraum über $\mathbb{K}$ aller *polynomialen Vektorfelder auf Z , die homogen vom Grad $n+1$ sind*. Die algebraische direkte Summe

$$\mathcal{P}(Z) := \sum_{n \geq -1}^{\oplus} \mathcal{T}_n(Z)$$

aller polynomialen Vektorfelder auf Z ist eine Unteralgebra von $\mathcal{T}(Z)$, versehen mit einer additiven Graduierung.

4.2.10 Fluss (UPMEIER [279](1985), Seiten 67, 70 und 74). Sei Ω eine offene Teilmenge einer Banach-Mannigfaltigkeit M über $\mathbb{K}$ und $\Sigma \subseteq \mathbb{R} \times \Omega$ offen. Sei $r : \Sigma \rightarrow M$ eine analytische Abbildung, so dass für jedes $m \in \Omega$ gilt:

$$\Sigma_m := \{t \in \mathbb{R} : (t, m) \in \Sigma\}$$

ist ein offenes Intervall mit $0 \in \Sigma_m$ und $r(0, m) = m$. Gilt dann für alle $m \in \Omega$, $t \in \Sigma_m$ mit $r(t, m) \in \Omega$ und für alle $s \in \Sigma_{r(t, m)} \cap (\Sigma_m - t)$ die Gleichung $r(s + t, m) = r(s, r(t, m))$, so heißt r ein *lokaler analytischer Fluss auf Ω* . Die analytische Abbildung $r_m : \Sigma_m \rightarrow M$, $t \mapsto r(t, m)$, heißt die *Auswertungsabbildung* (im Englischen *evaluation mapping*) bei $m \in \Omega$.

Das Vektorfeld X auf Ω definiert durch $X_m := T_0(r_m)$ für alle $m \in \Omega$ heißt der *infinitesimale Generator* (oder das *Differential*) des lokalen analytischen Flusses r auf Ω . Gilt $\Omega = M$, so heißt r ein *analytischer Fluss auf M* .

Ist speziell Z ein Banachraum über $\mathbb{K}$, sind $\Omega \subseteq C$ zwei offene Teilmengen von Z und ist $r : I \times \Omega \rightarrow C$ eine analytische Abbildung für ein offenes Intervall $I \subseteq \mathbb{R}$ mit $0 \in I$, so ist r genau dann ein lokaler analytischer Fluss auf Ω mit infinitesimalem Generator $X = h \frac{\partial}{\partial z} \in \mathcal{T}(C)$, wenn r die Differentialgleichung $\frac{\partial}{\partial t} r(t, z) = h(r(t, z))$ für alle $(t, z) \in I \times \Omega$, unter der Anfangsbedingung $r(0, z) = z$ für alle $z \in \Omega$, löst.

4.2.11 Vollständiges Vektorfeld (UPMEIER [279](1985), Seite 80). Sei M eine Banach-Mannigfaltigkeit über $\mathbb{K}$. Ein analytisches Vektorfeld X auf M heißt *vollständig* (im Englischen *complete*) oder auch *integrierbar*, wenn es der infinitesimale Generator eines analytischen Flusses $r_X: \mathbb{R} \times M \rightarrow M$ ist. Die Menge aller vollständigen analytischen Vektorfelder auf M wird mit $\text{aut}(M)$ bezeichnet. Wie bei $\text{Aut}(M)$ kann man auch hier die Bezeichnung $\text{aut}_{MF}(M)$ verwenden, wenn man den Bezug auf die Banach-Mannigfaltigkeits-Struktur deutlich machen möchte. Die Elemente von $\text{aut}(M)$ werden auch die *infinitesimalen Automorphismen von M* genannt. $\text{aut}(M)$ ist im Allgemeinen nicht abgeschlossen bezüglich der Addition oder des Kommutator-Produktes. $\text{aut}(M)$ ist in $\mathcal{T}(M)_{\mathbb{R}}$ ein punktierter Kegel (KAUP und UPMEIER [182](1977), Seite 383). Von Interesse sind die in $\text{aut}(M)$ enthaltenen Banach-Lie-Algebren von Vektorfeldern. In manchen Fällen kann man auch $\text{aut}(M)$ selbst mit der Struktur einer Banach-Lie-Algebra ausstatten, siehe 4.5.28.

Für $X \in \text{aut}(M)$ und $t \in \mathbb{R}$ definiert man die analytische Abbildung $\exp(tX): M \rightarrow M$ per $\exp(tX)(m) := r_X(t, m)$ für alle $m \in M$. Der Homomorphismus $\mathbb{R} \rightarrow \text{Aut}(M)$, $t \mapsto \exp(tX)$, heißt die zu $X \in \text{aut}(M)$ *assozierte Einparametergruppe*. Die Abbildung

$$\exp: \text{aut}(M) \rightarrow \text{Aut}(M), \quad X \mapsto \exp(1 \cdot X). \quad (4.4)$$

heißt die *Exponentialabbildung*; dabei wird das von ihr zu einem $X \in \text{aut}(M)$ zugeordnete $\exp(X) := \exp(1 \cdot X)$ das *Exponential* von X genannt.

Sei X eine Banach-Lie-Algebra über $\mathbb{K}$ und M eine Banach-Mannigfaltigkeit über $\mathbb{L}$, wobei $\mathbb{L}$ für $\mathbb{R}$ oder $\mathbb{C}$ steht. Eine *Darstellung von X auf M* ist definiert als ein reell-linearer Lie-Algebra-Homomorphismus $T: X \rightarrow \mathcal{T}(M)$ mit $T(X) \subseteq \text{aut}(M)$ (UPMEIER [279](1985), Seite 86).

4.2.12 Lie-Gruppe und deren Lie-Algebra (UPMEIER [279](1985), Seiten 92, 93, 95 und 99). Eine *Banach-Lie-Gruppe* über $\mathbb{K}$ ist eine Gruppe G , die zugleich eine Banach-Mannigfaltigkeit über $\mathbb{K}$ ist, so dass die Multiplikationsabbildung $G \times G \rightarrow G$, $(g, h) \mapsto gh$, analytisch ist. Das neutrale Element von G wird mit e bezeichnet. Die Links- und Rechtstranslationen $L_g: h \mapsto gh$ und $R_g: h \mapsto hg$ sind bianalytische Automorphismen der Banach-Mannigfaltigkeit G , das heißt, Elemente von $\text{Aut}_{MF}(G)$.

Die Lie-Unteralgebra von $\mathcal{T}(G)$

$$\mathfrak{g}(G) := \{X \in \mathcal{T}(G) : (R_g)_* X = X \text{ für alle } g \in G\}$$

aller rechtsinvarianten analytischen Vektorfelder auf G heißt die *Lie-Algebra von G* . Es gilt

$$\mathfrak{g}(G) \subseteq \text{aut}(G).$$

Sei $X \in \mathfrak{g}(G)$; setze $\chi := (\exp(X))(e)$. Dann gilt $\exp(X) = L_\chi$, das heißt, $\exp(X)$ ist die Linkstranslation $h \mapsto \chi h$. Die Abbildung

$$\mathfrak{g}(G) \rightarrow G, \quad X \mapsto (\exp(X))(e)$$

heißt die *Exponentialabbildung von G* und wird mit $\exp$ bezeichnet.

Die Inversenbildung $g \mapsto g^{-1}$ einer Banach-Lie-Gruppe G über $\mathbb{K}$ ist analytisch, siehe ebd., Korollar 6.6.

Wie bereits erwähnt, definieren manche Autoren das Kommutatorprodukt 4.2.7 zweier Vektorfelder mit umgekehrtem Vorzeichen und dementsprechend die Lie-Algebra $\mathfrak{g}(G)$ einer Banach-Lie-Gruppe G als die Menge aller linksinvarianten analytischen Vektorfelder auf G .

Sei G eine Banach-Lie-Gruppe mit Lie-Algebra $\mathfrak{g}$. Die Auswertungsabbildung $\rho_e: \mathfrak{g} \rightarrow T_e(G)$, $X \mapsto X_e$, ist ein linearer Isomorphismus mit der inversen Abbildung $T_e(G) \rightarrow \mathfrak{g}$, $v \mapsto X_v$ mit $(X_v)_g := T_e(R_g)v$ für alle $g \in G$.

4.2.13 (UPMEIER [279](1985), Seite 96). Ist A eine unitale Banachalgebra über $\mathbb{K}$, so ist die Gruppe $G(A)$ aller invertierbaren Elemente von A offen, also eine Banach-Mannigfaltigkeit, und eine Banach-Lie-Gruppe über $\mathbb{K}$. Da man den Tangentialraum $T_e(A)$ mit A identifizieren kann, ist die Auswertungsabbildung $\rho_e: \mathfrak{g}(G(A)) \rightarrow A$ ein Banach-Lie-Gruppen-Isomorphismus von $\mathfrak{g}(G(A))$ auf $A^\ominus$. Insbesondere gilt für jeden Banachraum X über $\mathbb{K}$, dass die Gruppe $G\mathcal{L}(X)$ eine Banach-Lie-Gruppe über $\mathbb{K}$ ist. Ihre Lie-Algebra $\mathfrak{g}(G\mathcal{L}(X))$ ist mittels der Auswertungsabbildung Banach-Lie-Gruppenisomorph zu $(\mathcal{L}(X))^\ominus$.

4.3 Operation und Darstellung einer Gruppe

4.3.1 Definition. Sei G eine Gruppe, e ihr neutrales Element und M eine Menge. Eine Abbildung $\circ: G \times M \rightarrow M$ heißt eine *Operation* (genauer eine *Links-Operation*) von der Gruppe G auf M , wenn die beiden folgenden Bedingungen gelten:

- (a) $(gh) \circ m = g \circ (h \circ m)$ für alle $g, h \in G, m \in M$.
- (b) $e \circ m = m$ für alle $m \in M$.

Man sagt dann auch, dass die Gruppe G auf der Menge M operiere. Im Allgemeinen wird $\circ$ multiplikativ geschrieben. Man siehe auch bei 4.3.7 bezüglich des Begriffes der Darstellung einer Gruppe.

4.3.2. Sei G eine Gruppe, die auf einer Menge M operiert. Dann ist auf M durch

$$m \sim n :\Leftrightarrow \exists g \in G : gm = n$$

eine Äquivalenzrelation erklärt. Sei $m \in M$. Die Äquivalenzklasse $[m] := Gm = \{gm : g \in G\}$ heißt *Bahn* (oder *Orbit* oder *Transitivitätsgebiet*) von m bei der gegebenen Operation. Die Operation von G auf M heißt *transitiv*, wenn nur genau eine Bahn existiert, das heißt, wenn $M \neq \emptyset$ ist und zu je zwei Elementen $m, n \in M$ stets ein $g \in G$ mit $gm = n$ existiert. Man sagt dann auch, G operiere *transitiv* auf M .

4.3.3 Fixsysteme. Sei G eine Gruppe, die auf einer Menge M operiere. Fasse die Potenzmenge von M als eine mit der leeren Menge als Basispunkt punktierte Menge $(\mathcal{P}(M), \emptyset)$ auf. Dann ist mittels der Abbildung

$$B: G \times M \rightarrow \mathcal{P}(M), \quad (g, m) \mapsto \{gm\} \setminus \{m\}$$

ein Annihilatorsystem $(G, M; B; \mathcal{P}(M))$ erklärt; es heißt *das Fixsystem der auf M operierenden Gruppe G* .

Für jede Teilmenge Z von M heißt der Linksannihilator von Z , also die Untergruppe (wegen siehe 2.1.4)

$$Z_\perp = \{g \in G : gm = m \text{ für alle } m \in Z\}$$

von G , die genau aus den Elementen von G besteht, die jedes $m \in Z$ festlassen — man sagt dazu auch: die Z elementweise festlassen —, der *Stabilisator von Z (in G)*; andere übliche Bezeichnungen für ihn sind *Stabilitätsuntergruppe*, *Stand(unter)gruppe*, *Fixgruppe* und *Isotropie-(Unter)gruppe von Z (in G)*; falls Z einelementig ist, etwa $Z = \{m\}$, so sagt man in den vorstehenden Formulierungen statt „von Z “ auch einfach „von m “.

Für jede Teilmenge H von G heißt der Rechtsannihilator von H ,

$$H^\perp = \{m \in M : gm = m \text{ für alle } g \in H\},$$

also die Teilmenge von M , die aus genau denjenigen Elementen $m \in M$ besteht, die unter allen $g \in H$ festbleiben, die *Fixpunktmenge* von H .

Man bemerke: Würde man ein Fixsystem statt mit einer auf eine Menge M operierenden Gruppe allgemeiner mit einer Menge von Selbstabbildungen einer Menge M definieren, so wäre die in 1.1.2 definierte Fixpunktmenge einer Selbstabbildung f gleich der hier definierten Fixpunktmenge der einelementigen Selbstabbildungsmenge $\{f\}$.

4.3.4 Galoissysteme. (a) Sei E ein Körper. Dann heißt das Fixsystem der auf E operierenden Automorphismengruppe $\text{Aut}(E)$ das *Galoissystem* von E . Da in der sogenannten Galoistheorie parallel zu dem Galoissystem eines Körpers E auch die Galoissysteme von etwa Unterkörpern L von E betrachtet werden, empfiehlt es sich dann, statt $\perp$ etwa $\perp(L)$ als Annihilatorsymbol zu verwenden.

Sei E ein Körper und betrachte das Galoissystem von E . Sei K ein Unterkörper von E . Mit E , aufgefasst als einen Vektorraum über K , wird die Dimension von E — wie in der Körpertheorie üblich — mit $[E : K]$ bezeichnet. Der Stabilisator von K , also $K_\perp$, heißt die *Galoisgruppe* von K (bezüglich E); sie wird auch mit $\text{Gal}(E : K)$ oder $\gamma_E(K)$ oder einfach $\gamma(K)$ bezeichnet. (Siehe etwa bei SCHULZE [258](2006), Seite 100, oder MEYBERG [213](1976), Seite 65, Definition 7.1.3.) Man überlegt sich leicht (siehe MEYBERG [213](1976), Seite 65), dass $K_\perp$ genau aus den K -linearen (siehe Definition 2.1.35) Elementen von $\text{Aut}(E)$ besteht.

Für jede Teilmenge G von $\text{Aut}(E)$ ist $\kappa := G^\perp$ ein Körper, der sogenannte *Fixpunktkörper* von G .

(HORNFECK [141](1976), Seiten 208, 209). Weist $G \subseteq \text{Aut}(E)$ nur endlich viele Elemente auf, wobei auch $\text{Id}_E \in G$ gelte, so gilt

$$[E : G^\perp] \geq |G|, \quad (4.5)$$

wobei $|G|$ die Elementanzahl von G bezeichnet; ist G dabei sogar eine Gruppe, so liegt in (4.5) Gleichheit vor.

(b) Betrachte das Galoissystem eines Körpers E . Für jeden Unterkörper K von E mit $[E : K] < \infty$ ist die Galoisgruppe $K_\perp$ endlich.

Beweis. Angenommen, $K_\perp$ wäre dann nicht endlich. Dann gibt es eine endliche Menge $M \subseteq K_\perp$ mit $\text{Id}_E \in M$ und $|M| > [E : K]$. Nach 1.1.19(d) ist die Inklusion $M \subseteq K_\perp$ äquivalent zu der Inklusion $K \subseteq M^\perp$. Somit gilt die Ungleichung $[E : M^\perp] \leq [E : K]$. Nach der Ungleichung (4.5) gilt aber $|M| \leq [E : M^\perp]$ und man hat den *Widerspruch* $|M| \leq [E : K]$. $\square$

(c) Betrachte das Galoissystem eines Körpers E . Jede endliche Untergruppe von $\text{Aut}(E)$ ist annihilatorabgeschlossen.

Beweis. Sei G eine endliche Untergruppe von $\text{Aut}(E)$. nach 1.1.19(e) gilt $G^\perp = G^{\perp\perp\perp}$, nach 1.1.19(d) gilt $G \subseteq G^{\perp\perp}$. Nach (4.5) gilt $|G| = [E : G^\perp]$, also $[E : G^\perp] < \infty$, und somit mit (b), dass $|G^{\perp\perp}| < \infty$. Wieder mit (4.5) also $|G| = [E : G^{\perp\perp\perp}] = |G^{\perp\perp}|$, das heißt, $G = G^{\perp\perp}$, und das war zu zeigen. $\square$

(d) Betrachte das Galoissystem eines Körpers E . Sei K ein Unterkörper von E . Dann heißt K ein *Galois(unter)körper* (von E) oder *galoissch*, wenn er annihilatorabgeschlossen ist. (Abweichend von der hier verwendeten Terminologie bezeichnet HORNFECK [141](1976), Definition auf Seite 211, für $[E : K] < \infty$ den Stabilisator $K_\perp$ erst genau dann als die *Galoisgruppe* $G(E|K)$, wenn K annihilatorabgeschlossen ist, und in diesem Fall nennt er dann E *normal* über K .)

(e) (HORNFECK [141](1976), Seite 213, Satz 66.1). Betrachte das Galoissystem eines Körpers E . Sei K ein Galoisunterkörper von E mit $[E : K] < \infty$. Dann ist auch jeder

Zwischenkörper L von E und K — soll heißen, L ist ein Unterkörper von E und enthält K als Unterkörper, $E \supseteq L \supseteq K$ — ein Galoisunterkörper von E .

(f) (HORNFECK [141](1976), Seiten 214, 215). Betrachte das Galoissystem eines Körpers E . Sei K ein Galoisunterkörper von E mit $[E : K] < \infty$. Sei L ein Zwischenkörper von E und K . Nun betrachte zusätzlich auch das Galoissystem von L . Genau dann ist K auch ein Galoisunterkörper von L , wenn die Menge aller Restriktionen $\sigma \upharpoonright L$ der $\sigma \in K_{\perp(E)}$ auf L die Gruppe $K_{\perp(L)}$ ist (symbolisch also etwa $K_{\perp(E)} \upharpoonright L = K_{\perp(L)}$); und dies ist bereits genau dann der Fall, wenn für jedes $\sigma \in K_{\perp(E)}$ die Gleichung $\sigma(L) = L$ gilt.

(g) (HORNFECK [141](1976), Seite 215). Betrachte das Galoissystem eines Körpers E . Sei K ein Galoisunterkörper von E mit $[E : K] < \infty$. Sei M_1 die Menge aller Untergruppen von $K_{\perp}$. Sei M_2 die Menge aller Zwischenkörper L von E und K . Dann ist die Bildung des Rechtsannihilators eine bijektive Abbildung von M_1 auf M_2 mit der Linksannihilatorbildung als ihre inverse Abbildung.

(h) In Fortsetzung von (f) gilt: Genau dann ist K ein Galoisunterkörper von L , wenn $L_{\perp(E)}$ ein Normalteiler von $K_{\perp(E)}$ ist; und in diesem Fall ist dann $K_{\perp(L)}$ isomorph zu der Faktorgruppe (siehe ebd., Seite 52) $K_{\perp(E)} / L_{\perp(E)}$.

4.3.5 Affine Varietäten und Ideale (COX, LITTLE und O'SHEA [54](1997)). Sei K ein Körper. Sei $n \in \mathbb{N}^{\times}$. Sei P die Algebra $K[x_1, \dots, x_n]$ über K der Polynome über n unabhängigen Unbestimmten $x_1, \dots, x_n$ mit Koeffizienten aus K . Sei X der Vektorraum K^n über K . Jedes Polynom $p \in P$ liefert kanonisch eine Funktion $p^* : X \rightarrow K$, indem in dem Polynom p für ein gegebenes $(a_1, \dots, a_n) \in X$ jedes x_i durch a_i , $i = 1, \dots, n$, ersetzt wird; falls keine Missverständnisse zu befürchten sind — also speziell bei unendlichdimensionalem K , wo dann $p = 0 \Leftrightarrow p^* = 0$ gilt, siehe ebd., Seite 3, Satz 5 —, wird diese Funktion auch wieder einfach mit p bezeichnet. Man definiere die folgende Abbildung:

$$B: X \times P \rightarrow K, \quad ((a_1, \dots, a_n), p) \mapsto p^*(a_1, \dots, a_n).$$

Dann betrachte man das Annihilatorsystem

$$(K^n, K[x_1, \dots, x_n]) := (X, P; B; K).$$

Es ist offensichtlich stets links trennend; ist K unendlichdimensional, so ist es auch rechts trennend. Für jede Teilmenge M von X ist der Rechtsannihilator von M , also $M^{\perp}$, ein Ideal von P , man nennt es *das Ideal von M* . Für jede Teilmenge N von P ist der Linksannihilator von N , also $N_{\perp}$, die *durch N definierte affine Varietät*; sie ist im Allgemeinen kein Untervektorraum von X . Man bemerke, dass sie aber wegen 2.1.57 und 2.1.55 stets gleich dem Linksannihilator einer endlichen Teilmenge von P ist. Des Weiteren bemerke man, dass wegen 1.1.19(b) der Durchschnitt beliebig vieler affiner Varietäten wieder eine affine Varietät ist. Die annihilatorabgeschlossene Hülle einer Teilmenge M von X , also $M^{\perp\perp}$, nennt man den *Zariski-Abschluss* von M .

4.3.6 Definition. Ist M eine Banach-Mannigfaltigkeit über $\mathbb{K}$ und operiert $\text{Aut}_{MF}(M)$ auf M transitiv, so heißt M *homogen*.

4.3.7 Darstellung einer Gruppe. Sei G eine Gruppe und M eine Banach-Mannigfaltigkeit über $\mathbb{K}$. Eine *Darstellung von G auf M* ist ein Homomorphismus $r: G \rightarrow \text{Aut}_{MF}(M)$. UPMEIER [279](1985) nennt die Darstellung von G auf M im Englischen *action of G on M* . Die Abbildung $r: G \times M \rightarrow M$, $(g, m) \mapsto r(g)(m)$ heißt die *mit der Darstellung r assoziierte Auswertungsabbildung*; sie ist gemäß Definition 4.3.1 die Operation der Gruppe G auf M . Die Abbildung $r_m: G \rightarrow M$, $g \mapsto r(g, m)$ ist die Auswertungsabbildung bei $m \in M$. Eine Darstellung r heißt *treu*, wenn sie injektiv

ist. Eine Darstellung r von einer topologischen Gruppe G nach M heißt *stetig*, wenn die mit der Darstellung r assoziierte Auswertungsabbildung stetig ist. Eine Darstellung r von einer Banach-Lie-Gruppe auf M heißt *analytisch*, wenn die mit der Darstellung r assoziierte Auswertungsabbildung analytisch ist.

4.3.8 Differential einer Lie-Gruppen-Darstellung (UPMEIER [279](1985), Seite 99). Sei r eine analytische Darstellung einer Banach-Lie-Gruppe G auf eine Banach-Mannigfaltigkeit M . Dann existiert genau eine analytische Darstellung ρ der Lie-Algebra $\mathfrak{g}(G)$ auf M , so dass das folgende Diagramm kommutiert:

$$\begin{array}{ccc} G & \xrightarrow{r} & \text{Aut}_{MF}(M) \\ \exp \uparrow & & \uparrow \exp \\ \mathfrak{g}(G) & \xrightarrow{\rho} & \text{aut}(M) ; \end{array}$$

ρ heißt das *Differential* von r .

4.3.9 Konjugationen einer Gruppe (STORCH und WIEBE [268]; SCHULZE [258](2006), Seite 27; UPMEIER [279](1985), Seite 103). Sei G eine Gruppe und $a \in G$. Dann heißt die Abbildung

$$\text{Int}(a): G \rightarrow G, \quad x \mapsto axa^{-1}$$

die *Konjugation von G mit a* oder der *von a induzierte innere Automorphismus von G* . Für eine Veranschaulichung sei bemerkt, dass für alle $x \in G$ stets die Gleichung $ax = \text{Int}(a)(x)a$ gilt. Die Abbildung $G \times G \rightarrow G$, $(g, x) \mapsto gxg^{-1}$ ist nach Definition 4.3.1 eine Operation von G auf sich. Bezüglich dieser Operation heißt für $x \in G$ die Bahn von x , also $\text{Int}(G)x$, die *Konjugationsklasse von x* . Falls für alle $x \in G \setminus \{e\}$ die Konjugationsklasse von x unendlich ist, heißt G eine *ICC-Gruppe* (vom englischen *Infinite Conjugacy Class*); Beispiele sind die Gruppe $\mathfrak{S}_\infty$ (siehe 2.1.3) und die freien Gruppen vom Rang n , $n > 1$ (bezüglich der Definition einer freien Gruppe siehe etwa GOLDBERGER und EHRLICH [115](1970), Kapitel 1, Abschnitt 11, oder HUNGERFORD [143](1980), Kapitel 1, Abschnitt 9). Für $a \in G$ heißt die Fixpunktmenge von $\text{Int}(a)$ der *Normalisator von a* ; er ist eine Untergruppe von G und sei hier mit $\{a\}'$ bezeichnet; man bemerke, dass der Normalisator von a übereinstimmt mit dem Stabilisator von a bezüglich der oben genannten Operation $(g, x) \mapsto gxg^{-1}$. Ist M eine Teilmenge von G , so heißt $M' := \bigcap_{a \in M} (\{a\}')$ der *Zentralisator von M* ; $M' \cap M = \{g \in M : \text{Int}(a)g = g \text{ für alle } a \in M\}$ heißt das *Zentrum von M* .

Die Abbildung $\text{Int}: G \rightarrow \text{Aut}(G)$, $g \mapsto \text{Int}(g)$ ist ein Gruppenhomomorphismus, dessen Kern das Zentrum von G ist. Ist G eine Banach-Lie-Gruppe über $\mathbb{K}$, dann ist $\text{Int}(a)$ für alle $a \in G$ bianalytisch.

Ist die Gruppe $(G, \cdot)$ mit multiplikativ geschriebener Verknüpfung derart, dass auf G die Struktur eines Ringes $(G, +, \cdot)$ existiert, so fällt offensichtlich der Begriff des Zentralisators mit dem der Kommutante zusammen; insbesondere gilt das für den Begriff des Zentrums.

4.3.10 (UPMEIER [279](1985), Seiten 104-107). Sei G eine Banach-Lie-Gruppe über $\mathbb{K}$. Da nach 4.3.9 für alle $g \in G$ die Konjugation $\text{Int}(g)$ ein analytischer Gruppen-Homomorphismus ist, kommutiert nach UPMEIER, ebd., Beispiel 6.21, für alle $g \in G$ das folgende Diagramm:

$$\begin{array}{ccc}
 G & \xrightarrow{\text{Int}(g)} & G \\
 \exp \uparrow & & \uparrow \exp \\
 \mathfrak{g}(G) & \xrightarrow{\text{Int}(g)_*} & \mathfrak{g}(G) \\
 \rho_e \downarrow & & \downarrow \rho_e \\
 T_e(G) & \xrightarrow{T_e(\text{Int}(g))} & T_e(G).
 \end{array} \tag{4.6}$$

Dabei gibt der untere Teil die in 4.2.12 erwähnten Identifikation von $\mathfrak{g}(G)$ mit $T_e(G)$ wieder.

Jedes $g \in G$ induziert einen Automorphismus

$$\text{Ad}(g) := (L_g)_* \upharpoonright \mathfrak{g}(G) \in \text{Aut}(\mathfrak{g}(G)),$$

wobei L_g die Links-Translation $L_g: h \mapsto gh$ von G ist. Da nach Definition die Vektorfelder $X \in \mathfrak{g}(G)$ rechtsinvariant sind, gilt für alle $X \in \mathfrak{g}(G)$: $\text{Ad}(g)X = (L_g)_*X = (L_g)_*(R_{g^{-1}})_*X = (L_g R_{g^{-1}})_*X = (L_g R_{g^{-1}})_*X = (L_g R_{g^{-1}})_*X = (\text{Int}(g))_*X$. Das heißt, es gilt

$$\text{Ad}(g) = \text{Int}(g)_* \upharpoonright \mathfrak{g}(G).$$

Gemäß Diagramm (4.6) kann $\text{Ad}(g)$ als das Differential des Automorphismus $\text{Int}(g) = L_g R_g^{-1}$ von G betrachtet werden: $(\text{Ad}(g)X)_e = (T_e(\text{Int}(g)))(X_e)$. Es gilt $\text{Ad}(g) \in \text{GL}(\mathfrak{g}(G))$.

Wegen Gleichung (4.3) aus 4.2.8 ist die Abbildung $\text{Ad}: G \rightarrow \text{Aut}(\mathfrak{g}(G))$ ein Homomorphismus und definiert daher eine Darstellung von G auf $\mathfrak{g}(G)$ per Automorphismen. Diese Darstellung heißt die *adjungierte Darstellung* (im Englischen *adjoint action*) von G auf $\mathfrak{g}(G)$; sie ist analytisch. Das folgende Diagramm kommutiert:

$$\begin{array}{ccc}
 G & \xrightarrow{\text{Ad}} & \text{Aut}(\mathfrak{g}(G)) \\
 \exp \uparrow & & \uparrow \exp \\
 \mathfrak{g}(G) & \xrightarrow{\text{ad}} & \text{aut}(\mathfrak{g}(G)),
 \end{array}$$

wobei $\text{aut}(\mathfrak{g}(G))$ sich auf die Definition 2.7.3 bezieht.

4.4 Lokal linearisierbare Abbildungen

4.4.1 (KAUP [180], Seite 25). Sei $g: M \rightarrow N$ eine analytische Abbildung zwischen Banach-Mannigfaltigkeiten über $\mathbb{K}$ und $m \in M$. Können die Karten (Ω, p, X) und (Σ, q, Y) mit $m \in \Omega$ und $g(m) \in \Sigma$ so gewählt werden, dass $q \circ g \circ p^{-1} = T|_p(\Omega)$ für ein $T \in \mathcal{L}(X, Y)$ gilt, so heißt die Abbildung g in m *lokal linearisierbar*. Zwei Spezialfälle solcher lokal linearisierbarer Abbildungen bekommen einen eigenen Namen:

4.4.2 Submersion (KAUP [180], Seite 25; UPMEIER [279](1985), Seite 121; BOURBAKI, Variétés, 5.9.1). Eine analytische Abbildung $g: M \rightarrow N$ zwischen zwei Banach-Mannigfaltigkeiten über $\mathbb{K}$ heißt eine *Submersion im Punkt* $m \in M$, falls eine der folgenden äquivalenten Aussagen zutrifft:

- (a) g ist in m lokal linearisierbar und mit den Bezeichnungen von 4.4.1 gilt: Die lokale Linearisierung T in m lässt sich so wählen, dass Y in X topologisch komplementiert ist, etwa $X = U \dot{+} Y$, und $T(u + y) = y$ für alle $u \in U$, $y \in Y$ gilt. (Man bemerke, dass T surjektiv ist.)

- (b) Es gibt offene Umgebungen Ω von m , Σ von $g(m)$ und eine analytische Abbildung $h: \Sigma \rightarrow \Omega$ mit $g(\Omega) \subseteq \Sigma$, $h(g(m)) = m$ und $g \circ h = \text{Id}_\Sigma$.
- (c) Es gibt eine offene Umgebung Σ von $g(m)$ und eine analytische Abbildung $h: \Sigma \rightarrow M$ mit $h(g(m)) = m$ und $g \circ h = \text{Id}_\Sigma$.
- (d) Das Differential $T_m(g)$ von g bei m ist surjektiv und dessen Kern $\ker T_m(g)$ ist im Tangentialraum $T_m(M)$ topologisch komplementiert.

4.4.3 Immersion (KAUP [180], Seite 25; UPMEIER [279](1985), Seite 121; BOURBAKI [33], 5.7.1). Eine analytische Abbildung $g: M \rightarrow N$ zwischen zwei Banach-Mannigfaltigkeiten über $\mathbb{K}$ heißt eine *Immersion im Punkt* $m \in M$, falls eine der folgenden äquivalenten Aussagen zutrifft:

- (a) g ist in m lokal linearisierbar und mit den Bezeichnungen von 4.4.1 gilt: Die lokale Linearisierung T in m lässt sich so wählen, dass X in Y topologisch komplementiert ist, etwa $Y = X \dot{+} V$, und $T(x) = x$ für alle $x \in X$ gilt. (Man bemerke, dass T injektiv ist.)
- (b) Es gibt offene Umgebungen Ω von m , Σ von $g(m)$ und eine analytische Abbildung $h: \Sigma \rightarrow \Omega$ mit $g(\Omega) \subseteq \Sigma$, $h(g(m)) = m$ und $h \circ (g|_\Omega) = \text{Id}_\Omega$.
- (c) Es gibt offene Umgebungen Ω von m , Σ von $g(m)$ und eine analytische Abbildung $h: \Sigma \rightarrow M$ mit $g(\Omega) \subseteq \Sigma$ und $h \circ (g|_\Omega) = \text{Id}_\Omega$.
- (d) Das Differential $T_m(g)$ von g bei m ist injektiv und dessen Bildraum $\text{ran}(T_m(g)) = T_m(g)T_m(M)$ ist im Tangentialraum $T_{g(m)}(N)$ topologisch komplementiert.

4.5 Beschränkte und symmetrische Gebiete

4.5.1 Tangentialnorm (UPMEIER [279](1985), Seite 200). Sei M eine Banach-Mannigfaltigkeit über $\mathbb{K}$. Eine (*stetige*) *Norm* auf M (auch genauer (*stetige*) *Tangentialnorm auf* $T(M)$) ist eine (stetige) Abbildung $\|\cdot\|: T(M) \rightarrow \mathbb{R}^+$, bei der für jedes $m \in M$ die Einschränkung $\|\cdot\|_m := \|\cdot\|_{T_m(M)}$ eine Norm auf $T_m(M)$ ist.

Eine analytische Abbildung $g: M \rightarrow N$ zwischen Banach-Mannigfaltigkeiten, die jeweils eine Norm, $\|\cdot\|_M$ beziehungsweise $\|\cdot\|_N$, tragen, heißt *Kontraktion*, wenn für alle $v \in T(M)$ die Ungleichung $\|T(g)v\|_N \leq \|v\|_M$ gilt; liegt hierbei für alle $v \in T(M)$ Gleichheit vor und ist g bianalytisch, so heißt g eine *Isometrie*.

Eine stetige Norm $\|\cdot\|$ auf M heißt *kompatibel*, wenn für jedes $a \in M$ eine Karte (Ω, p, X) von M um a und Konstanten $0 < c \leq C$ existieren, so dass gilt:

$$c\|T(p)v\| \leq \|v\| \leq C\|T(p)v\| \quad \text{für alle } v \in T(\Omega);$$

das heißt mit anderen Worten — da v von der Gestalt $\left(m, x \frac{\partial}{\partial p}\right)$ für $m \in \Omega$ und $x \in X$ ist —, so dass gilt:

$$c\|x\| \leq \left\|x \frac{\partial}{\partial p}\right\|_m \leq C\|x\| \quad \text{für alle } m \in \Omega, x \in X.$$

Eine Banach-Mannigfaltigkeit M über $\mathbb{K}$, versehen mit einer kompatiblen Tangentialnorm $\|\cdot\|$, wird als eine *normierte Banach-Mannigfaltigkeit* $(M, \|\cdot\|)$ bezeichnet. Die Gruppe aller bianalytischen Isometrien $g: M \rightarrow M$ wird mit $\text{Aut}_{MF}(M, \|\cdot\|)$ bezeichnet.

Sei $(M, \|\cdot\|)$ eine normierte Banach-Mannigfaltigkeit über $\mathbb{K}$. Dann ist für jede glatte Kurve $\gamma: [0, 1] \rightarrow M$ die *Bogenlänge bezüglich $\|\cdot\|$* definiert als $L(\gamma) := \int_{[0,1]} \|T_t(\gamma)\| dt$. Ist M außerdem zusammenhängend, dann definiert

$$d(m, n) := \inf \{L(\gamma) : \gamma \text{ eine stückweise glatte Kurve in } M \text{ von } m \text{ nach } n\},$$

$m, n \in M$, eine Metrik auf M , die invariant unter $\text{Aut}_{MF}(M, \|\cdot\|)$ ist.

4.5.2. Die Definition einer Norm in 4.5.1 ist gemäß Upmeyer; insbesondere seine Unterscheidung zwischen einer stetigen Norm und einer Norm. KAUP [173](1977) definiert eine Norm wie hier, aber mit dem Unterschied, dass er sie als unterhalbstetig voraussetzt.

In der Funktionentheorie einer komplexen Variable hat man den

4.5.3 Identitätssatz. *Sei $\Omega \subseteq \mathbb{C}$ ein Gebiet. Zwei holomorphe Funktionen $f, g: \Omega \rightarrow \mathbb{C}$ sind identisch, wenn sie auf einer Teilmenge von Ω , welche in Ω einen Häufungspunkt besitzt, übereinstimmen.*

und das

4.5.4 Lemma von Schwarz. *Sei $f: \Omega \rightarrow \Omega$ eine holomorphe Abbildung der offenen Einheitskreisscheibe Ω von $\mathbb{C}$ in sich. Gilt dann $f(0) = 0$ und $|Df(0)| = 1$, so ist f durch $Df(0)$ bestimmt, genauer ist dann f eine Drehung: $f(z) = \exp(i\theta)z$ für ein $\theta \in \mathbb{R}$ und für alle $z \in \Omega$.*

4.5.5 Korollar (HARRIS [127](1969)). *Seien X und Y zwei Banachräume über $\mathbb{C}$ und sei $f: \text{int}(B_X) \rightarrow \text{int}(B_Y)$ eine surjektive biholomorphe Abbildung zwischen den offenen Einheitskugeln von X und Y mit $f(0) = 0$. Dann ist $f = F \upharpoonright \text{int}(B_X)$ für ein surjektives isometrisches $F \in \mathcal{L}(X, Y)$.*

Sind X und Y zwei Banachräume über $\mathbb{C}$ und ist die offene Einheitskugel $\text{int}(B_X)$ von X homogen, so kann man jeder surjektiven biholomorphen Abbildung $f: \text{int}(B_X) \rightarrow \text{int}(B_Y)$ eine surjektive biholomorphe Abbildung $\varphi: \text{int}(B_X) \rightarrow \text{int}(B_X)$ voranstellen, so dass die Komposition $f \circ \varphi$ die $0 \in X$ auf die $0 \in Y$ abbildet. Somit:

4.5.6 Korollar (HARRIS [128](1974), Seite 15, Korollar 1). *Seien X und Y zwei Banachräume über $\mathbb{C}$ und sei die offene Einheitskugel $\text{int}(B_X)$ von X homogen. Dann sind die offenen Einheitskugeln $\text{int}(B_X)$ und $\text{int}(B_Y)$ genau dann biholomorph äquivalent, wenn $X \cong Y$ gilt.*

Für eine schärfere Fassung dieses Satzes siehe 4.5.32.

Als eine Verallgemeinerung des Lemmas von Schwarz 4.5.4 hat man den

4.5.7 Satz. *Sei Ω ein beschränktes Gebiet eines Banachraumes über $\mathbb{C}$. Seien $f, g \in \text{Aut}_{MF}(\Omega)$. Gelten für ein $m \in \Omega$ die beiden Gleichungen $f(m) = g(m)$ und $Df(m) = Dg(m)$, so gilt $f = g$.*

Beweis. Siehe zum Beispiel ISIDRO und STACHÓ [149](1984), Seite 19. □

Als Korollar hat man die Banachraumversion des ursprünglich von Henri Cartan 1931 für den $\mathbb{C}^n$, $n \in \mathbb{N}$, gezeigten Satzes:

4.5.8 Eindeutigkeits-Satz von Cartan (VIGUÉ [282](1976), Satz 1.2.1). *Sei Ω ein beschränktes Gebiet eines Banachraumes X über $\mathbb{C}$ und $f: \Omega \rightarrow \Omega$ holomorph. Wenn ein Fixpunkt $m \in \Omega$ von f mit $Df(m) = \text{Id}_X$ existiert, dann ist $f = \text{Id}_\Omega$.*

4.5.9. VIGUÉ [283](1982), Seite 53, Theorem 2.3, zeigt die Gültigkeit einer dem Eindeutigkeits-Satz 4.5.8 von Cartan entsprechenden Aussage für normierte Banach-Mannigfaltigkeiten über $\mathbb{C}$.

Als infinitesimale Version folgt das

4.5.10 Korollar (KAUP und UPMEIER [182](1977); [149], Seite 65). *Sei Ω ein beschränktes Gebiet eines Banachraumes Z über $\mathbb{C}$ und $m \in \Omega$. Gelten für ein Vektorfeld $X = h \frac{\partial}{\partial z} \in \text{aut}(\Omega)$ die beiden Gleichungen $h(m) = 0$ und $Dh(m) = 0 \in \mathcal{L}(Z)$, so gilt $X = 0$.*

Für beschränkte Gebiete Ω eines Banachraumes Z über $\mathbb{C}$ sind also sowohl die $f \in \text{Aut}_{MF}(\Omega)$ als auch die analytischen $h: \Omega \rightarrow Z$ der $X = h \frac{\partial}{\partial z} \in \text{aut}(\Omega)$ eindeutig durch ihren Wert und ihrer ersten Ableitung an einem beliebig gegebenen Punkt $m \in \Omega$ bestimmt.

4.5.11 Korollar (VIGUÉ [282](1976), Theorem 1.2.3; ISIDRO und STACHÓ [149] (1984); UPMEIER [279](1985), Seite 220, Korollar 13.23). *Sei Ω ein beschränktes Gebiet eines Banachraumes Z über $\mathbb{C}$ und $m \in \Omega$. Dann ist die Abbildung*

$$\text{aut}(\Omega) \rightarrow Z \times \mathcal{L}(Z), \quad h \frac{\partial}{\partial z} \mapsto (h(m), Dh(m))$$

injektiv und stetig.

4.5.12 Korollar (UPMEIER [280](1987), Seite 12). *Seien Ω und Σ zwei beschränkte, zirkuläre Gebiete eines Banachraumes X über $\mathbb{C}$ und sei $f: \Omega \rightarrow \Sigma$ eine biholomorphe Abbildung mit 0 als Fixpunkt. Dann ist $f = F \upharpoonright \Omega$ für ein $F \in \mathcal{L}(X)$.*

4.5.13 Definition. Sei M ein topologischer Raum und $T: M \rightarrow M$ eine Selbstabbildung von M . Sei F die Menge der Fixpunkte von T , also $F = \{x \in M : Tx = x\}$. Dann heißt ein Fixpunkt $x \in F$ von T ein *isolierter Fixpunkt* von T , wenn er ein isolierter Punkt von F ist, wenn es also eine Umgebung U von x in M gibt mit $F \cap U = \{x\}$.

4.5.14 Definition (UPMEIER [279](1985), Seiten 281 und 343). Sei M eine Banach-Mannigfaltigkeit über $\mathbb{K}$. Eine *Symmetrie* von M in $a \in M$ ist ein $s \in \text{Aut}_{MF}(M)$ mit Periode 2 und a als isolierten Fixpunkt.

Eine zusammenhängende, normierte Banach-Mannigfaltigkeit M über $\mathbb{C}$ heißt *symmetrisch*, wenn in jedem $a \in M$ eine Symmetrie s_a existiert.

4.5.15. HARRIS [128](1974), Seite 14, nennt ein Gebiet eines Banachraumes über $\mathbb{C}$ symmetrisch, wenn für alle $a \in \Omega$ eine biholomorphe Abbildung $h: \Omega \rightarrow \Omega$ mit Periode 2 existiert derart, dass a der einzige Fixpunkt von h ist. Offensichtlich ist ein solches Gebiet auch symmetrisch im Sinne von Definition 4.5.14.

4.5.16 (UPMEIER [279](1985), Seite 281, Lemma 17.3). Sei M eine Banach-Mannigfaltigkeit über $\mathbb{K}$, $a \in M$ und s_a eine Symmetrie in a . Dann gilt $T_a s_a = -\text{Id}_{T_a(M)}$. Ist also insbesondere X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes und symmetrisches Gebiet in X , so gilt für jedes $a \in \Omega$ die Gleichung $Ds_a(a) = -\text{Id}_X$. Nach dem Eindeutigkeits-Satz 4.5.7 ist somit für jedes $a \in \Omega$ die Symmetrie s_a in a eindeutig bestimmt. Diese Eindeutigkeit gilt auch allgemeiner für symmetrische Banach-Mannigfaltigkeiten über $\mathbb{C}$; sie folgt letztlich wieder aus dem Eindeutigkeits-Satz 4.5.8 von Cartan, siehe VIGUÉ [283](1982), Seite 67, Satz 4.1.

VIGUÉ [282](1976), Seite 249, Satz 3.1.1, zeigt: Sei X ein Banachraum über $\mathbb{C}$, Ω ein beschränktes Gebiet in X , $a \in \Omega$, $j \in \text{Aut}_{MF}(\Omega)$ mit a als Fixpunkt. Dann ist j genau dann eine Symmetrie in a , wenn $Dj(a) = -\text{Id}_X$ gilt.

4.5.17 (KAUP [178](2002); UPMEIER [280](1987), Seite 14; UPMEIER [279] (1985), Seite 345; MCCRIMMON [210], Seite 26). Sei X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes und symmetrisches Gebiet in X . Die Abbildung $\Omega \rightarrow \text{Aut}_{MF}(\Omega)$, die jedem $a \in \Omega$ die Symmetrie s_a in a zuordnet, ist $\mathbb{R}$ -analytisch. Da jedes $g \in \text{Aut}_{MF}(\Omega)$ die Symmetrie s_a in $a \in \Omega$ zu einer Symmetrie $\text{Int}(g)(s_a) = g \circ s_a \circ g^{-1}$ in $g(a)$ konjugiert, ist ein beschränktes Gebiet Ω eines Banachraumes X über $\mathbb{C}$ genau dann symmetrisch, wenn es einerseits einen Punkt $x \in \Omega$ gibt, in dem eine Symmetrie $j \in \text{Aut}_{MF}(\Omega)$ existiert und andererseits Ω homogen ist; x heißt dann ein *Basispunkt* von Ω .

4.5.18. Die von KAUP [173](1977) verwendete Definition, wann eine zusammenhängende, normierte Banach-Mannigfaltigkeit $(M, \|\cdot\|)$ über $\mathbb{C}$ als symmetrisch zu bezeichnen ist, lautet wie folgt: Bezeichne $\text{Aut}_{MF}(M, \|\cdot\|)$ die Gruppe aller biholomorphen Isometrien $g: M \rightarrow M$. (Man beachte den Unterschied zu der Bezeichnung der analytischen Automorphismengruppe $\text{Aut}_{MF}(M)$.) Upmeier zeigte 1975, dass eine natürlich aus der Theorie der Lie-Gruppen entspringende Topologie auf $\text{Aut}_{MF}(M, \|\cdot\|)$ existiert, mit der $\text{Aut}_{MF}(M, \|\cdot\|)$ eine Banach-Lie-Gruppe G über $\mathbb{R}$ ist; siehe etwa UPMEIER [279](1985), Korollar 13.17, oder auch ISIDRO und STACHÓ [149](1984), Seite 109, Theorem 6.57, wo sie als die *analytische Topologie* bezeichnet wird. (Für $M = B_X$, X ein Banachraum über $\mathbb{C}$, siehe ARAZY [9](1987), Seite 133, für eine Definition der analytischen Topologie auf $\text{Aut}_{MF}(B_X)$.) Sei $\text{Aut}_{MF}(M, \|\cdot\|)$ mit dieser Topologie ausgestattet. Dann heißt M *symmetrisch*, wenn es ein $a \in M$ gibt, so dass einerseits ein $s \in \text{Aut}_{MF}(M, \|\cdot\|)$ mit Periode 2 und a als isolierten Fixpunkt existiert und andererseits die Auswertungsabbildung $G \rightarrow M$, $g \mapsto g(a)$, eine Submersion ist.

Da die letzte der beiden Bedingungen impliziert, dass die Gruppe $\text{Aut}_{MF}(M, \|\cdot\|)$ transitiv auf M operiert — für einen Beweis siehe OUCHEVA [228](2001), Seite 18, Satz 3.12 —, ist eine nach dieser Definition symmetrische, zusammenhängende Banach-Mannigfaltigkeit über $\mathbb{C}$ auch symmetrisch gemäß der Definition 4.5.14. Es war bekannt, dass die Umkehrung im endlichdimensionalen Fall gilt und Vigué zeigte 1976, dass die Umkehrung auch für den Fall von beschränkten Gebieten von Banachräumen über $\mathbb{C}$ gilt. In VIGUÉ [283](1982), Seite 72, wird schließlich gezeigt, dass die hier in 4.5.14 gegebene Definition einer symmetrischen Banach-Mannigfaltigkeit über $\mathbb{C}$ mit der Definition in KAUP [173](1977) äquivalent ist.

4.5.19 (KAUP [173](1977)). Die symmetrischen Banach-Mannigfaltigkeiten über $\mathbb{C}$ bilden eine Kategorie; dabei heißt eine holomorphe Abbildung $g: M \rightarrow N$ zwischen zwei symmetrischen Banach-Mannigfaltigkeiten über $\mathbb{C}$ ein *Morphismus*, wenn gilt:

$$g \circ s_a = s_{g(a)} \circ g \quad \text{für alle } a \in M.$$

Ebenso bilden die symmetrischen Banach-Mannigfaltigkeiten über $\mathbb{C}$ mit einem als Basispunkt bezeichneten festgewählten Punkt $a \in M$ eine Kategorie, wobei die eben erwähnten holomorphen Abbildungen g jetzt noch zusätzlich die Basispunkte aufeinander abbilden.

In der Funktionentheorie einer komplexen Variablen hat man den

4.5.20 Riemann'schen Abbildungssatz. *Jedes nicht leere, einfach zusammenhängende Gebiet in $\mathbb{C}$, welches nicht $\mathbb{C}$ ist, ist biholomorph äquivalent zu der offenen Einheitskreisscheibe in $\mathbb{C}$.*

4.5.21 (KAUP und KAUP [172](1983)). Sei $(r_1, \dots, r_n) \in \mathbb{R}^{+\times n}$ und $(a_1, \dots, a_n) \in \mathbb{C}^n$ für ein $n \in \mathbb{N}^\times$. Setze $r := (r_1, \dots, r_n)$ und $a := (a_1, \dots, a_n)$. Dann heißt $\text{Zyl}(a, r) := \text{Zyl}^n(a, r) := \{(z_1, \dots, z_n) \in \mathbb{C}^n : |z_j - a_j| < r_j, 1 \leq j \leq n\}$ eine (n -dimensionale) *offene Polyscheibe* (im Englischen *polydisk*) oder auch ein (n -dimensionaler) *offener Polyzylinder* (im Englischen *polycylinder*) um a mit *Polyradius* r .

Der (n -dimensionale) *offene Einheitspolyzylinder* ist als $\text{Zyl} := \text{Zyl}^n := \text{Zyl}^n(0, (1, \dots, 1))$ definiert; also $\text{Zyl}^n = (\text{int} B_{\mathbb{C}})^n$, die Einheitskugel des mit der Maximumsnorm versehenen $\mathbb{C}^n$. Im Fall von $n = 2$ wird anstelle von *Poly* das Wort *Di* verwendet. Im Englischen nennt man den offenen Dizylinder auch den offenen *biball*. Für jedes $n \in \mathbb{N}^\times$ ist Zyl^n homogen.

Sind $f_i: M_i \rightarrow N_i$, $i = 1, 2$, zwei holomorphe Abbildungen zwischen offenen Teilmengen von $\mathbb{C}$, so ist auch die Abbildung $f_1 \times f_2: M_1 \times M_2 \rightarrow N_1 \times N_2$, $(x_1, x_2) \mapsto (f_1(x_1), f_2(x_2))$ holomorph. Als eine Folgerung daraus und aus dem Riemann'schen Abbildungssatz 4.5.20 hat man die folgende Verallgemeinerung auf den n -dimensionalen Fall:

4.5.22 Riemann'scher Abbildungssatz II (KAUP und KAUP [172](1983)). *Sei $G := G_1 \times \dots \times G_n$ ein Produkt von einfach zusammenhängenden, nicht leeren Gebieten $G_i \subseteq \mathbb{C}$, $i = 1, \dots, n$. Dann existiert ein $k \in \mathbb{N}$, so dass G biholomorph äquivalent ist zu dem Produkt $\text{Zyl}^k \times \mathbb{C}^{n-k}$.*

4.5.23. Während der Riemann'sche Abbildungssatz 4.5.20 für $n = 1$ aussagt, dass alle nicht leeren, von $\mathbb{C}^n$ verschiedenen, einfach zusammenhängenden Gebiete biholomorph äquivalent sind, gilt dies nicht mehr für $n \geq 2$.

Seien Ω und Σ zwei Gebiete eines Banachraumes über $\mathbb{C}$ und sei φ eine biholomorphe Abbildung $\Omega \rightarrow \Sigma$. Dann gilt

$$\text{Aut}(\Sigma) = \{\varphi \circ g \circ \varphi^{-1} : g \in \text{Aut}(\Omega)\}; \quad (4.7)$$

betrachte dazu nur $\varphi g \varphi^{-1} \Sigma = \varphi g \Omega = \varphi \Omega = \Sigma$ und entsprechend für Ω . Damit also zwei Gebiete eines Banachraumes über $\mathbb{C}$ biholomorph äquivalent sein können, ist es notwendig, dass ihre Automorphismengruppen isomorph sind.

Poincaré [242](1907) bestimmt explizit (ebd., Seiten 207-212) die Gruppe der Abbildungen, die die Einheitssphäre des $\mathbb{C}^2$ bijektiv auf sich abbilden und auf einem Gebiet, das diese Einheitssphäre enthält, holomorph sind. Es wird ein Theorem von Hartogs verwendet (ebd., Seite 213), womit gezeigt wird, dass die besagte Gruppe gleich der Automorphismengruppe $\text{Aut}(\text{Zyl}^2)$ des offenen Dizylinders des $\mathbb{C}^2$ ist. Da diese nicht isomorph zu der Automorphismengruppe $\text{Aut}(\text{int } B_{\mathbb{C}^2})$ der offenen Einheitskugel des $\mathbb{C}^2$ ist (ebd., Seite 190), hat somit Poincaré gezeigt, dass die offene Einheitskugel des $\mathbb{C}^2$ nicht biholomorph äquivalent zu dem offenen Dizylinder des $\mathbb{C}^2$ ist. Siehe auch NARASIMHAN [222](1971), Seite 70, und DINEEN [69](1981), Seiten 389 und 390.

Allgemein gilt: Ist für ein $n \in \mathbb{N}$ eine Abbildung $\varphi: \text{int } B_{\mathbb{C}^n} \rightarrow \text{Zyl}^n$ biholomorph, so ist $n = 1$. Denn: Da Zyl^n homogen ist, kann man $\varphi(0) = 0$ annehmen. Nach dem Korollar 4.5.12 kann man φ als eine lineare Abbildung des $\mathbb{C}^n$ annehmen. Insbesondere ist $\varphi(S_{\mathbb{C}^n}) = \partial(\text{Zyl}^n)$. Angenommen, es wäre $n \geq 2$. Dann enthält $\partial(\text{Zyl}^n)$ das Segment $\{(1, t, 0, \dots, 0) \in \mathbb{C}^n : -1 \leq t \leq 1\}$. Da aber φ^{-1} auch linear ist, müsste $S_{\mathbb{C}^n} = \varphi^{-1}(\partial(\text{Zyl}^n))$ ebenfalls ein Segment enthalten, was aber im Widerspruch dazu steht, dass dies für $n \geq 2$ nicht der Fall ist.

(KRANTZ [189](1982)). Eine andere Möglichkeit dies zu beweisen verallgemeinert die von Poincaré verwendete Idee: Man versehe die Automorphismengruppen für Gebiete des $\mathbb{C}^n$ mit der Topologie der gleichmäßigen Konvergenz auf Kompakta (sie ist hier gleich der kompakt-offen-Topologie, siehe zum Beispiel KELLEY [184](1955)). (Siehe auch NARASIMHAN [222](1971), Seite 5.) Ist G eine topologische Gruppe, so bezeichne G_0 die Zusammenhangskomponente von G , die das neutrale Element enthält. Seien nun Ω und Σ zwei Gebiete des $\mathbb{C}^n$ und sei φ eine biholomorphe Abbildung $\Omega \rightarrow \Sigma$. In (4.7) hat man gesehen, dass $\omega_\varphi: g \mapsto \varphi g \varphi^{-1}$ ein Gruppenisomorphismus von $\text{Aut}(\Omega)$ nach $\text{Aut}(\Sigma)$ ist. ω_φ und ω_φ^{-1} sind stetig. Bezeichnet man nun für $p \in \Omega$ mit K_p die Untergruppe von Elementen aus $\text{Aut}(\Omega)$, die p als Fixpunkt haben — also die *Isotropie-Untergruppe*

in p (natürlich bezüglich der kanonischen Operation der Gruppe $\text{Aut}(\Omega)$) —, so bildet mit $q := \varphi(p)$ die Abbildung ω_φ die Untergruppe K_p isomorph auf K_q ab. Ebenfalls wird $(\text{Aut}(\Omega))_0$ isomorph auf $(\text{Aut}(\Sigma))_0$ abgebildet. $(K_0)_0 \subseteq \text{Aut}(\text{Zyl}^n)$, also die Zusammenhangskomponente der Isotropie-Untergruppe von $\text{Aut}(\text{Zyl}^n)$ in 0, die die Identität enthält, ist gleich $\{(z_1, \dots, z_n) \mapsto (\exp(i\theta_1)z_1, \dots, \exp(i\theta_n)z_n) : 0 \leq \theta_1, \dots, \theta_n \leq 2\pi\}$, insbesondere ist dieses K_0 also eine abelsche Gruppe. $(K_0)_0 \subseteq \text{Aut}(\text{int } B_{\mathbb{C}^n})$ enthält alle unitären Automorphismen und daher ist dieses K_0 nicht abelsch.

(PINCHUK [240](1989), Seite 180). Es gilt: Im $\mathbb{C}^n$ gibt es für $n < 7$ bis auf biholomorphe Äquivalenz nur eine endliche Anzahl von verschiedenen homogenen Gebieten, aber für $n \geq 7$ hat die Menge aller solcher Äquivalenzklassen die Mächtigkeit des Kontinuums.

Um für Gebiete im $\mathbb{C}^n$, $n \in \mathbb{N}$, $n > 1$, Aussagen analog zum Riemann'schen Abbildungssatz zu erhalten, muss man gewisse Zusatzannahmen an die Gebiete stellen. Für $n = 2$ hat Élie Cartan per Automorphismengruppen gezeigt:

4.5.24 Satz. *Jedes beschränkte homogene Gebiet in $\mathbb{C}^2$ ist biholomorph abbildbar entweder auf die offene Einheitskugel des $\mathbb{C}^2$ oder auf den Dizylinder $\{(w, z) \in \mathbb{C}^2 : |w| < 1, |z| < 1\}$.*

Beweis. CARTAN [44](1935), Seite 160. □

4.5.25. Nach dem Einbettungssatz von Harish-Chandra (siehe zum Beispiel ISIDRO und STACHÓ [149](1984), Seite 192, wo auch das diesem Satz vorausgegangene Klassifizierungstheorem von É. Cartan, 1935, zu finden ist) ist jedes beschränkte, symmetrische Gebiet Ω eines Vektorraumes über $\mathbb{C}$, welches endlichdimensional ist, soll heißen, $\Omega \subseteq \mathbb{C}^n$ für ein $n \in \mathbb{N}$, biholomorph äquivalent zu einem beschränkten zirkulären Gebiet des $\mathbb{C}^n$. Diese sogenannte Harish-Chandra-Realisierung ist immer homogen und konvex, das heißt, bezüglich einer passenden Norm des $\mathbb{C}^n$ ist sie gleich der offenen Einheitskugel des $\mathbb{C}^n$. (KAUP [175](1983); MCCRIMMON [210](2004), Seite 22.)

4.5.26. Eine *Riemann'sche* (bzw. *hermitesche*) Metrik auf einer C^∞ -Mannigfaltigkeit M über $\mathbb{R}$ (bzw. $\mathbb{C}$) ist ein glatt variierendes (im Englischen *smoothly varying*) Skalarprodukt $\langle \cdot, \cdot \rangle_a$ auf dem Tangentialraum $T_a(M)$ bei jedem $a \in M$.

4.5.27 (KAUP [178](2002)). Sei X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes und symmetrisches Gebiet in X . Während im endlichdimensionalen Fall (also $\Omega \subseteq \mathbb{C}^n$ für ein $n \in \mathbb{N}$) immer eine $\text{Aut}_{MF}(\Omega)$ -invariante hermitesche Metrik auf Ω existiert, zum Beispiel die Bergmann-Metrik (HELGASON [131], Seite 298), kann im unendlichdimensionalen Fall X im Allgemeinen nicht zu einem Hilbertraum umgenormt werden; es gibt also im Allgemeinen keine hermitesche Metrik auf Ω .

Anstelle einer Hilbertraum-Norm gibt es aber die folgende kanonische Norm:

4.5.28 (UPMEIER [279](1985), Seite 202, Satz 12.23). Sei X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes Gebiet in X . Dann ist per

$$\|v\| := \sup \{ \|T_a(f)v\| : f: \Omega \rightarrow \text{int}(B_{\mathbb{C}}) \text{ holomorph} \} \text{ für alle } v \in T_a(\Omega), a \in \Omega,$$

eine $\text{Aut}_{MF}(\Omega)$ -invariante Banachraum-Norm auf Ω definiert, die als die *Carathéodory-Norm* (bezüglich a) bezeichnet wird. Diese Definition einer Norm formuliert sich ganz entsprechend auch allgemeiner für den Fall, dass man Ω direkt als eine Banach-Mannigfaltigkeit M über $\mathbb{C}$ vorliegen hat (also nicht zwingend als eine Teilmenge eines Banachraumes) (siehe ebd.) und bezüglich dieser somit erklärten Norm ist jede holomorphe Abbildung $f: M \rightarrow N$ zwischen Banach-Mannigfaltigkeiten über $\mathbb{C}$ eine Kontraktion; insbesondere ist jede biholomorphe Abbildung eine Isometrie!

Somit ist $\text{aut}(\Omega)$ nach UPMEIER [279](1985), Seite 219, Korollar 13.17, eine Unteralgebra über $\mathbb{R}$ der Lie-Algebra $\mathcal{T}(\Omega)_{\mathbb{R}}$ und $\text{aut}(\Omega)$ ist bezüglich einer passenden Norm eine Banach-Lie-Algebra. Die in 4.2.11 erklärte Abbildung (4.4) $\exp: \text{aut}(\Omega) \rightarrow \text{Aut}_{MF}(\Omega)$ induziert auf $\text{Aut}_{MF}(\Omega)$ die Struktur einer Banach-Mannigfaltigkeit, mit der $\text{Aut}_{MF}(\Omega)$ eine Banach-Lie-Gruppe über $\mathbb{R}$ wird (siehe die dazu in 4.5.18 angegebenen Literaturhinweise), deren Lie-Algebra $\mathfrak{g}(\text{Aut}_{MF}(\Omega))$ über das Differential der Lie-Gruppen-Darstellung (siehe 4.3.8) $r = \text{Id}_{\text{Aut}_{MF}(\Omega)}$ mit ganz $\text{aut}(\Omega)$ identifiziert werden kann, womit die Lie-Algebra der Banach-Lie-Gruppe $\text{Aut}_{MF}(\Omega)$ modulo Identifikation konkret angegeben ist. $\text{Aut}_{MF}(\Omega)$ operiert auf Ω reell analytisch, das heißt, die Abbildung $\text{Aut}_{MF}(\Omega) \times \Omega \rightarrow \Omega$, $(f, x) \mapsto f(x)$ ist reell-analytisch (ISIDRO und STACHÓ [149](1984), Seite 111, Theorem 6.58).

Man beachte, dass zwar nach Cartan für endlichdimensionale Banachräume über $\mathbb{C}$ die Gruppe $\text{Aut}_{MF}(\Omega)$, versehen mit der Topologie der gleichmäßigen Konvergenz auf Kompakta — die in diesem Fall mit der Topologie der sogenannten lokalen gleichmäßigen Konvergenz übereinstimmt —, für jedes beschränkte Gebiet $\Omega \subseteq X$ eine Banach-Lie-Gruppe über $\mathbb{R}$ ist, dass dies aber für unendlichdimensionale Banachräume über $\mathbb{C}$ im Allgemeinen nicht stimmt. Zum Beispiel ist zwar dann $\text{Aut}_{MF}(\Omega)$ in der Topologie der sogenannten lokalen gleichmäßigen Konvergenz eine topologische Gruppe, aber im Allgemeinen nicht vollständig (ISIDRO und STACHÓ [149](1984), Seite 39). Siehe auch VIGUÉ [282](1976), Seiten 203 und 229; insbesondere zeigt Vigué dort mit seinen beiden Theoremen 2.4.7 und 2.4.8 zwei konkrete Beispiele, in denen $\text{Aut}_{MF}(\Omega)$ keine Lie-Gruppe ist.

4.5.29. Sei Z ein Banachraum über $\mathbb{C}$ und Ω ein zirkuläres, beschränktes Gebiet in Z . Dann lässt sich die Lie-Algebra von $\text{Aut}_{MF}(\Omega)$ noch genauer spezifizieren: Es gilt

$$\text{aut}(\Omega) \subseteq \mathcal{T}_{-1}(Z) \oplus \mathcal{T}_0(Z) \oplus \mathcal{T}_1(Z),$$

das heißt, jedes Vektorfeld $X \in \text{aut}(\Omega)$ ist polynomial vom Grad höchstens 2. Es existiert eine sogenannte Cartan-Zerlegung $\text{aut}(\Omega) = \mathfrak{p} \oplus \mathfrak{k}$, wobei $\mathfrak{p} := \text{aut}(\Omega) \cap (\mathcal{T}_{-1}(Z) \oplus \mathcal{T}_1(Z))$ der abgeschlossene Unterraum aller sogenannten *infinitesimalen Transvektionen* und $\mathfrak{k} := \text{aut}(\Omega) \cap \mathcal{T}_0(Z)$ die abgeschlossene Unteralgebra aller sogenannten *infinitesimalen Drehungen* ist. (KAUP und UPMEIER [181](1976): $\text{aut}(B_Z) \cap \mathcal{T}_{-1}(Z) = \text{aut}(B_Z) \cap \mathcal{T}_1(Z) = \{0\}$.)

Für jedes Vektorfeld $X = h \frac{\partial}{\partial z} \in \mathfrak{p}$ ist die Abbildung h von der Form $z \mapsto a + q(z)$ mit $q(z) = b(z, z)$ für alle $z \in Z$ mit von X abhängigen $a \in Z$ und $b := \tilde{q} \in \mathcal{L}^2(Z)$. (Bezüglich der eben verwendeten Schreibweise mit der aufgesetzten Tilde sei an 2.2.40 erinnert.) Daher ist nach Korollar 4.5.10 die Auswertungsabbildung $\mathfrak{p} \rightarrow T_0(\Omega)$, $X \mapsto X(0) := X_0 = h(0) \frac{\partial}{\partial z}$ injektiv. Das heißt, jedem Element von $\mathfrak{p}$ entspricht genau einem Element aus dem kanonisch eingebetteten Unterraum $\mathfrak{p}(0) = \{X(0) : X \in \mathfrak{p}\}$ über $\mathbb{R}$ des Vektorraumes $Z_{\mathbb{R}}$. Dass dieser Raum abgeschlossen ist, wird in KAUP und UPMEIER [181](1976), Seite 130, Lemma 2 gezeigt; man beachte dazu auch die dortige Bemerkung auf Seite 132. Somit existiert eine Abbildung $q: \mathfrak{p}(0) \rightarrow \mathcal{P}^2(Z)$, $a \mapsto q_a$, so dass

$$\mathfrak{p} = \left\{ \left(z \mapsto a - q_a(z) \right) \frac{\partial}{\partial z} : a \in \mathfrak{p}(0) \right\}$$

gilt. (Dabei wurde hier das h in $h \frac{\partial}{\partial z} \in \mathfrak{p}$ in der Form $z \mapsto a - q(z)$ mit $q(z) = -b(z, z)$ für alle $z \in Z$ mit $a \in Z$, $b \in \mathcal{L}^2(Z)$, geschrieben.) Die Abbildung q ist stetig. $\mathfrak{p}(0)_{\mathbb{C}}$ ist ein abgeschlossener Unterraum über $\mathbb{C}$ von Z , der wieder mit $\mathfrak{p}(0)$ bezeichnet wird, und q ist konjugiert-linear; speziell für $\Omega = \text{int}(B_Z)$ heißt dieser Unterraum der *symmetrische*

Teil des Banachraumes Z und wird mit Z_{sym} bezeichnet. $Z_{\text{sym}} \cap \text{int}(B_Z)$ heißt der *symmetrische Teil* der offenen Einheitskugel des Banachraumes Z .

Es gilt $\mathfrak{p} = \left\{ h(z) \frac{\partial}{\partial z} \in \text{aut}(\Omega) : Dh(0) = 0 \right\}$. Setze $G := \{h \in GL(Z) : h(\Omega) = \Omega\}$, also $G \subseteq \text{Aut}_{MF}(\Omega)$ in kanonischer Weise. Dann gilt $\mathfrak{k} = \left\{ h(z) \frac{\partial}{\partial z} \in \text{aut}(\Omega) : h(0) = 0 \right\} = \left\{ T \in \mathcal{L}(Z) : \exp(\mathbb{R}T) \subseteq G \right\}$. Dabei drückt man das erste Gleichheitszeichen in Worten aus, indem man sagt, dass $\mathfrak{k}$ die *Isotropie-Unteralgebra in 0* von $\text{aut}(\Omega)$ sei. (KAUP [180], Seite 44, Satz 16.5.)

Der folgende Satz veranschaulicht, was der symmetrische Teil der offenen Einheitskugel eines Banachraumes über $\mathbb{C}$ ist.

4.5.30 Satz. *Sei X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes, zirkuläres Gebiet in X . Ist $\mathfrak{p}(0) \cap \Omega$ zusammenhängend (also insbesondere, wenn Ω kreisförmig ist), so gilt*

$$\text{Aut}_{MF}(\Omega)(0) = \mathfrak{p}(0) \cap \Omega.$$

Beweis. Siehe zum Beispiel ISIDRO und STACHÓ [149](1984), Seite 125. $\square$

4.5.31. Als ein Korollar von Satz 4.5.30 folgt, dass die Bahn des Ursprungs unter der Gruppe $\text{Aut}_{MF}(\text{int}(B_X))$ unter einer surjektiven, biholomorphen Abbildung $f: \text{int}(B_X) \rightarrow \text{int}(B_Y)$, X und Y Banachräume über $\mathbb{C}$, erhalten bleibt, das heißt, es gilt $f\left(\text{Aut}_{MF}(\text{int}(B_X))(0)\right) = \text{Aut}_{MF}(\text{int}(B_Y))(0)$. (Zur Begründung siehe die Anmerkung zu: ARAZY [9](1987), Seite 134, Hauptlemma 4.2.) Daraus wiederum folgt mit Korollar 4.5.5, dass dessen Korollar 4.5.6 auch ohne die Voraussetzung der Homogenität gilt; also:

4.5.32 Satz (KAUP und UPMEIER [181](1976)). *Zwei Banachräume über $\mathbb{C}$ sind genau dann linear isometrisch isomorph, wenn es eine surjektive, biholomorphe Abbildung zwischen ihren offenen Einheitskugeln gibt.*

4.5.33. Sei X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes, symmetrisches und zirkuläres Gebiet in X . Nach 4.5.17 ist Ω homogen und folglich gilt dann mit Satz 4.5.30, dass $\mathfrak{p}(0)$ gleich ganz X sein muss. Somit gilt $\mathfrak{p} \cong X_{\mathbb{R}}$. Bezüglich der Banachräume über $\mathbb{C}$, die über diese Eigenschaft charakterisiert sind, siehe 4.7.24.

4.5.34 (KAUP [179](2007), Seite 49). Sei X ein Banachraum über $\mathbb{C}$, Ω ein beschränktes Gebiet in X und $a \in \Omega$. Dann sind äquivalent:

- (a) Ω ist homogen. (Definition 4.3.6)
- (b) $\text{Aut}(\Omega)(a)$ ist offen in Ω .
- (c) Die Auswertungsabbildung $\text{aut}(\Omega) \rightarrow X$, $f \frac{\partial}{\partial z} \mapsto f(a)$ ist surjektiv.

4.5.35 (ISIDRO [146](1989)). Nach VIGUÉ (1976) gilt folgende Banachraum-Version der sogenannten *Harish-Chandra-Realisierung* für beschränkte, symmetrische Gebiete des $\mathbb{C}^n$:

4.5.36 Satz. *Sei X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes, symmetrisches Gebiet in X . Dann existiert ein kreisförmiges, beschränktes, symmetrisches Gebiet Σ in X und eine surjektive biholomorphe Abbildung $\varphi: \Omega \rightarrow \Sigma$.*

Beweis. Einen elementaren Beweis findet man bei ISIDRO und STACHÓ [149] (1984): Die Definition von φ und Σ : Seite 221, Definition 9.34; die Beschränktheit von Σ : Seite 221, Lemma 9.35; Σ zusammenhängend, offen und zirkulär: Seite 223, Satz 9.38; Σ homogen und kreisförmig: Seite 226, Korollar 9.41. $\square$

Man kann also annehmen, dass jedes beschränkte symmetrische Gebiet eines Banachraumes über $\mathbb{C}$ als ein zirkuläres Gebiet realisiert ist.

4.5.37. Sei X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes, symmetrisches und zirkuläres Gebiet in X . Dann ist gemäß 4.5.33 die in 4.5.29 definierte Abbildung q auf ganz X erklärt. Somit ist es möglich, die folgende Abbildung zu definieren:

$$X \times X \times X \rightarrow X, (x, y, z) \mapsto \tilde{q}_y(x, z).$$

Ausgestattet mit dieser als ein Produkt interpretierbaren Abbildung, ist X ein sogenanntes, im Abschnitt 4.7 definiertes JB^* -Tripel über $\mathbb{C}$; siehe auch 4.7.24.

4.5.38 (KAUP [178](2002)). Sei X ein Banachraum über $\mathbb{C}$ und Ω ein beschränktes und symmetrisches Gebiet in X . Betrachte Ω als mit seiner Carathéodory-Norm versehen. Jeder Tangentialraum $T_a(\Omega)$, $a \in \Omega$, ist homöomorph zu X und wegen der $\text{Aut}_{MF}(\Omega)$ -Invarianz der Carathéodory-Norm kann für jeden beliebigen Basispunkt $p \in \Omega$ jeder Tangentialraum $T_a(\Omega)$, $a \in \Omega$, mit $T_p(\Omega)$ identifiziert werden. Somit kann man ohne Einschränkung annehmen, dass X der Tangentialraum $T_p(\Omega)$ von Ω am Basispunkt p ist, und dass seine Norm die Carathéodory-Norm von Ω ist.

Nach KAUP [175](1983) gilt:

4.5.39 Satz. *Sei X ein Banachraum über $\mathbb{C}$, Ω ein beschränktes, symmetrisches Gebiet in X und $p \in \Omega$. Dann existiert eine biholomorphe Abbildung φ von Ω auf die offene Einheitskugel des Tangentialraumes $X = T_p(\Omega)$ mit $\varphi(p) = 0$, und je zwei solcher Abbildungen unterscheiden sich in einer surjektiven linearen Isometrie von X .*

4.5.40. Man kann also jedes beschränkte, symmetrische Gebiet Ω eines Banachraumes über $\mathbb{C}$ als die offene Einheitskugel eines Banachraumes X über $\mathbb{C}$ realisiert auffassen, wobei X bis auf lineare Isometrie eindeutig durch Ω bestimmt ist; dabei kann man dann $0 \in \Omega$ als Basispunkt annehmen.

4.6 Tripelsysteme

4.6.1 Definition. Ein *lineares Tripel* (oder auch *trilineares Tripel*, *lineares Tripelsystem*) über $\mathbb{K}$ ist ein Vektorraum X über $\mathbb{K}$ mit einer $\mathbb{K}$ -trilinearen Abbildung $X^3 \rightarrow X$.

Ein *Tripel* (oder auch ein *Tripelsystem*) über $\mathbb{K}$ ist ein Vektorraum X über $\mathbb{K}$ mit einer Abbildung $X^3 \rightarrow X$, die linear in den beiden äußeren Variablen und semilinear in der inneren Variable ist. Diese Art der Linearität heißt nach dem lateinischen Wort *sesterti* für zweieinhalb *sestertilinear*.

Diese Abbildungen werden als $\{ \} : (x, y, z) \mapsto \{x, y, z\}$ geschrieben, wobei die Kommas in $\{x, y, z\}$ weggelassen werden, wenn keine Missverständnisse möglich sind. Die Tripel und linearen Tripel werden aufgefasst als mit einem ternären Produkt versehene Vektorräume.

Ein *allgemeines Tripel* über $\mathbb{K}$ ist ein Tripel über $\mathbb{K}$ oder ein lineares Tripel über $\mathbb{K}$.

4.6.2 Bemerkung. In der Literatur findet man auch die Bezeichnung $*$ -Tripel für die hier definierten Tripel und mancherorts werden Tripel als die hier als lineare Tripel bezeichneten Objekte definiert.

4.6.3 Definition. Seien X und Y beide ein lineares Tripel über $\mathbb{K}$ oder beide ein Tripel über $\mathbb{K}$.

(a) Ein $x \in X$ heißt *tripotent* („Tripel-idempotent“), falls $\{xxx\} = x$ gilt. Ein tripotentes $x \in X$ wird auch als ein *Tripotent* von X bezeichnet. Die Menge aller tripotenten Elemente von X wird mit $\text{Tri}(X)$ bezeichnet. Falls X eine Topologie trägt, so ist, wenn nicht ausdrücklich etwas anderes vereinbart ist, $\text{Tri}(X)$ mit der von X induzierten Relativtopologie ausgestattet.

Ist X auch ein Banachraum über $\mathbb{K}$ und ist dann die Abbildung $\{ \}$ eine stetige Abbildung, so heißt X ein (*lineares*) *Banach-Tripel* über $\mathbb{K}$. Ein *allgemeines Banach-Tripel* über $\mathbb{K}$ ist ein Banach-Tripel über $\mathbb{K}$ oder ein lineares Banach-Tripel über $\mathbb{K}$.

(b) Mit D wird die Abbildung $X \times X \rightarrow L(X)$, $(x, y) \mapsto D(x, y)$ mit $D(x, y)(z) = \{xyz\}$ für alle $z \in X$, bezeichnet. Ebenfalls mit D wird die Abbildung $X \rightarrow L(X)$, $x \mapsto D(x, x)$ bezeichnet. Mit Q wird die Abbildung $X \times X \rightarrow L(X)$, $(x, z) \mapsto Q(x, z)$ mit $Q(x, z)(y) = \{xyz\}$ für alle $y \in X$, bezeichnet. Die zu Q assoziierte quadratische Abbildung wird ebenfalls mit Q bezeichnet. (Man vergleiche mit der Abbildung U_a im Unterpunkt 3.4.22(a) über Kompressionen.) Für $a \in X$ wird die Abbildung $X \times X \rightarrow X$, $(x, y) \mapsto \{xay\}$ die *a-Multiplikation* genannt. X wird *kommutativ* genannt, wenn $D(X, X)$ in $L(X)$ kommutativ ist. Ein $x \in \text{Tri}(X)$ heißt *regulär* (im Englischen *regular*), falls $D(x) \in GL(X)$ gilt. Für beliebige, aber feste $x, y \in X$ ist die Abbildung

$$B(x, y) := \text{Id}_X - 2D(x, y) + Q(x)Q(y)$$

ein Element von $L(X)$ und wird der zu dem Paar (x, y) assoziierte *Bergmann-Operator* genannt; anstelle von $B(x, x)$ schreibt man $B(x)$.

(c) Als *Jordan-Tripel-Gleichung* (in der Literatur auch *Jordan-Gleichung*, *Hauptgleichung*, *5-lineare Gleichung*) wird folgendes bezeichnet:

$$D(x, y)\{abc\} = \{D(x, y)a, bc\} - \{a, D(y, x)b, c\} + \{ab, D(x, y)c\}$$

für alle $x, y, a, b, c \in X$. Man beachte, dass in dem mittleren Term auf der rechten Seite der Gleichung $D(y, x)$ und nicht $D(x, y)$ steht.

(d) Ist die Abbildung $\{ \}$ außen kommutativ, das heißt, gilt $\{xyz\} = \{zyx\}$ für alle $x, y, z \in X$, so heißt $\{ \}$ ein *Tripelprodukt* auf X . Erfüllt ein Tripelprodukt die Jordan-Tripel-Gleichung, so nennt man es ein *Jordan-Tripelprodukt* und X heißt dann ein (*lineares*) *Jordan-Tripel* über $\mathbb{K}$. Ist X ein (*lineares*) Banach-Tripel über $\mathbb{K}$ und ist die Abbildung $\{ \}$ ein Jordan-Tripelprodukt, so heißt X ein (*lineares*) *Banach-Jordan-Tripel* über $\mathbb{K}$. Ein *allgemeines Jordan-Tripel* über $\mathbb{K}$ ist ein Jordan-Tripel über $\mathbb{K}$ oder ein lineares Jordan-Tripel über $\mathbb{K}$. *Allgemeines Banach-Jordan-Tripel* über $\mathbb{K}$ entsprechend. Ist X ein allgemeines Jordan-Tripel über $\mathbb{K}$, so heißt ein $x \in X$ *unitär*, falls $D(x) = \text{Id}_X$ gilt. (e) Ein *Untertripel* von X ist ein Untervektorraum U von X mit $\{UUU\} \subseteq U$. Ein *inneres Ideal* (im Englischen *inner ideal*) in X ist ein Untervektorraum U von X mit $\{UXU\} \subseteq U$. Ein *Ideal* (genauer ein *Tripelideal*) von X ist ein Untervektorraum U von X mit $\{UXX\} + \{XUX\} + \{XXU\} \subseteq U$; ist die Abbildung $\{ \}$ ein Tripelprodukt, so genügt es natürlich, $\{XUX\} + \{XXU\} \subseteq U$ als Idealeigenschaft zu fordern.

(f) Eine lineare Abbildung $T: X \rightarrow Y$ heißt ein *Homomorphismus* (genauer ein *Tripel-Homomorphismus*), wenn für alle $x, y, z \in X$ die Gleichung

$$T\{xyz\} = \{(Tx)(Ty)(Tz)\} \quad (4.8)$$

gilt. Ist solch ein Homomorphismus invertierbar, so heißt er ein *Isomorphismus* (genauer ein *Tripel-Isomorphismus*). Im Fall $\mathbb{K} = \mathbb{C}$ heißt eine konjugiert-lineare Abbildung $T: X \rightarrow Y$, die die Gleichung (4.8) für alle $x, y, z \in X$ erfüllt, ein *konjugiert-linearer Tripel-Homomorphismus*. Entsprechend ist ein *konjugiert-linearer Tripel-Isomorphismus* definiert.

(g) Eine lineare Abbildung $T: X \rightarrow X$ heißt eine *Tripel-Derivation*, wenn für alle $x, y, z \in X$ die folgende Gleichung erfüllt ist

$$T\{xyz\} = \{(Tx)y\} + \{x(Ty)z\} + \{xy(Tz)\}.$$

(h) Via dem System $(X, X; D; L(X))$ ist für jede Teilmenge M von X der Links- und Rechtsannihilator von M erklärt. Insofern dieses System orthosymmetrisch ist, ist somit auch der Begriff der Orthogonalität auf X erklärt und, wenn nicht ausdrücklich etwas anderes vereinbart ist, auch immer in diesem Sinne aufzufassen. Siehe auch 4.6.20, 4.7.3 und 4.7.61.

4.6.4. Sei X ein allgemeines Tripel über $\mathbb{K}$. Dann ist offensichtlich jedes Ideal von X ein inneres Ideal von X und jedes innere Ideal von X ist ein Untertripel von X .

4.6.5. Sei X ein allgemeines Tripel über $\mathbb{K}$ und U ein Untervektorraum von X . Des Weiteren nehme man an, dass auf X ein Tripelprodukt $\{ \}$ definiert sei. Dann ist U genau dann ein Untertripel von X , wenn $\{xyx\} \in U$ für alle $x, y \in U$ gilt. Im Fall $\mathbb{K} = \mathbb{C}$ ist U genau dann ein Untertripel von X , wenn $\{xxx\} \in U$ für alle $x \in U$ gilt.

Beweis. Für die erste Aussage beachte man nur, dass gemäß 2.2.41 für alle $x, z \in X$ die Gleichung

$$Q(x, z) = \frac{1}{2} (Q(x + z) - Q(x) - Q(z)) \quad (4.9)$$

gilt, also:

$$2\{xyz\} = \{(x + z)y(x + z)\} - \{xyx\} - \{zyz\} \quad \text{für alle } x, y, z \in X.$$

Für die Aussage im Fall $\mathbb{K} = \mathbb{C}$ wende man die Bemerkung 2.2.42 mit folgender Besetzung an: $X := X$, $Y := L(X)$ und $B(x, y) := \{xyx\}$ für alle $x, y \in X$ $\square$

4.6.6 Bemerkungen. (a) (UPMEIER [279](1985), Seite 329). Da $D(x, y)$ in der Literatur auch oft als $x \square y$ geschrieben wird, wird $D(x, y)$ auch *Box-Operator* genannt. Der Operator $D(x, y)$ ist ein von links operierender linearer Multiplikationsoperator (oder auch *Jordan-Multiplikationsoperator*). Siehe auch unter (c).

(b) Die endlichdimensionalen Jordan-Tripel über $\mathbb{C}$ heißen bei LOOS [202] (1977), Seite 2.9, hermitesche Jordan-Tripel-Systeme.

(c) Sei X ein Tripel über $\mathbb{C}$. Dann ist es bei einigen Autoren üblich, anstelle von $\{xyz\}$ die Schreibweise $\{xy^*z\}$ zu verwenden. Der Stern am y soll dabei nur daran erinnern, dass die Abbildung $\{ \}$ in der mittleren Variable konjugiert-linear ist und hat nichts mit einer Involution zu tun, die beispielsweise eventuell noch erst auf y angewendet werden würde. Dieselben Autoren, die dann auch noch die in (a) erwähnte Box-Operator-Schreibweise verwenden wollen, schreiben dann auch anstelle von $D(x, y)$ das Symbol $x \square y^*$, wobei auch hier der Stern wie eben nur als ein Merkzeichen dienen soll. Es sei aber auf die algebraisch allgemeinere Theorie der *Jordan-Paare* verwiesen, siehe LOOS [201](1975) und LOOS [202](1977). Jedes Jordan-Tripel X über $\mathbb{K}$ kann als ein Jordan-Paar $(X, \bar{X})$ geschrieben werden. Das Jordan-Tripelprodukt entspricht in Jordan-Paaren zwei trilinearen Abbildungen $X \times \bar{X} \times X \rightarrow X$ und $\bar{X} \times X \times \bar{X} \rightarrow \bar{X}$ (LOOS [202](1977), Seite 2.9). Bezeichne $\{ \}_{\text{Jordan-Paar}}$ die erstgenannte trilineare Abbildung und $^*: X \rightarrow \bar{X}$ die identische Abbildung. Dann gilt $\{xyz\}_{\text{Jordan-Tripel}} = \{xy^*z\}_{\text{Jordan-Paar}}$ für alle $x, y, z \in X$, wobei auf der rechten Seite nun wirklich erst auf y die konjugiert-lineare Abbildung * angewendet wird, bevor die trilineare Abbildung $\{ \}_{\text{Jordan-Paar}}$ agiert.

Im Anhang von LOOS [202](1977) findet man eine Liste von Gleichungen. Bei deren Verwendung ist zu beachten, dass Loos $Q(x)y$ als $\frac{1}{2}\{xyx\}$ definiert; da er mit diesem

seinem einparametrischen Q aber wiederum sein $Q(x, z)$ als $Q(x+z) - Q(x) - Q(z)$ definiert, stimmt andererseits wegen Gleichung (4.9) sein zweiparametrisches Q mit dem hier in 4.6.3(b) definierten zweiparametrischen Q überein. Siehe dazu auch LOOS, ebd., Seite 0.3, Bemerkung 0.8. In FARAUT, KANEYUKI, KORÁNYI, LU und ROOS [92](2000), Seite 430, findet man direkt für allgemeine Jordan-Tripel über $\mathbb{K}$ eine Auswahl von Gleichungen mit Beweisen, wobei auch dort nur zu beachten ist, dass das einparametrische Q wie bei LOOS, ebd., definiert ist.

4.6.7 Homotopie- und Fundamentalformel (UPMEIER [279](1985), Seite 297, Satz 18.2; siehe auch MEYBERG [212](1972), Seiten 69, 70, 81, 86, 93). Sei $(X, \{ \})$ ein allgemeines Tripel über $\mathbb{K}$ und liege der Fall vor, dass $\{ \}$ ein Tripelprodukt auf X sei. Dann ist $\{ \}$ genau dann ein Jordan-Tripelprodukt, wenn für alle $x, y \in X$ sowohl die sogenannte *Homotopieformel*

$$D(x, y)Q(x) = Q(x)D(y, x) = Q(Q(x)y, x) \quad (4.10)$$

als auch die Gleichung

$$D(Q(x)y, y) = D(x, Q(y)x) \quad (4.11)$$

gilt, also genau dann, wenn für alle $x, y, z \in X$ sowohl die Gleichungskette

$$\{xy\{xzx\}\} = \{x\{yxz\}x\} = \{\{xyx\}zx\}$$

als auch die Gleichung

$$\{\{xyx\}yz\} = \{x\{xyx\}z\}$$

gilt. Falls $\{ \}$ ein Jordan-Tripelprodukt ist, gilt für alle $x, y \in X$ die sogenannte *Fundamentalformel*

$$Q(Q(x)y) = Q(x)Q(y)Q(x), \quad (4.12)$$

ausgeschrieben also für alle $x, y, z \in X$ die Gleichung

$$\{\{xyx\}z\{xyx\}\} = \{x\{y\{xzx\}y\}x\}.$$

4.6.8 Bemerkung. Sei X ein allgemeines Tripel über $\mathbb{K}$. Die Jordan-Tripel-Gleichung lässt sich auch wie folgt schreiben:

$$[D(a, b), D(x, y)] = D(\{abx\}, y) - D(x, \{bay\})$$

für alle $x, y, a, b \in X$.

Beweis. Sei $a, b, x, y, v \in X$. Es ist $D(a, b)D(x, y)(v) = D(a, b)\{xyv\}$ und dies ist genau bei Gültigkeit der Jordan-Tripel-Gleichung gleich $\{D(a, b)x, y, v\} - \{x, D(b, a)y, v\} + \{x, y, D(a, b)v\}$. Mit anderen Worten gilt $D(a, b)D(x, y) = D(\{abx\}, y) - D(x, \{bay\}) + D(x, y)D(a, b)$, das heißt, die behauptete Gleichung. Gilt umgekehrt die behauptete Gleichung, so folgt aus der gleichen Argumentation, nur rückwärts gelesen, dass dann auch die Jordan-Tripel-Gleichung gilt. $\square$

Folglich ist $\text{span}\{D(x, y) \in L(X) : x, y \in X\}$ unter dem Kommutator von Definition 2.1.29 von $L(X)$ abgeschlossen und somit eine Lie-Unteralgebra von $L(X)^\ominus$.

4.6.9 Bemerkung. Bei Tripeln X über $\mathbb{C}$ erfüllt $D: X \times X \rightarrow L(X)$ die Jordan-Tripel-Gleichung genau dann, wenn die durch $\delta := iD: X \rightarrow L(X)$ definierte Abbildung die Eigenschaft hat, dass für alle $x \in X$ die Abbildung $\delta(x)$ eine *Tripel-Derivation* ist, wenn also gilt:

$$\delta(x)\{abc\} = \{\delta(x)a, bc\} + \{a, \delta(x)b, c\} + \{ab, \delta(x)c\}$$

für alle $x, a, b, c \in X$. Siehe auch Bemerkung 4.6.12(b).

Beweis. (Gemäß FRIEDMAN und RUSSO [101](1986), Seite 142 und LOOS [202] (1977), Seite 2.6.) Wenn $D: X \times X \rightarrow L(X)$ die Jordan-Tripel-Gleichung erfüllt, dann ist klar, dass die Abbildung δ die angegebene Eigenschaft hat.

Für die Umkehrung argumentiert man wie folgt. Für alle $x, y \in X$ gilt $D(x+y) - D(x) - D(y) = D(x, y) + D(y, x)$, also $i(D(x, y) + D(y, x)) = \delta(x+y) - \delta(x) - \delta(y)$. Setze $\Delta(x, y) := i(D(x, y) + D(y, x))$ für alle $x, y \in X$. Ist nun für jedes $x \in X$ die Abbildung $\delta(x)$ eine Tripel-Derivation, so gilt für alle $a, b, x, y, z \in X$: $\Delta(a, b)\{xyz\} = \delta(a+b)\{xyz\} - \delta(a)\{xyz\} - \delta(b)\{xyz\} = \{\delta(a+b)x, yz\} + \{x, \delta(a+b)y, z\} + \{xy, \delta(a+b)z\} - \{\delta(a)x, yz\} - \{x, \delta(a)y, z\} - \{xy, \delta(a)z\} - \{\delta(b)x, yz\} - \{x, \delta(b)y, z\} - \{xy, \delta(b)z\} = \{\Delta(a, b)x, yz\} + \{x, \Delta(a, b)y, z\} + \{xy, \Delta(a, b)z\}$. Dies zeigt, dass $\Delta(x, y)$ für alle $x, y \in X$ ebenfalls eine Tripel-Derivation ist. Wiederauflösen des Symbols Δ gibt: Für alle $a, b, x, y, z \in X$ gilt $iD(a, b)\{xyz\} + iD(b, a)\{xyz\} = \{iD(a, b)x, yz\} + \{x, iD(b, a)y, z\} + \{xy, iD(a, b)z\} + \{iD(b, a)x, yz\} + \{x, iD(a, b)y, z\} + \{xy, iD(b, a)z\}$. Befördert man alle Terme dieser Gleichung, die linear in a sind, auf die linke Seite der Gleichung und entsprechend alle Terme, die konjugiert-linear in a sind, auf die rechte Seite, so erhält man: Für alle $a, b, x, y, z \in X$ gilt $iD(a, b)\{xyz\} - \{iD(a, b)x, yz\} - \{x, iD(b, a)y, z\} - \{xy, iD(a, b)z\} = -iD(b, a)\{xyz\} + \{iD(b, a)x, yz\} + \{x, iD(a, b)y, z\} + \{xy, iD(b, a)z\}$. Bei beliebig gewählten, aber festen $b, x, y, z \in X$ ist die linke Seite dieser Gleichung eine lineare Funktion in der Variablen a , und wegen des Gleichheitszeichens aber auch eine konjugiert-lineare Funktion in der Variablen a , das heißt, es muss die Nullfunktion auf X sein. Somit gilt für alle $a, b, x, y, z \in X$ die Gleichung $iD(a, b)\{xyz\} = \{iD(a, b)x, yz\} + \{x, iD(b, a)y, z\} + \{xy, iD(a, b)z\}$ und Multiplikation mit $(-i)$ liefert für alle $a, b, x, y, z \in X$ die Jordan-Tripel-Gleichung für D . $\square$

4.6.10 Jordan-Algebra (ISIDRO und STACHÓ [149](1984), Seite 231; LOOS [200](1971), Satz 1.4; UPMEIER [279](1985), Seiten 317-320).

Sei X ein allgemeines Jordan-Tripel über $\mathbb{K}$ und $a \in X$. Dann ist X , ausgestattet mit der a -Multiplikation, eine Jordan-Algebra über $\mathbb{K}$ (siehe Definition 2.3.15); sie wird gelegentlich der Deutlichkeit halber auch mit $X^{(a)}$ bezeichnet. Für alle $x \in X$ gilt also in dieser Jordan-Algebra für die durch x bestimmte Links-Multiplikation L_x in $X^{(a)}$ die Gleichung $L_x = D(x, a)$. Existiert überdies in X ein unitäres Element $p \in X^\times$, so hat die Jordan-Algebra $X^{(p)}$ die Eins p . Liegt dann auch noch der Fall vor, dass X ein Jordan-Tripel X über $\mathbb{K}$ ist, so ist $Q(p)$ eine Involution dieser nichtassoziativen Algebra. Liegt dagegen der Fall vor, dass X ein lineares Jordan-Tripel X über $\mathbb{C}$ ist, so ist $Q(p)$ im Allgemeinen keine Involution dieser nichtassoziativen Algebra $X^{(p)}$, da $Q(p)$ wegen $Q(p)(ix) = iQ(p)(x)$ nicht unbedingt konjugiert-linear sein wird.

Sei nun umgekehrt X eine Jordan-Algebra über $\mathbb{K}$. Stattet man X mit dem Tripelprodukt

$$\{xyz\} := x(yz) - y(zx) + z(xy) \quad \text{für alle } x, y, z \in X \quad (4.13)$$

aus, so ist X ein lineares Jordan-Tripel über $\mathbb{K}$. Ist auf X eine Vektorraum-Involution $*$ definiert, so ist X mit dem Tripelprodukt

$$\{xyz\} := x(y^*z) - y^*(zx) + z(xy^*) \quad \text{für alle } x, y, z \in X \quad (4.14)$$

ein Jordan-Tripel über $\mathbb{K}$.

4.6.11 (KAUP [174](1981), Seite 465). Sei X ein Banach-Tripel über $\mathbb{C}$. Nach Bemerkung 2.2.42 gilt für alle $x, y \in X$ die Gleichung $D(x, y) = \frac{1}{4} \sum_{\epsilon^4=1} \epsilon D(x + \epsilon y)$. Wenn für alle $x \in X$ der Operator $D(x)$ ein hermitesches Element aus $\mathcal{L}(X)$ ist (Definition 2.11.22), dann ist wegen dieser Gleichung $D(X, X) \subseteq \mathcal{L}(X)_{\text{herm}, \mathbb{R}(\mathbb{C})}$ und mit $\bar{}$ als die kartesische

Involution auf $\mathcal{L}(X)_{\text{herm},\mathbb{R}(\mathbb{C})}$ gilt $\overline{D(x,y)} = \frac{1}{4}(D(x+y) - D(x-y) - iD(x+iy) + iD(x-iy)) = \frac{1}{4}(D(y+x) - D(y-x) - iD(y-ix) + iD(y+ix)) = D(y,x)$ für alle $x, y \in X$. Gilt andererseits $D(x,y) \in \mathcal{L}(X)_{\text{herm},\mathbb{R}(\mathbb{C})}$ für alle $x, y \in X$, wobei dann jeweils auch die Gleichung $\overline{D(x,y)} = D(y,x)$ erfüllt sein soll, dann gilt mit $x = y$ wegen $\mathcal{L}(X)_{\text{herm},\mathbb{R}(\mathbb{C})}^{(-)} = \mathcal{L}(X)_{\text{herm},\mathbb{R}}$ (siehe 3.1.8), dass für alle $x \in X$ der Operator $D(x)$ hermitesch ist. Somit sind äquivalent:

- (a) Für alle $x \in X$ ist $D(x)$ ein hermitesches Element aus $\mathcal{L}(X)$. (Definition 2.11.22.)
- (b) Es ist $D(x,y) \in \mathcal{L}(X)_{\text{herm},\mathbb{R}(\mathbb{C})}$ für alle $x, y \in X$ und die sesquilineare Abbildung D ist bezüglich der kartesischen Involution auf $\mathcal{L}(X)_{\text{herm},\mathbb{R}(\mathbb{C})}$ hermitesch. (Definition 3.1.17.)

Für eine weitere Äquivalenz für den Fall, dass X ein Banach-Jordan-Tripel über $\mathbb{C}$ ist, siehe Satz 4.6.15.

4.6.12 Bemerkungen. (a) Äquivalent zu der Definition 4.6.3(d) kann man ein Banach-Jordan-Tripel über $\mathbb{K}$ definieren als einen Banachraum X über $\mathbb{K}$ mit einer sesquilinearen Abbildung $D: X \times X \rightarrow \mathcal{L}(X)$, so dass mit $\{xyz\} := D(x,y)(z)$, $x, y, z \in X$, die Abbildung $\{ \}$ außen symmetrisch ist (also $\{ \}$ ein Tripelprodukt ist) und die Jordan-Tripel-Gleichung gilt.

(b) Anknüpfend an Bemerkung 4.6.9 hat man: Ist X ein Banach-Tripel über $\mathbb{C}$ mit einem Tripel-Produkt, so ist aufgrund der Stetigkeit des Tripel-Produktes der Operator $\exp(tD(x)) \in \mathcal{L}(X)$ für alle $t \in \mathbb{C}$, $x \in X$ wohldefiniert und das Tripel-Produkt ist dann genau dann ein Jordan-Tripelprodukt, X also ein Banach-Jordan-Tripel über $\mathbb{C}$, wenn für alle $t \in \mathbb{R}$, $x \in X$ die Operatoren $\exp(itD(x))$ Tripel-Isomorphismen sind (STACHÓ und ISIDRO [266](1990)); nach KAUP [179](2007) ist dies wiederum genau dann der Fall, wenn $\exp(iD(x)) \in G\mathcal{L}(X)$ für alle $x \in X$ ein Tripel-Isomorphismus ist.

4.6.13 Spur III (MCCRIMMON [210](2004), Seiten 23, 27). Ein endlichdimensionales Jordan-Tripel über $\mathbb{C}$ heißt *positiv*, wenn die Spurform $(x,y) \mapsto \text{sp}(D(x,y))$ (siehe 3.4.1) ein Skalarprodukt ist.

OTTMAR LOOS zeigte, dass zwischen beschränkten, homogenen, zirkulären Gebieten des $\mathbb{C}^n$ und den endlichdimensionalen positiven Jordan-Tripeln über $\mathbb{C}$ eine natürliche 1-1-Korrespondenz besteht.

In unendlichdimensionalen Banachräumen lässt sich im Allgemeinen weder die übliche Schreibweise $T = T^*$ für selbstadjungierte Operatoren auf Hilberträumen anwenden noch eine Spur finden. Zum Beispiel gibt es keine stetige Spur $\neq 0$ — und dies sogar im Sinne der Definition 3.2.24(d) — weder auf $\mathcal{L}(H)$ noch auf $HS(H)$, wenn H ein unendlichdimensionaler Hilbertraum über $\mathbb{C}$ ist (siehe DIEUDONNÉ [68](1987), Band 2, Seiten 322 und 324). Man fordert also, dass die Operatoren $D(x,y)$ selber in gewissen Sinne positiv definit sind. Dieses Konzept ist mit den weiter unten definierten JB^* -Tripeln verwirklicht; siehe 4.7.4.

4.6.14 Definition (UPMEIER [279] (1985), Seite 305). Ein Banach-Jordan-Tripel X über $\mathbb{K}$ wird *hermitesch* genannt, wenn für alle $x, y \in X$ der Operator $D(x,y) - D(y,x) \in \mathcal{L}(X)$ eine infinitesimale Isometrie ist, das heißt, wenn gilt:

$$\left\| \exp(t(D(x,y) - D(y,x))) \right\| \leq 1 \quad \text{für alle } x, y \in X, t \in \mathbb{R}.$$

4.6.15 Satz (UPMEIER [279](1985), Seite 305). Ist X ein Banach-Jordan-Tripel über $\mathbb{C}$, so sind äquivalent:

(a) X ist $\neg$ hermitesch.

(b) $D(x)$ ist hermitesch für alle $x \in X$. (Definition 2.11.22)

Beweis. Wegen $D(x, zx) - D(zx, x) = (\bar{z} - z)D(x) = -2i\operatorname{Im}(z)D(x)$ für alle $z \in \mathbb{C}$ und $x \in X$, hat man per $z = -\frac{1}{2}i$ die Gleichung $D(x, zx) - D(zx, x) = iD(x)$ für alle $x \in X$. Ist nun X $\neg$ hermitesch, so gilt also $\|\exp(itD(x))\| \leq 1$ für alle $t \in \mathbb{R}$ und $x \in X$, das heißt, gemäß Satz 2.11.27 gilt (b). Um (a) aus (b) zu folgern, beachte man, dass für alle $x, y \in X$ gilt: $D(x + iy) - D(x) - D(y) = D(x + iy, x + iy) - D(x) - D(y) = D(x, iy) + D(iy, x) = i(D(y, x) - D(x, y))$, also $D(x, y) - D(y, x) = i(D(x + iy) - D(x) - D(y))$. Gilt nun (b), so hat man demnach für alle $t \in \mathbb{R}$ und $x, y \in X$: $\|\exp(t(D(x, y) - D(y, x)))\| = \|\exp(it(D(x + iy) - D(x) - D(y)))\| = \|\exp(itD(x + iy)) \exp(-itD(x)) \exp(-itD(y))\| \leq \|\exp(itD(x + iy))\| \|\exp(-itD(x))\| \|\exp(-itD(y))\| \leq 1$ und X ist $\neg$ hermitesch. $\square$

4.6.16 Definition (UPMEIER [279] (1985), Seite 331). Ein Banach-Jordan-Tripel X über $\mathbb{K}$ wird $^+$ hermitesch genannt, wenn für alle $x, y \in X$ die beiden folgenden Aussagen gelten:

(a) $\sigma(D(x, y) + D(y, x)) \subseteq \mathbb{R}$,

(b) $\|D(x)\| = \rho(D(x))$. (Definition 2.12.3)

Merke: Bei $\neg$ hermitesch sind die beiden Terme $D(x, y)$ und $D(y, x)$ mit einem Minuszeichen, bei $^+$ hermitesch mit einem Pluszeichen verbunden.

4.6.17 Definition. Ein Banach-Jordan-Tripel X über $\mathbb{K}$ wird *hermitesch* genannt, wenn es sowohl $\neg$ hermitesch als auch $^+$ hermitesch ist.

4.6.18. Sei X ein $\neg$ hermitesches Banach-Jordan-Tripel über $\mathbb{C}$. Da auf jedem Tripel A für alle $x, y \in A$ gilt: $D(x + y) = D(x + y, x + y) = D(x) + D(x, y) + D(y, x) + D(y)$, also $D(x, y) + D(y, x) = D(x + y) - D(x) - D(y)$, hat man hier für alle $x, y \in X$ und für alle $t \in \mathbb{R}$: $\|\exp(it(D(x, y) + D(y, x)))\| \leq \|\exp(itD(x + y))\| \|\exp(-itD(x))\| \|\exp(-itD(y))\| \leq 1$ nach Satz 4.6.15 und Satz 2.11.27(a)/(c). Das heißt, wieder nach Satz 2.11.27(a)/(c), dass der Operator $D(x, y) - D(y, x)$ hermitesch ist, also nach Satz 2.11.27(a)/(a.i) und Satz 2.12.8 ein reelles Spektrum hat und somit (a) von der Definition 4.6.16 der $^+$ Hermitezität gilt; nach Satz 4.6.15 und nach dem Satz 2.12.15 für $\lambda = 0$ gilt auch (b) von der Definition 4.6.16 der $^+$ Hermitezität. Somit ist jedes $\neg$ hermitesche Banach-Jordan-Tripel über $\mathbb{C}$ automatisch $^+$ hermitesch und somit hermitesch. Damit stimmt auch die in der Literatur oft gefundene Definition der Hermitezität von Banach-Jordan-Tripeln über $\mathbb{C}$ — in der üblicherweise nur die $\neg$ Hermitezität gefordert wird — mit der hier gegebenen Definition überein.

4.6.19 Definition. Ein *lineares J^* -Tripel* ist ein lineares Banach-Jordan-Tripel über $\mathbb{C}$, bei dem für alle $x \in X$ der Operator $D(x)$ hermitesch (Definition 2.11.22) ist. Ein J^* -Tripel über $\mathbb{K}$ ist ein hermitesches Banach-Jordan-Tripel über $\mathbb{K}$.

4.6.20 Bemerkung. Sei X ein J^* -Tripel über $\mathbb{C}$ oder ein lineares J^* -Tripel. Dann folgt aus 4.6.18, Satz 4.6.15 und 4.6.11, dass das System $(X, X; D; L(X))$ orthosymmetrisch ist.

KAUP [173](1977) hat gezeigt: Die Kategorie der einfach zusammenhängenden, symmetrischen Banach-Mannigfaltigkeiten über $\mathbb{C}$ mit Basispunkt ist äquivalent zu der Kategorie der J^* -Tripel über $\mathbb{C}$.

4.6.21 (KAUP [173](1977), Seite 40). Sei $\{X_\nu : \nu \in I\}$ eine Familie von J^* -Tripeln über $\mathbb{C}$. Setzt man $X := \ell_\infty(\bigoplus_{\nu \in I} X_\nu)$ und erklärt man für alle $u, v, w \in X$ den Ausdruck $\{uvw\}$ als das $x \in X$ mit $x_\nu := \{u_\nu v_\nu w_\nu\}$ für alle $\nu \in I$, so ist auch X ein J^* -Tripel über $\mathbb{C}$.

4.6.22 Definition (UPMEIER [279](1985), Seite 336). Sei X ein J^* -Tripel über $\mathbb{K}$ oder ein lineares J^* -Tripel. Dann heißt X *positiv*, wenn $\sigma(D(x)) \subseteq \mathbb{R}^+$ für alle $x \in X$ gilt.

4.6.23. Sei X ein lineares J^* -Tripel X oder ein J^* -Tripel über $\mathbb{C}$. Da nach Satz 4.6.15 für alle $x \in X$ die Operatoren $D(x)$ hermitesch sind, also nach Bemerkung 2.11.22 alle einen reellen numerischen Wertebereich haben, haben einerseits nach Satz 2.12.8 für alle $x \in X$ die Operatoren $D(x)$ ein reelles Spektrum und andererseits ist X genau dann positiv, wenn $D(X) \subseteq \text{Pos}(V\sigma; \mathcal{L}(X))$ gilt.

4.6.24 Bemerkung (KAUP [175](1983), Seite 518; UPMEIER [279](1985), Seite 343). Auf jedem positiven J^* -Tripel X über $\mathbb{C}$ ist

$$\|\cdot\|_\infty : x \mapsto \sqrt{\|D(x)\|}$$

eine stetige Halbnorm. Sie wird *spektrale Halbnorm* genannt.

4.6.25 Definition (KAUP [175](1983), Seite 508). Für jedes J^* -Tripel X über $\mathbb{C}$ und jedes $x \in X$ bezeichnet X_x den kleinsten abgeschlossenen Untervektorraum von X , der x enthält und invariant bezüglich $D(x)$ ist.

4.6.26 Bemerkung (KAUP [175](1983), Seiten 508, 516 und 518). Sei X ein positives J^* -Tripel über $\mathbb{C}$. Dann ist X_x ein Jordan-Untertripel und wird das *von x erzeugte Untertripel* genannt. Für alle $x \in X$ ist $\sigma(D(x)|X_x)$ eine kompakte Teilmenge von $\mathbb{R}$ und es gilt:

$$\|x\|_\infty = \left(\max \sigma(D(x)|X_x) \right)^{1/2} \quad \text{für alle } x \in X.$$

4.6.27 (KAUP [175](1983), Seiten 505 und 522). $\mathbb{C}$, versehen mit dem Tripel-Produkt $\{xyz\} := x\bar{y}z$ für alle $x, y, z \in \mathbb{C}$, ist ein J^* -Tripel über $\mathbb{C}$.

4.7 JB*-Tripel

4.7.1 Definition. Ein J^* -Tripel X über $\mathbb{K}$ erfüllt die *C^* -Bedingung für Tripel*, wenn für alle $x \in X$ die Gleichung $\|D(x)\| = \|x\|^2$ gilt. Ein *JB*-Tripel* über $\mathbb{K}$ ist ein positives J^* -Tripel über $\mathbb{K}$, das die C^* -Bedingung für Tripel erfüllt.

4.7.2 Bemerkungen. Ausformuliert ist also ein *JB*-Tripel* über $\mathbb{K}$ ein positives hermitesches Banach-Jordan-Tripel über $\mathbb{K}$, das die C^* -Bedingung für Tripel erfüllt. W. Kaup bezeichnet die *JB*-Tripel* über $\mathbb{C}$ in [173](1977) noch als *C^* -Tripel*. DINEEN [71](1986) bezeichnet abgeschlossene Tripelideale eines *JB*-Tripels* über $\mathbb{C}$ als *J^* -Ideale*.

4.7.3 C^* -Algebra. Sei A eine C^* -Algebra über $\mathbb{K}$. Dann ist A , versehen mit dem Tripelprodukt $\{xyz\} := \frac{1}{2}(xy^*z + zy^*x)$, ein *JB*-Tripel* über $\mathbb{K}$; es wird mit A^Δ bezeichnet. Zum Beweis siehe UPMEIER [279](1985), Seite 337. Dabei ist es auch üblich, von A als einem *JB*-Tripel* über $\mathbb{K}$ zu reden, auch wenn genau genommen A^Δ gemeint ist.

Es sei an die Definition der Orthogonalität in C^* -Algebren in 3.5.21 erinnert; gemäß dieser gilt im Einklang mit Definition 4.6.3(h): $x, y \in A$ sind genau dann in der C^* -Algebra A orthogonal, wenn in A^Δ die Gleichung $D(x, y) = 0$ erfüllt ist.

In der Theorie der Tripel-Systeme spielen die Tripotente die gleiche Rolle wie die Projektionen bei den C^* -Algebren (ISIDRO [146](1989), Seite 150). Siehe auch 4.7.33

und 4.7.66. Mit Bemerkung 3.5.20 gilt offensichtlich: Genau die partiellen Isometrien (Definition 3.5.19) von A sind die tripotenten Elemente von A^Δ . Vergleiche 3.5.36.

Eine normierte Algebra A über $\mathbb{C}$ ist genau dann eine C^* -Algebra über $\mathbb{C}$, wenn A sowohl linear isometrisch isomorph zu einem JB^* -Tripel über $\mathbb{C}$ ist als auch eine durch 1 beschränkte approximative Eins besitzt (EL AMIN, MORALES CAMPOY und RODRÍGUEZ PALACIOS [89](2001), Korollar 3.4).

Sei A eine C^* -Algebra über $\mathbb{C}$. Dann ist ein $x \in A$ genau dann unitär, wenn $x \in A^\Delta$ unitär ist. Allgemeiner als nur für C^* -Algebren über $\mathbb{C}$ gilt die Aussage (3.15) von 3.5.33 auch für jedes JB^* -Tripel A über $\mathbb{C}$ (RODRÍGUEZ PALACIOS [248](2010), Theorem 3.1).

4.7.4 Skalarprodukt. Sei X ein JB^* -Tripel über $\mathbb{C}$. Nach den Äquivalenzaussagen in 4.6.11 und 4.6.23 ist die auf $X \times X$ definierte Abbildung D ein $\mathcal{L}(X)$ -wertiges Skalarprodukt auf X (Definition 3.1.17). Man beachte, dass dabei die positive Definitheit — die ja bereits in 4.6.13 gefordert wurde — eine unmittelbare Folgerung der C^* -Bedingung für Tripel ist.

4.7.5 Reelle Strukturierung. Ist Y ein JB^* -Tripel über $\mathbb{C}$, dann ist $Y_{\mathbb{R}}$ ein JB^* -Tripel über $\mathbb{R}$.

Die entsprechende Aussage gilt auch, wenn Y ein positives (lineares) J^* -Tripel über $\mathbb{C}$, ein (lineares) J^* -Tripel über $\mathbb{C}$, ein $^+$ hermitesches allgemeines Banach-Jordan-Tripel über $\mathbb{C}$, ein allgemeines Banach-Jordan-Tripel über $\mathbb{C}$, ein allgemeines Jordan-Tripel über $\mathbb{C}$, ein allgemeines Banach-Tripel über $\mathbb{C}$ oder ein allgemeines Tripel über $\mathbb{C}$ ist.

Beweis. Sei Y ein JB^* -Tripel über $\mathbb{C}$, also ein positives, hermitesches Banach-Jordan-Tripel über $\mathbb{C}$, das die C^* -Bedingung für Tripel erfüllt. Bezeichne X den Banachraum $Y_{\mathbb{R}}$ über $\mathbb{R}$ und ι die identische Abbildung von der Menge X auf Y . Versieht man X mit der Abbildung

$$\{ \}_X: X^3 \rightarrow X, \quad (x, y, z) \mapsto \iota^{-1}\{\iota(x), \iota(y), \iota(z)\}_Y,$$

so ist X ein Tripel über $\mathbb{R}$.

Mit $\{ \}_Y$ als eine stetige Abbildung ist auch die Abbildung $\{ \}_X$ stetig. Somit ist X ein Banach-Tripel über $\mathbb{R}$.

Sei D_Y die Abbildung $Y \times Y \rightarrow L(Y)$ mit $D_Y(y_1, y_2)(y_3) = \{y_1 y_2 y_3\}_Y$ für alle $y_1, y_2, y_3 \in Y$. Entsprechend ist $D_Y \rightarrow L(Y)$ die durch $D_Y(y) := D_Y(y, y)$, $y \in Y$, definierte Abbildung. D_X seien die beiden vollkommen analog definierten Abbildungen für X . Mit D_Y erfüllt auch D_X die Jordan-Tripel-Gleichung und mit $\{ \}_Y$ ist auch $\{ \}_X$ außen kommutativ; das heißt, $\{ \}_X$ ist ein Tripelprodukt auf X und X ist ein Banach-Jordan-Tripel über $\mathbb{R}$.

Sei $x, y \in X$, $t \in \mathbb{R}$. Setze $T := t(D_Y(\iota(x), \iota(y)) - D_Y(\iota(y), \iota(x)))$. Sei γ die in 2.9.8 betrachtete isometrische Einbettung $\gamma: (\mathcal{L}(Y))_{\mathbb{R}} \hookrightarrow \mathcal{L}(X)$. Nach Definition von D_X und D_Y gilt $\gamma(T) = S$ mit $S = t(D_X(x, y) - D_X(y, x))$. Für alle $n \in \mathbb{N}$ gilt: $\gamma\left(\sum_{k=1}^n \frac{T^k}{k!}\right) = \sum_{k=1}^n \gamma\left(\frac{T^k}{k!}\right)$ und $\left\|\sum_{k=1}^n \gamma\left(\frac{T^k}{k!}\right)\right\| \leq \sum_{k=1}^n \left\|\gamma\left(\frac{T^k}{k!}\right)\right\| = \sum_{k=1}^n \left\|\frac{T^k}{k!}\right\| \leq \sum_{k=1}^n \frac{\|T\|^k}{k!} \leq \exp(\|T\|)$. Es folgt $\gamma(\exp(T)) = \exp(\gamma(T)) = \exp(S)$ und S ist $^+$ hermitesch.

Setze nun $T := D_Y(\iota(x), \iota(y)) + D_Y(\iota(y), \iota(x))$. Dann gilt $\gamma(T) = S$ mit $S = D_X(x, y) + D_X(y, x)$. Nach Bemerkung 2.12.6 gilt $\sigma(S) = \sigma(T) \cup \overline{\sigma(T)}$, hier also $\sigma(S) = \sigma(T) \subseteq \mathbb{R}$; ebenso hat man $\sigma(D_X(x)) = \sigma(D_Y(\iota(x)))$. Somit: $\|D_X(x)\| = \|\gamma(D_Y(\iota(x)))\| = \|D_Y(\iota(x))\| = \rho(D_Y(\iota(x))) = \sup\{|\lambda| : \lambda \in \sigma(D_Y(\iota(x)))\} = \sup\{|\lambda| : \lambda \in \sigma(D_X(x))\} = \rho(D_X(x))$ und X ist auch $^+$ hermitesch, also hermitesch, also ein J^* -Tripel über $\mathbb{R}$.

Klar, wegen der Spektrumsgleichheit: X ist positiv. Auch klar: X erfüllt die C^* -Bedingung für Tripel. $\square$

4.7.6 Bemerkung. Sei X sowohl ein Tripel über $\mathbb{K}$, als auch ein Banachraum über $\mathbb{K}$, Y ein JB^* -Tripel über $\mathbb{K}$ und $T: X \rightarrow Y$ ein isometrischer Tripel-Isomorphismus von X auf Y . Dann ist auch X ein JB^* -Tripel über $\mathbb{K}$. Speziell gilt $D_X(x, y) = T^{-1} \circ D_Y(Tx, Ty) \circ T$ und $\sigma_X(D_X(x, y)) = \sigma_Y(D_Y(Tx, Ty))$ für alle $x, y \in X$.

Beweis. Wegen $\{xyz\}_X = T^{-1}\{Tx Ty Tz\}_Y$ ist X ein Banach-Tripel über $\mathbb{K}$. Es gilt $D_X(x, y) = T^{-1} \circ D_Y(Tx, Ty) \circ T$ für alle $x, y \in X$. D_X erfüllt die Jordan-Tripel-Gleichung und X ist ein Banach-Jordan-Tripel über $\mathbb{K}$. Man überprüft leicht, dass die folgende Äquivalenz gilt: $\forall N \in \mathbb{N} \forall t \in \mathbb{R} \forall x, y \in Y : \|\text{Id}_Y + \sum_{n=1}^N (n!)^{-1}(t(D(x, y) - D(y, x)))^n\|_Y \leq 1 \Leftrightarrow \forall N \in \mathbb{N} \forall t \in \mathbb{R} \forall x, y \in X : \|\text{Id}_X + \sum_{n=1}^N (n!)^{-1}(t(D(x, y) - D(y, x)))^n\|_X \leq 1$. Somit gilt auch die folgende Äquivalenz: $\forall t \in \mathbb{R} \forall x, y \in Y : \|\exp(t(D(x, y) - D(y, x)))\|_Y \leq 1 \Leftrightarrow \forall t \in \mathbb{R} \forall x, y \in X : \|\exp(t(D(x, y) - D(y, x)))\|_X \leq 1$. Das heißt, X ist hermitesch.

Seien nun $x, y \in X$ und sei $\lambda \in \mathbb{C}$ jeweils beliebig, aber fest. Dann gelten folgende Äquivalenzen: $\lambda \cdot \text{Id}_X - D(x, y)$ ist in $\mathcal{L}(X)$ invertierbar $\Leftrightarrow \exists G \in \mathcal{L}(X)$ mit $\text{Id}_X = (\lambda \cdot \text{Id}_X - D(x, y)) \circ G = G \circ (\lambda \cdot \text{Id}_X - D(x, y)) \Leftrightarrow \exists G \in \mathcal{L}(X) : \text{Id}_Y = (\lambda \cdot T \circ \text{Id}_X - D_Y(Tx, Ty) \circ T) \circ G \circ T^{-1} = T \circ G \circ (\lambda \cdot \text{Id}_X \circ T^{-1} - T^{-1} \circ D_Y(Tx, Ty)) \Leftrightarrow \exists G \in \mathcal{L}(X) : \text{Id}_Y = (\lambda \cdot \text{Id}_Y - D_Y(Tx, Ty)) \circ T \circ G \circ T^{-1} = T \circ G \circ T^{-1}(\lambda \cdot \text{Id}_Y - D_Y(Tx, Ty)) \Leftrightarrow \lambda \cdot \text{Id}_Y - D_Y(Tx, Ty)$ ist in $\mathcal{L}(Y)$ invertierbar. Einerseits folgt speziell für alle $x \in X$ und allen $\lambda \in \mathbb{C}$ die Äquivalenz $\lambda \in \sigma_X(D(x)) \Leftrightarrow \lambda \in \sigma_Y(D(Tx))$. Und andererseits gilt also für alle $\lambda \in \mathbb{C}$ die Äquivalenz: $\lambda \cdot \text{Id}_X - D(x, y)$ ist für alle $x, y \in X$ in $\mathcal{L}(X)$ invertierbar $\Leftrightarrow \lambda \cdot \text{Id}_Y - D_Y(Tx, Ty)$ ist für alle $x, y \in Y$ in $\mathcal{L}(Y)$ invertierbar. Damit gilt in X für alle $x, y \in X$: $\sigma(D(x, y) + D(y, x)) \subseteq \mathbb{R}$.

Sei $x \in X$. Dann gilt: $\|D_X(x)\| = \sup_{y \in S_X} \|\{xxy\}_X\|_X = \sup_{y \in S_X} \|T\{xxy\}_X\|_Y = \sup_{y \in S_Y} \|\{Tx Tx y\}_Y\|_Y = \|D_Y(Tx)\|$. Somit gilt $\|D_X(x)\| = \|D_Y(Tx)\| = \rho_Y(D_Y(Tx)) = \sup\{|\lambda| : \lambda \in \sigma_Y(D_Y(Tx))\} = \sup\{|\lambda| : \lambda \in \sigma_X(D_X(x))\} = \rho_X(D_X(x))$ und X ist $^+$ hermitesch, also hermitesch.

Wegen $\sigma_X(D_X(x)) = \sigma_Y(D_Y(Tx))$ für alle $x \in X$ sieht man sofort, dass X positiv ist. C^* -Bedingung für Tripel: $\|D(x)\| = \sup_{y \in S_X} \|\{xxy\}\|_X = \sup_{y \in S_X} \|\{Tx Tx Ty\}\|_Y = \sup_{y \in S_Y} \|\{Tx Tx y\}\|_Y = \|D(Tx)\| = \|Tx\|^2 = \|x\|^2$ und X ist ein JB^* -Tripel über $\mathbb{K}$. $\square$

4.7.7 Satz (KAUP [175](1983), Seite 523). *Sei X ein J^* -Tripel über $\mathbb{C}$. Dann sind äquivalent:*

- (a) X ist ein JB^* -Tripel über $\mathbb{C}$.
- (b) X ist positiv und $\|\{xxx\}\| = \|x\|^3$ für alle $x \in X$.
- (c) Für jedes $x \in X$ existiert eine kommutative C^* -Algebra A über $\mathbb{C}$ und ein isometrischer Tripel-Isomorphismus von X_x auf A^Δ .
- (d) X ist positiv und $\|\cdot\| = \|\cdot\|_\infty$. (Siehe Bemerkung 4.6.24.)

Beweis. Siehe auch:

(a) $\Leftrightarrow$ (b): UPMEIER [279](1985), Seite 336 und Seite 348, Satz 20.25.

(a) $\Leftrightarrow$ (d): Dies gilt per Definition der spektralen Halbnorm in Bemerkung 4.6.24. $\square$

4.7.8 Bemerkung. Da in einem JB^* -Tripel über $\mathbb{C}$ nach Satz 4.7.7 das Tripelprodukt und die Norm sich gegenseitig eindeutig bestimmen, nennt man die Norm eines JB^* -Tripels über $\mathbb{C}$ auch *die JB^* -Norm* von $\{\cdot\}$. Für JB^* -Tripel über $\mathbb{R}$ liegen die Verhältnisse anders! Siehe 4.7.63.

4.7.9 Satz (HORN [140](1987), Satz 2.4). *Seien X und Y zwei JB*-Tripel über $\mathbb{K}$. Sei $T: X \rightarrow Y$ ein Tripel-Isomorphismus der zugrunde liegenden Tripel über $\mathbb{K}$. (Die Abbildung T wird also nicht als stetig vorausgesetzt.) Dann ist T eine Isometrie.*

Beweis. Sei $x \in X$. Sei $\lambda \notin \sigma(D(x))$. Dann existiert also ein $S \in L(X)$ mit $S \circ (\lambda - D(x)) = \text{Id}_X \Leftrightarrow ST^{-1}T(\lambda - D(x)) = \text{Id}_X \Leftrightarrow ST^{-1}(\lambda - D(Tx)) \circ T = \text{Id}_X \Leftrightarrow TST^{-1}(\lambda - D(Tx)) = TT^{-1} \Leftrightarrow TST^{-1}(\lambda - D(Tx)) = \text{Id}_Y$. Es folgt $\lambda \notin \sigma(D(Tx))$. Also gilt $\sigma(D(x)) = \sigma(D(Tx))$, womit das dritte Gleichheitszeichen in der folgenden Gleichungskette gerechtfertigt wird. $\|x\|^2 = \|D(x)\|$ (C^* -Bedingung für Tripel) $= \sup \{|\lambda| : \lambda \in \sigma(D(x))\}$ ($^+$ Hermitezität) $= \sup \{|\lambda| : \lambda \in \sigma(D(Tx))\} = \|D(Tx)\|$ ($^+$ Hermitezität) $= \|Tx\|^2$ (C^* -Bedingung für Tripel). $\square$

4.7.10 Satz (UPMEIER [279](1985), Seite 336). *Sei X ein JB*-Tripel über $\mathbb{K}$. Dann gilt für alle $x \in X$ die Gleichung $\|\{xxx\}\| = \|x\|^3$.*

4.7.11 Untertripel und JC^* -Tripel (UPMEIER [279](1985), Seiten 301, 305 und 337). Als Korollar von Satz 4.7.10 hat man: Sei X ein JB*-Tripel über $\mathbb{K}$ und U ein abgeschlossenes Untertripel von X . Dann ist auch U ein JB*-Tripel über $\mathbb{K}$. Möchte man dies betonen, so sagt man in dieser Situation auch, dass U ein Unter-JB*-Tripel über $\mathbb{K}$ von X sei.

Betrachte $\mathcal{L}(H_1, H_2)$ für Hilberträume H_1, H_2 über $\mathbb{K}$. (Es sei an 3.4.8 erinnert.) Bezeichne für $x \in \mathcal{L}(H_1, H_2)$ die Hilbertraumadjungierte mit x^* . Dann ist $\mathcal{L}(H_1, H_2)$, versehen mit dem Tripelprodukt $\{xyz\} := \frac{1}{2}(xy^*z + zy^*x)$, ein JB*-Tripel über $\mathbb{K}$; in Anlehnung an 4.7.3 wird es mit $\mathcal{L}(H_1, H_2)^\Delta$ bezeichnet, oder einfach auch wieder als $\mathcal{L}(H_1, H_2)$, wenn der Zusammenhang klar ist. Für die Bergmann-Operatoren gilt dann:

$$B(x, y)z = (\text{Id}_{H_2} - xy^*)z(\text{Id}_{H_1} - y^*x) \quad \text{für alle } x, y, z \in \mathcal{L}(H_1, H_2);$$

man vergleiche mit Gleichung (3.14) von 3.5.33. Man bemerke, dass wegen der für jeden Hilbertraum H über $\mathbb{K}$ bestehenden isometrischen Isomorphie $H \cong \mathcal{L}(\mathbb{K}, H)$ jeder Hilbertraum über $\mathbb{K}$ als ein JB*-Tripel über $\mathbb{K}$ aufgefasst werden kann, siehe 4.7.63 für eine Konkretisierung. Die abgeschlossenen Untertripel von $\mathcal{L}(H_1, H_2)^\Delta$, die ja nach dem eingangs genannten Korollar JB*-Tripel über $\mathbb{K}$ sind, werden JC^* -Tripel über $\mathbb{K}$ genannt. (Bei HARRIS [128](1973) und ISIDRO [146](1989) heißen sie J^* -Algebren.)

Als Beispiele von JC^* -Tripeln über $\mathbb{C}$ führt KAUP [179](2007), Seite 50, an, die abgeschlossenen linearen Hüllen in $\mathcal{L}(H)$, H ein Hilbertraum über $\mathbb{C}$, von allen Teilmengen ABC mit $B \subseteq \mathcal{L}(H)$ eine C^* -Algebra über $\mathbb{C}$ und $A, C \subseteq B$ beliebig.

4.7.12 Eindeutige Hahn-Banach-Fortsetzung (EDWARDS und RÜTTIMANN [85](1992), Theorem 2.6). Sei X ein JB*-Tripel über $\mathbb{C}$. Sei U ein abgeschlossenes Untertripel von X . Dann ist U genau dann ein inneres Ideal von X , wenn die für jedes lineare, stetige Funktional auf U durch den Satz von Hahn-Banach garantierte normgleiche lineare Fortsetzung auf X eindeutig bestimmt ist.

4.7.13 Die Norm eines Tripotents. Sei X ein JB*-Tripel über $\mathbb{K}$. Als eine unmittelbare Folgerung aus dem Satz 4.7.10 gilt $\text{Tri}(X)^\times \subseteq S_X$. Insbesondere bemerke man, dass jedes unitäre Element von X stets tripotent ist.

4.7.14 Definition (UPMEIER [279](1985), Seite 349). Sei X eine unitale Banach-Jordan-Algebra über $\mathbb{C}$ (siehe Definition 2.3.15), versehen mit einer Involution $*$. Sei X mit dem in 4.6.10 definierten Tripelprodukt (4.14) ausgestattet. Dann heißt X eine JB*-Algebra (oder auch (unitale) Jordan C^* -Algebra, siehe EDWARDS [82](1980)), wenn für alle $x \in X$ die Gleichung $\|\{xxx\}\| = \|x\|^3$ gilt. Eine JBW*-Algebra (oder auch Jordan W^* -Algebra) ist eine JB*-Algebra, die ein Prädual besitzt.

Eine JC^* -Algebra ist eine abgeschlossene unital Jordan $*$ -Unteralgebra von A^Δ , A eine unital C^* -Algebra über $\mathbb{C}$.

Manche Autoren wie zum Beispiel ISIDRO, KAUP und RODRÍGUEZ PALACIOS in [147](1995), Seite 312, und KAUP [179](2007), Seite 50, definieren eine JC^* -Algebra als eine abgeschlossene $*$ -Unteralgebra von $\mathcal{L}(H)$, H ein Hilbertraum über $\mathbb{C}$, die bezüglich des reellen Jordanproduktes $(x, y) \mapsto \frac{1}{2}(xy + yx)$ (siehe Definition 3.2.14) abgeschlossen ist.

4.7.15 JB*-Algebra-Ordnung. Sei A eine JB^* -Algebra.

(a) (HANCHE-OLSEN und STØRMER [125](1984), Seite 82). Als partielle Ordnung $\leq$ auf der Menge $\text{Proj}(A)$ verwendet man die in 2.1.45 definierte Ordnung. Auf dem selbstadjungierten Teil $A_{(*)}$ von A gilt $\text{Pos}(*; A_{(*)}) = \text{Pos}(*\sigma; A_{(*)})$. (Falls man eine JB^* -Algebra nicht als unital definiert, gilt im Allgemeinen nur $\supseteq$.) $\text{Pos}(*; A_{(*)})$ ist ein punktierter, spitzer, konvexer Kegel in $A_{(*)}$, induziert also gemäß 2.2.12 auf $A_{(*)}$ eine partielle Ordnung.

(b) (EDWARDS, FERNÁNDEZ POLO, HOSKIN und PERALTA [83](2010), Seite 126). Eingeschränkt auf $\text{Proj}(A_{(*)})$ stimmt die von $\text{Pos}(*; A_{(*)})$ auf $A_{(*)}$ induzierte Ordnung überein mit der Ordnung von $\text{Proj}(A)$. (Als Mengen gilt natürlich $\text{Proj}(A_{(*)}) = \text{Proj}(A)$.)

(c) Für lineare Funktionale auf A wird gemäß 2.3.33 die Definition 3.2.24 mit der Menge $\text{Pos}(\text{idem}; A)$ — sprich, der Menge der idempotenten Elemente von A — anstatt mit $\text{Pos}(*; A)$ formuliert.

(d) (EDWARDS [82](1980), Seite 400). *Normale* lineare Funktionale auf $A_{(*)}\mathbb{R}$ werden ganz entsprechend wie in 3.7.37 definiert. Ein $\ell \in A^\#$ heißt *normal*, wenn $\text{Re}(\ell)$ und $\text{Im}(\ell)$ normal auf $A_{(*)}$ sind.

4.7.16 JB*-Algebra und JC*-Algebra (UPMEIER [279](1985), Seite 351). Mit dem Tripelprodukt, mit dem per Definition 4.7.14 eine JB^* -Algebra ausgestattet ist, ist jede JB^* -Algebra ein JB^* -Tripel über $\mathbb{C}$. Für JB^* -Tripel über $\mathbb{C}$, die ein unitäres Element haben, gilt auch die Umkehrung: Ist X ein JB^* -Tripel über $\mathbb{C}$ und $e \in X$ unitär, so wähle dazu als Produkt die e -Multiplikation und als Involution die Abbildung $Q(e)$, also $x \mapsto x^* := \{exe\}$; es liegt dann eine JB^* -Algebra mit e als Eins vor; siehe KAUP [175](1983), Seite 525 und auch UPMEIER [279](1985), Seite 320, Satz 19.13; des Weiteren siehe auch hier bei 4.7.20.

Sei X eine JC^* -Algebra. Dann ist X nach 4.7.11 ein JB^* -Tripel über $\mathbb{C}$ und daher nach Satz 4.7.7(a)(b) auch eine JB^* -Algebra.

Bezüglich der von KAUP [179](2007) definierten Version einer JC^* -Algebra X gilt: X ist offensichtlich eine Banach-Jordan-Algebra über $\mathbb{C}$; ist X zusätzlich unital und mit dem in 4.6.10 definierten Tripelprodukt (4.14) versehen, so ist X eine JB^* -Algebra. Des Weiteren stimmt dann das Tripelprodukt überein mit dem auf X eingeschränkten Tripelprodukt von $\mathcal{L}(H)^\Delta$ (ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seite 312). Bezeichnet $\circ$ das reelle Jordanprodukt von X , so ist X wegen $xyx = 2x \circ (x \circ y) - (x \circ x) \circ y$ auch ein Beispiel eines JC^* -Tripels über $\mathbb{C}$ (KAUP [179](2007), Seite 50).

4.7.17 Satz. Seien X und Y zwei JB^* -Tripel über $\mathbb{C}$. Dann gilt:

- (a) Sei $T \in L(X, Y)$ bijektiv. Dann ist T genau dann eine Isometrie, wenn T ein Tripel-Isomorphismus ist.
- (b) Sei $T \in \mathcal{L}(X)$. Dann ist T genau dann hermitesch, wenn iT eine Tripel-Derivation ist.
- (c) Für alle $x, y, z \in X$ gilt $\|\{xyz\}\| \leq \|x\|\|y\|\|z\|$, das heißt, für alle $x, y \in X$ gilt $\|D(x, y)\| \leq \|x\|\|y\|$.

Beweis. Die wenn-Aussage von (a) ist im Satz 4.7.9 enthalten. (b) und die genau-dann-Aussage von (a) findet man bei KAUP [175](1983), Seite 523. Für einen Beweis per Tripotenten der genau-dann-Aussage von (a) siehe FERNÁNDEZ-POLO et al. [94](2004), Seite 439, Theorem 2.2; dort wird diese Aussage auch *Kaup's Banach-Stone-Theorem* genannt. (c) ist ein Korollar des Gelfand-Naimark-Theorems für JB^* -Tripel über $\mathbb{C}$, siehe FRIEDMAN und RUSSO [101](1986), Seite 145. An dieser Stelle sei bemerkt, dass KAUP [175](1983) bereits $\|D(x, y)\| \leq 2\|x\|\|y\|$ für alle $x, y \in X$ zeigte. $\square$

4.7.18 JB^* -Algebra. Da hier JB^* -Algebren per Definition unital sind, folgt, dass die per Definition vorhandene Involution einer JB^* -Algebra stets eine Isometrie ist.

Beweis. Sei X eine JB^* -Algebra und $*$ ihre Involution. Sei $x \in X$ beliebig. Nach Gleichung (4.14) in 4.6.10 gilt $x^* = Q(e)x$. Nach 4.7.16 ist X ein JB^* -Tripel über $\mathbb{C}$ und per der somit möglichen Anwendung des Satzes 4.7.17(c) hat man $\|x^*\| \leq \|e\|\|x\|\|e\| = \|x\|$. Analog demnach auch $\|x\| = \|x^{**}\| \leq \|x^*\|$. $\square$

Definiert man allgemeiner eine JB^* -Algebra ohne die Forderung der Unitalität, so wird die Involution zusätzlich explizit als isometrisch gefordert.

4.7.19 Reell-lineare Isometrie. Die entsprechende Aussage von Satz 4.7.17(a) für $\mathbb{R}$ -lineare Isometrien stimmt im Allgemeinen nicht. Für das dafür Gültige siehe DANG [57](1992): Satz 2.6 und Theorem 3.1.

Zumindest gilt aber (ebd.): Seien X und Y zwei JB^* -Tripel über $\mathbb{C}$ und sei $T: X \rightarrow Y$ eine $\mathbb{R}$ -lineare surjektive Isometrie. Dann gilt für alle $x \in X$ die Gleichung $T\{xxx\} = \{(Tx)(Tx)(Tx)\}$. Von den weiter unten definierten $*JB$ -Tripeln weiß man, dass auch die Umkehrung gilt, siehe Satz 4.7.54.

4.7.20. Wegen Satz 4.7.17(c) bemerke man: Ist X ein JB^* -Tripel über $\mathbb{C}$ und $a \in S_X$, so ist X , versehen mit der a -Multiplikation, eine Banach-Jordan-Algebra über $\mathbb{C}$. Siehe auch 4.6.10, 4.7.16 und 4.7.38.

4.7.21 (KAUP [178](2002)). Die beschränkten, symmetrischen Gebiete (mit Basispunkt) von Banachräumen über $\mathbb{C}$ bilden eine Kategorie $\mathcal{D}$: Sei Ω ein beschränktes, symmetrisches Gebiet eines Banachraumes über $\mathbb{C}$. Bezeichne $s_a, a \in \Omega$, die Symmetrien in Ω bei $a \in \Omega$. Dann ist per

$$\mu: \Omega \times \Omega \rightarrow \Omega, (y, z) \mapsto yz := s_y(z) \quad (4.15)$$

eine Multiplikation auf Ω definiert. Ein Morphismus der Kategorie ist eine holomorphe Abbildung $\varphi: \Omega_1 \rightarrow \Omega_2$ zwischen beschränkten, symmetrischen Gebieten (mit Basispunkt) für die $\varphi(yz) = \varphi(y)\varphi(z)$ für alle $y, z \in \Omega_1$ gilt (und die die Basispunkte aufeinander abbildet).

Dann ist jede biholomorphe Abbildung $\varphi: \Omega_1 \rightarrow \Omega_2$ zwischen beschränkten, symmetrischen Gebieten ein Isomorphismus der Kategorie.

Man kann zeigen, dass ein Morphismus zwischen symmetrischen offenen Einheitskugeln von Banachräumen, der die Ursprünge aufeinander abbildet, gleich der Restriktion einer kontraktiven, linearen Abbildung zwischen den besagten Banachräumen ist.

Bezeichne $\mathcal{E}$ die Kategorie aller Banachräume über $\mathbb{C}$ mit linearen Kontraktionen als Morphismen. Unter der in 4.5.40 gemachten Voraussetzung hat man dann den folgenden Funktor: $\mathcal{F}: \mathcal{D} \rightarrow \mathcal{E}, \Omega \mapsto \mathcal{F}(\Omega) := T_p(\Omega)$ für jedes beschränkte, symmetrische Gebiet $\Omega \in \mathcal{D}$ mit Basispunkt $p, \varphi \mapsto \mathcal{F}(\varphi) := T_p(\varphi)$ für jeden Morphismus φ in $\mathcal{D}$. Da jede Einheitskugel im Ursprung die Symmetrie $x \mapsto -x$ erlaubt, besteht wegen der in 4.5.17 erwähnten Äquivalenz das Bild des Funktors $\mathcal{F}$ aus genau den Banachräumen über $\mathbb{C}$, dessen offene Einheitskugel homogen ist. (STACHÓ und ISIDRO [266](1990), Seite 172, nennen solche Banachräume *Jordan-Banachräume*, abgekürzt *JB-Räume*.) Siehe auch 4.5.34.

Die offene Einheitskugel des einem JB^* -Tripel über $\mathbb{K}$ zugrunde liegenden Banachraumes über $\mathbb{K}$ ist symmetrisch (UPMEIER [279](1985), Korollar 20.13). Im Fall von $\mathbb{K} = \mathbb{C}$ gilt auch die Umkehrung; damit hat man: Ein Banachraum über $\mathbb{C}$ ist genau dann ein JB^* -Tripel über $\mathbb{C}$, wenn seine offene Einheitskugel symmetrisch (oder wegen 4.5.17 dazu äquivalent: homogen) ist; das Tripelprodukt ist dann eindeutig bestimmt.

Das Bild $\mathcal{F}(\mathcal{D})$ besteht also genau aus allen JB^* -Tripeln über $\mathbb{C}$. Im Detail kann man dabei die Werte von $\mathcal{F}$ wie folgt Jordan-algebraisch charakterisieren: Sei Ω ein beschränktes, symmetrisches Gebiet eines Banachraumes über $\mathbb{C}$. Weiterhin gemäß der Anmerkung 4.5.40 sei Ω als die offene Einheitskugel eines Banachraumes X über $\mathbb{C}$ realisiert. Die auf Ω in (4.15) definierte Multiplikation μ ist eine $\mathbb{R}$ -analytische Abbildung und die dritte partielle Ableitung von μ bei $(0, 0) \in \Omega \times \Omega$

$$\Lambda := \frac{\partial^3}{\partial z \partial y \partial z} \left((y, z) \mapsto \mu(y, z) \right) (0, 0)$$

ist eine $\mathbb{R}$ -trilineare Abbildung $X^3 \rightarrow X$, die in den äußeren Variablen kommutativ und $\mathbb{C}$ -bilinear ist. Mit

$$\{xyz\} := -\frac{1}{4}\Lambda(x, y, z)$$

als Jordan-Tripelprodukt ist dann X ein JB^* -Tripel über $\mathbb{C}$.

4.7.22 (KAUP [178](2002) und [179](2007)). Die JB^* -Tripel über $\mathbb{C}$ bilden eine Kategorie. Die Morphismen dieser Kategorie sind die Tripel-Homomorphismen. Jeder solcher Morphismus $T: X \rightarrow Y$ ist kontraktiv, also insbesondere stetig, hat in Y ein abgeschlossenes Bild und induziert eine lineare Isometrie von $X/\ker(T)$ auf $T(X)$. Für jedes abgeschlossene Ideal I eines JB^* -Tripels X über $\mathbb{C}$ ist der Quotient X/I in kanonischer Weise wieder ein JB^* -Tripel über $\mathbb{C}$. Die Kategorie der JB^* -Tripels über $\mathbb{C}$ ist abgeschlossen bezüglich der Bildung von ℓ_∞ -Summen im Sinne von 4.6.21.

Ganz mit 4.7.21 hat man den folgenden Satz, der bereits in KAUP [175](1983), aufbauend auf die gleich hier in 4.7.26 erwähnte ursprüngliche Funktorkonstruktion, bewiesen wurde:

4.7.23 Satz. *Die Kategorie $\mathcal{D}$ ist per $\mathcal{F}$ äquivalent zu der Kategorie der JB^* -Tripel über $\mathbb{C}$.*

$\mathcal{F}$ heißt daher auch der *Jordan-Funktor*.

4.7.24 (DINEEN [70](1986), Seite 133, Satz 4; ARAZY [9](1987); siehe auch CHU und MELLON [50](1998)). Sei X ein Banachraum X über $\mathbb{C}$. X ist genau dann ein JB^* -Tripel über $\mathbb{C}$, wenn der symmetrische Teil X_{sym} von X gleich ganz X ist, das heißt, wenn in der Situation von 4.5.29 die Gleichung $X = \left\{ h(0) : h \frac{\partial}{\partial z} \in \mathfrak{p} \right\}$ gilt. In Verallgemeinerung von 4.5.37 definiert man ein sogenanntes *partiell Jordan-Tripelprodukt* auf $X \times X_{\text{sym}} \times X$ per $\{xyz\} := \tilde{q}_y(x, z)$. X_{sym} ist der größte Unterraum von X , dessen offene Einheitskugel ein beschränktes, symmetrisches Gebiet ist; mit anderen Worten ist X_{sym} das größte in X liegende JB^* -Tripel über $\mathbb{C}$. Es sei an Satz 4.5.30 erinnert.

Man sagt (ARAZY, ebd., Seite 145), X habe die *lineare biholomorphe Eigenschaft* (im Englischen *linear biholomorphic property*, *LBP*), falls $\text{Aut}(\text{int}B_X) \subseteq L(X)$ gilt. X hat genau dann die LBP, wenn $h \in L(X)$ für alle $h \frac{\partial}{\partial z} \in \text{aut}(\text{int}B_X)$ gilt und dies wiederum ist äquivalent zu $X_{\text{sym}} = \{0\}$.

4.7.25 (MEYBERG [212](1972), Seite 43). Sei X ein allgemeines Tripel über $\mathbb{K}$. X heißt ein *Lie-Tripel* über $\mathbb{K}$, wenn gilt:

- (a) $\{xyx\} = 0$ für alle $x, y \in X$;
- (b) $\{xyz\} + \{yzx\} + \{zxy\} = 0$ für alle $x, y, z \in X$; (Jacobi-Gleichung)
- (c) $\{uv\{xyz\}\} = \{\{uvx\}yz\} + \{x\{uvy\}z\} + \{xy\{uvz\}\}$ für alle $u, v, x, y, z \in X$.

Ist X ein Lie-Tripel über $\mathbb{K}$, so ist es üblich, die ansonsten mit geschweiften Klammern geschriebene Abbildung $\{ \}$ mit eckigen Klammern, also als $[\]$ zu schreiben. $[\]$ nennt man das *Lie-Tripel-Produkt* des Lie-Tripels X .

4.7.26 (KAUP [178](2002) und [179](2007); OUCHEVA [228](2001)). Anders als die in 4.7.21 per der dritten partiellen Ableitung beschriebene Konstruktion des Jordan-Funktors, geht die ursprüngliche Konstruktion zurück auf die in KAUP [173](1977) durchgeführte Konstruktion des Funktors von der Kategorie der einfach-zusammenhängenden, symmetrischen Banach-Mannigfaltigkeiten über $\mathbb{C}$ mit Basispunkt auf die Kategorie der J^* -Tripel über $\mathbb{C}$, angewandt auf den Spezialfall, dass man als symmetrische Banach-Mannigfaltigkeit über $\mathbb{C}$ mit Basispunkt ein beschränktes, symmetrisches Gebiet mit Basispunkt eines Banachraumes über $\mathbb{C}$ betrachtet:

Sei $\Omega \in \mathcal{D}$ mit Basispunkt p , Ω also ein beschränktes, symmetrisches Gebiet eines Banachraumes Y über $\mathbb{C}$ mit $p \in \Omega$. Über die Identifikation von $T_p(\Omega)$ mit Y renormiere man den Banachraum Y zu einem Banachraum X , indem man $T_p(\Omega)$ mit der Carathéodory-Norm bezüglich p ausstattet. Damit ist die Norm von X äquivalent zu der Norm von Y . Die folgenden Aussagen beziehen sich auf den Banachraum X . Sei K_p die *Isotropie-Untergruppe* in p , soll heißen,

$$K_p := \{g \in \text{Aut}_{MF}(\Omega) : g(p) = p\}.$$

Für jedes $g \in K_p$ ist $Dg(p) \in GL(X)$ eine Isometrie. Die Symmetrie $s_p \in K_p$ in p ist im Zentrum von K_p (siehe 4.5.17 und Definition 4.3.9 und betrachte für alle $h \in K_p$ die Fixpunktgleichung $\text{Int}(h)s_p = s_p$ in den beiden Formulierungsvarianten $hs_p = s_ph$ und $s_{h(p)} = hs_ph^{-1}$). Nach 4.5.28 und 4.3.10 operiert die Symmetrie s_p per ihrer adjungierten Darstellung $\text{Ad}(s_p)$ der Banach-Lie-Gruppe $\text{Aut}_{MF}(\Omega)$ auf deren Lie-Algebra $\mathfrak{g}(\text{Aut}_{MF}(\Omega))$, die, wie bereits in 4.5.28 erwähnt, mit $\text{aut}(\Omega)$ identifiziert werden kann, das heißt, wenn man für alle $g \in \text{Aut}_{MF}(\Omega)$ und für alle $m \in \Omega$ kanonisch den Ausdruck $L_g(m)$ als $g(m)$ und den Ausdruck $T(L_g)$ als $T(g)$ interpretiert, dass speziell für $g = s_p$ gilt: $\text{Ad}(s_p) = (s_p)_*|_{\text{aut}(\Omega)} \in \text{Aut}(\text{aut}(\Omega))$.

Wegen $s_p^2 = \text{Id}_\Omega$ gilt $(\text{Ad}(s_p))^2 = \text{Ad}(s_p^2) = \text{Ad}(\text{Id}_\Omega) = \text{Id}_{\text{aut}(\Omega)}$. Wendet man demnach $\text{Ad}(s_p)$ in der Eigenwertgleichung $\text{Ad}(s_p)V = \lambda V$, $\lambda \in \mathbb{C}$, $V \in \text{aut}(\Omega)^\times$, an, so sieht man, dass $\text{Ad}(s_p)$ die beiden Eigenwerte $\lambda = -1$ und $\lambda = +1$ hat. Man kann also die Lie-Algebra $\text{aut}(\Omega)$ mittels der beiden Eigenräume

$$\mathfrak{p} := \{V \in \text{aut}(\Omega) : \text{Ad}(s_p)V = -V\}$$

und

$$\mathfrak{k} := \{V \in \text{aut}(\Omega) : \text{Ad}(s_p)V = V\}$$

als

$$\text{aut}(\Omega) = \mathfrak{p} \dot{+} \mathfrak{k}$$

darstellen. $\mathfrak{k}$ kann identifiziert werden mit der zu der Banach-Lie-Untergruppe $K_p \subseteq \text{Aut}(\Omega)$ korrespondierenden Lie-Unteralgebra von $\text{aut}(\Omega)$; des Weiteren gilt $\mathfrak{k} = \{V \in \text{aut}(\Omega) : V(p) = 0\}$, $[\mathfrak{k}, \mathfrak{p}] \subseteq \mathfrak{p}$ und $[\mathfrak{p}, \mathfrak{p}] \subseteq \mathfrak{k}$. $\mathfrak{p}$ ist ein Lie-Tripel mit dem Lie-Tripel-Produkt $[xyz] := [[x, y], z]$ für alle $x, y, z \in \mathfrak{p}$. Die Auswertungsabbildung $\text{aut}(\Omega) \rightarrow X$, $h(z) \xrightarrow{\frac{\partial}{\partial z}} h(p)$ ist nach 4.5.34 surjektiv und gibt somit einen $\mathbb{R}$ -linearen Isomorphismus von $\mathfrak{p}$ auf $X_{\mathbb{R}}$. Damit hat man auf X ein Lie-Tripel-Produkt $(x, y, z) \mapsto [xyz]$

erhalten, welches $\mathbb{R}$ -trilinear und $\mathbb{C}$ -linear in z ist. Dann ist das Jordan-Tripelprodukt definiert als der Teil von $(x, y, z) \mapsto \frac{1}{2}[xyz]$, welcher konjugiert-linear in y ist (oder dazu äquivalent: welcher linear in x ist), genauer:

$$\{xyz\} := \frac{1}{4}([xyz] + i[x(iy)z]) \quad \left(= \frac{1}{4}([xyz] - i[(ix)yz]) \right) \quad \text{für alle } x, y, z \in X.$$

X , ausgestattet mit diesem Tripelprodukt, ist ein JB^* -Tripel über $\mathbb{C}$. (Wegen $[\mathfrak{p}, \mathfrak{p}] \subseteq \mathfrak{k}$ hat man mittels $\mathfrak{p} \cong X_{\mathbb{R}}$ auf X das $\mathfrak{k}$ -wertige Klammerprodukt $[\cdot, \cdot]: X \times X \rightarrow \mathfrak{k}$ und für jedes $x \in X$, $t \in \mathbb{R}$ ist der Operator $itD(x)$ die Ableitung des Automorphismus $\exp(t[x, ix]) \in K_p$.)

Wie bereits in 4.5.40 erwähnt, kann man Ω mit der offenen Einheitskugel von X identifizieren, wobei der Basispunkt von Ω dem Ursprung von X entspricht.

Man kann den Jordan-Funktor auch ohne Verwendung von Lie-Gruppen oder einer Realisation als offene Einheitskugel konstruieren, womit dann die Abhängigkeit des Jordan-Tripelproduktes von der komplexen Struktur des Tangentialraumes $T_p(\Omega)$ vermieden werden kann.

4.7.27 (KAUP [176](1984); KAUP [179](2007), Seite 53). Sei X ein JB^* -Tripel über $\mathbb{C}$ und bezeichne $\{\cdot\}_X$ sein Tripelprodukt. Sei $P \in \mathcal{L}(X)$ eine Projektion mit $\|P\| = 1$. Auf dem Bild $Y := \text{ran}(P)$ sei das Tripelprodukt $\{\cdot\}_Y$ per $\{xyz\}_Y := P\{xyz\}_X$ für alle $x, y, z \in Y$ definiert; es wird das *projizierte Tripelprodukt* genannt. Im Allgemeinen ist dann zwar Y kein Untertripel von X . Aber es gilt die folgende Aussage: $(Y, \{\cdot\}_Y)$ ist ein JB^* -Tripel über $\mathbb{C}$. Für JB^* -Tripel über $\mathbb{R}$ gilt diese Aussage im Allgemeinen nicht, siehe das Gegenbeispiel in 4.7.63.

4.7.28 Bidual. Sei X ein JB^* -Tripel über $\mathbb{C}$, I eine beliebige Indexmenge und $\mathcal{U}$ ein Ultrafilter auf I . Setze $X_\nu := X$ für alle $\nu \in I$ und bezeichne Y die ℓ_∞ -Summe $\ell_\infty\left(\bigoplus_{\nu \in I} X_\nu\right)$ im Sinne von 4.6.21. DINEEN [71](1986) bemerkt zunächst, dass wegen der Stetigkeit der Abbildung $D: Y \times Y \rightarrow Y$ die Teilmenge

$$N_{\mathcal{U}} := \left\{ (x_\nu)_{\nu \in I} \in Y : \lim_{\mathcal{U}} \|x_\nu\| = 0 \right\},$$

ein abgeschlossenes Tripelideal des JB^* -Tripels Y über $\mathbb{C}$ ist, also die Ultrapotenz $X_{\mathcal{U}}$ ein JB^* -Tripel über $\mathbb{C}$ ist und folgert durch eine einfache Kombination bestehender Ergebnisse, dass der Bidualraum X^{**} ein JB^* -Tripel über $\mathbb{C}$ ist. Derselbe Autor hat dies auch in DINEEN [70](1986) gezeigt, dort aber aufwendiger unter Verwendung von Vektorfeldern mittels 4.7.24, konnte dafür aber auch folgern, dass die kanonische Inklusionsabbildung i_X eines JB^* -Tripels X über $\mathbb{C}$ ein isometrischer Tripel-Isomorphismus auf ein Norm-abgeschlossenes w^* -dichtes Untertripel von X^{**} ist, das heißt, jedes $f \in \text{Aut}_{MF}(\text{int}(B_X))$ setzt sich zu einem $\tilde{f} \in \text{Aut}_{MF}(\text{int}(B_{X^{**}}))$ fort.

Der Beweis in DINEEN [70](1986) der Aussage, dass mit X auch der Bidualraum X^{**} ein JB^* -Tripel über $\mathbb{C}$ ist, fußt auf der Ultrapotenz-Formulierung des Prinzips der lokalen Reflexivität; bezüglich dieser Formulierung siehe HEINRICH [130](1980), Seite 86, Satz 6.7. BARTON und TIMONEY [12](1986) haben die Ultrafilter-Argumentation in Dineen's Beweis ausgebaut und zeigen dann damit, dass das Tripelprodukt von X^{**} getrennt w^* -stetig ist. Damit zeigen sie weiter, dass jedes JB^* -Tripel X über $\mathbb{C}$, dessen zugrunde liegender Banachraum ein Dualraum ist, ein Prädual besitzt, bezüglich dessen das Tripelprodukt von X getrennt w^* -stetig ist; siehe dazu aber HORN [140](1987), Bemerkung 3.5 und ebd. Seite 127.

HORN, ebd., zeigt:

(a) Jedes JB^* -Tripel X über $\mathbb{C}$, dessen zugrunde liegender Banachraum ein Dualraum ist, das ein Prädual besitzt, bezüglich dessen das Tripelprodukt von X getrennt w^* -stetig ist, hat ein stark eindeutiges Prädual (ebd., Theorem 3.21).

(b) Jedes JB^* -Tripel X über $\mathbb{C}$, das ein stark eindeutiges Prädual besitzt, hat ein getrennt w^* -stetiges Tripelprodukt (ebd., Satz 3.24).

Folglich haben BARTON und TIMONEY [12](1986) als Theorem, dass jedes JB^* -Tripel X über $\mathbb{C}$, dessen zugrunde liegender Banachraum ein Dualraum ist, ein stark eindeutiges Prädual besitzt und, dass das Tripelprodukt von X getrennt w^* -stetig ist.

Es liegt also eine Situation wie bei den C^* -Algebren vor, siehe die Sätze 3.7.22 und 3.7.42. Dementsprechend definiert man:

4.7.29 Definition. Ein JB^* -Tripel über $\mathbb{C}$ heißt ein JBW^* -Tripel über $\mathbb{C}$, wenn der zugrunde liegende Banachraum ein Dualraum ist.

4.7.30. Somit ist wegen Satz 3.7.22 offensichtlich jede W^* -Algebra über $\mathbb{C}$ ein JBW^* -Tripel über $\mathbb{C}$. Siehe 4.7.50 für den reellen Fall.

4.7.31. Aus 4.7.28 zusammen mit dem Bipolarensatz 2.5.6 und Satz 2.4.36 folgt: Jedes schwach*-abgeschlossene Untertripel eines JBW^* -Tripels über $\mathbb{C}$ ist ein JBW^* -Tripel über $\mathbb{C}$. Siehe Korollar 4.7.51 für einen reellen analogen Fall.

4.7.32 (STACHÓ [265](2007), Seite 3). Das JB^* -Tripel $\left(C_0([-1, 1], \mathbb{C})\right)^\Delta$ hat keine von Null verschiedenen Tripotenten.

4.7.33. Ähnlich wie in Satz 3.7.45(b) hat man: Sei X ein JBW^* -Tripel über $\mathbb{C}$. Dann gilt $X = \text{span Tri}(X)$. (Man beachte, dass hier kein cl nötig ist.)

4.7.34 Ideale. Sei X ein JB^* -Tripel über $\mathbb{C}$. Dann gelten die folgenden Aussagen.

(a) Jeder abgeschlossene Untervektorraum U von X mit $\{XUU\} \subseteq U$ ist ein Tripelideal.

(b) Ist I ein Tripelideal von X , so ist $I^\perp$ (im Sinne von 4.6.3(h)) ein abgeschlossenes Tripelideal von X .

(c) Genau die abgeschlossenen Tripelideale von X sind die M -Ideale von X .

(d) Falls X ein JBW^* -Tripel über $\mathbb{C}$ ist, dann sind genau die schwach*-abgeschlossenen Tripelideale von X die M -Summanden von X . Siehe auch 4.7.56.

Beweis. (a) Siehe Zitierung „[6]“ in BUNCE, CHU und ZALAR [40](2000), Seite 18. (b) BUNCE, CHU und ZALAR [40](2000), Seite 18. (c) BARTON und TIMONEY [12](1986), Seite 187, Theorem 3.2. (d) HORN [140](1987), Theorem (4.2)(4), Lemmata (4.3) und (4.4). $\square$

4.7.35 Satz (DANG und RUSSO [58](1994), Theorem 1.2). Sei X ein Banach-Jordan-Tripel über $\mathbb{C}$, das die folgenden drei Bedingungen erfüllt:

- (a) $\|\{xxx\}\| = \|x\|^3$ für alle $x \in X$;
- (b) $\|\{xyz\}\| \leq \|x\|\|y\|\|z\|$ für alle $x, y, z \in X$;
- (c) $\sigma(D(x)) \subseteq \mathbb{R}^+$ für alle $x \in X$.

Dann ist X ein JB^* -Tripel über $\mathbb{C}$.

4.7.36. Die JB^* -Tripel sind ursprünglich als JB^* -Tripel über $\mathbb{C}$ entstanden. Die hier gegebene, von UPMEIER [279](1985) stammende Definition von JB^* -Tripeln umfasst auch direkt ein reelles Analogon, nämlich die JB^* -Tripel über $\mathbb{R}$. Es gibt aber noch andere Möglichkeiten, reelle Analoga von JB^* -Tripeln über $\mathbb{C}$ zu definieren, siehe unten die Definition 4.7.39. Dazu vorab folgende Definition:

4.7.37 Definition (DANG und RUSSO [58](1994)). Ein J^*B -Tripel ist ein Banach-Jordan-Tripel X über $\mathbb{R}$, das die folgenden vier Bedingungen erfüllt.

- (a) $\|\{xxx\}\| = \|x\|^3$ für alle $x \in X$.
- (b) $\|\{xyz\}\| \leq \|x\|\|y\|\|z\|$ für alle $x, y, z \in X$.
- (c) $\sigma(D(x)) \subseteq \mathbb{R}^+$ für alle $x \in X$.
- (d) $\sigma(D(x, y) - D(y, x)) \subseteq i\mathbb{R}$ für alle $x, y \in X$.

4.7.38. Es sei an die Aussage 4.7.20 erinnert. Analog gilt hier wegen 4.7.37(b): Ist X ein J^*B -Tripel und $a \in S_X$, so ist X , versehen mit der a -Multiplikation, eine Banach-Jordan-Algebra über $\mathbb{R}$.

In der einschlägigen Literatur sind hauptsächlich drei reelle Analoga von JB^* -Tripeln über $\mathbb{C}$ verbreitet:

4.7.39 Definition. (a) Ein *reelles JB^* -Tripel nach Upmeyer* ist ein JB^* -Tripel über $\mathbb{R}$. (UPMEIER [279](1985))

(b) Ein *reelles JB^* -Tripel nach Isidro, Kaup und Rodríguez Palacios*, im Folgenden durchweg als **JB -Tripel* bezeichnet, ist ein Banachraum über $\mathbb{R}$, versehen mit der Struktur eines Tripels über $\mathbb{R}$, für den ein JB^* -Tripel Y über $\mathbb{C}$ und ein isometrischer Tripel-Homomorphismus $T: X \rightarrow Y_{\mathbb{R}}$ existiert. (ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995))

(c) Ein *reelles JB^* -Tripel nach Russo* ist ein J^*B -Tripel. (DANG und RUSSO [58](1994))

4.7.40 Bemerkung. Nach dem Erscheinen des Artikels ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995) versteht man in der Literatur üblicherweise unter reellen JB^* -Tripeln die hier in 4.7.39(b) definierten *JB -Tripel und unter JB^* -Tripeln ohne einer Angabe eines Skalarenkörpers weiterhin wie gehabt die hier definierten JB^* -Tripel über $\mathbb{C}$.

4.7.41 Untertripel und reelle Strukturierung. Ist somit X ein JB^* -Tripel über $\mathbb{C}$, so ist $X_{\mathbb{R}}$ ein *JB -Tripel und jedes abgeschlossene Untertripel von $X_{\mathbb{R}}$ ist ein *JB -Tripel. Siehe auch Bemerkung 4.7.55. Nach Bemerkung 4.7.6 ist jedes *JB -Tripel ein JB^* -Tripel über $\mathbb{R}$.

4.7.42 J^*B -Tripel (DANG und RUSSO [58](1994)). Ein Banach-Jordan-Tripel X über $\mathbb{C}$ ist genau dann ein JB^* -Tripel über $\mathbb{C}$, wenn $X_{\mathbb{R}}$ ein J^*B -Tripel ist (ebd., Satz 1.4). Jedes abgeschlossene Untertripel eines J^*B -Tripels ist ein J^*B -Tripel. Insbesondere ist somit für jedes JB^* -Tripel X über $\mathbb{C}$ jedes abgeschlossene Untertripel von $X_{\mathbb{R}}$ ein J^*B -Tripel (ebd., Bemerkung 1.5). Für jede C^* -Algebra X über $\mathbb{R}$ ist X^{Δ} ein J^*B -Tripel (ebd., Seite 137).

4.7.43 Konjugation auf komplexen JB^* -Tripeln (ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seiten 316 und 317). Sei Y ein JB^* -Tripel über $\mathbb{C}$, auf dessen zugrunde liegenden normierten Raum über $\mathbb{C}$ eine Konjugation $*$ definiert sei (Definition 3.1.9). Mit Y aufgefasst als Banachraum wird mit $\bar{Y}$ gemäß Bemerkung 2.9.7 der zu Y konjugiert-komplexe Banachraum bezeichnet. Sei κ die identische Abbildung von Y auf $\bar{Y}$; sie ist konjugiert-linear und isometrisch. Man setze kanonisch

$$\{xyz\} := \kappa\{\kappa^{-1}x, \kappa^{-1}y, \kappa^{-1}z\} \quad \text{für alle } x, y, z \in \bar{Y}; \quad (4.16)$$

ausgestattet mit diesem Produkt ist auch $\bar{Y}$ ein JB^* -Tripel über $\mathbb{C}$. Mit

$$y^* := \kappa\left(\left(\kappa^{-1}y\right)^*\right) \quad \text{für alle } y \in \bar{Y} \quad (4.17)$$

ist $*$ eine Konjugation auf dem normierten Raum $\overline{Y}$. Sei T die Abbildung

$$T: Y \rightarrow \overline{Y}, y \mapsto \kappa(y^*) \quad \left(= (\kappa(y))^* \right).$$

Da T eine surjektive lineare Isometrie zwischen zwei JB^* -Tripeln über $\mathbb{C}$ ist, ist T nach Satz 4.7.17(a) ein Tripel-Isomorphismus. Somit gilt für alle $x, y, z \in Y$:

$$\kappa(\{xyz\}^*) = \{\kappa(x^*), \kappa(y^*), \kappa(z^*)\}, \quad (4.18)$$

also

$$\begin{aligned} \{xyz\}^* &= \kappa^{-1} \kappa(\{xyz\}^*) \\ &= \kappa^{-1} \{\kappa(x^*), \kappa(y^*), \kappa(z^*)\} && \text{(wegen (4.18))} \\ &= \kappa^{-1} \kappa \left\{ \kappa^{-1} \kappa(x^*), \kappa^{-1} \kappa(y^*), \kappa^{-1} \kappa(z^*) \right\} && \text{(wegen (4.16))} \\ &= \{x^*, y^*, z^*\}. \end{aligned}$$

Somit ist $*$ auf Y ein konjugiert-linearer Tripel-Automorphismus. Da dann für alle $x, y, z \in Y_{(*)}$ die Gleichung $\{xyz\}^* = \{xyz\}$ gilt, ist $Y_{(*)}$ nach 3.1.10 ein abgeschlossenes Untertripel des reellen JB^* -Tripels $Y_{\mathbb{R}}$ und somit einerseits per 4.7.11 ein JB^* -Tripel über $\mathbb{R}$ und andererseits ein $*JB$ -Tripel.

Ist umgekehrt $*$: $Y \rightarrow Y$ ein konjugiert-linearer Tripel-Automorphismus mit Periode 2, so ist T ein Tripel-Isomorphismus und wieder nach Satz 4.7.17(a) eine surjektive lineare Isometrie $Y \rightarrow \overline{Y}$, womit auch $*$ eine Isometrie ist, also eine Konjugation auf dem Banachraum Y über $\mathbb{C}$.

Man hat also (vergleiche auch mit Satz 4.7.17(a)): Ist X ein JB^* -Tripel über $\mathbb{C}$ und $T: X \rightarrow X$ eine konjugiert-lineare Abbildung mit Periode 2, so ist T genau dann eine Isometrie (also eine Konjugation), wenn T ein konjugiert-linearer Tripel-Isomorphismus ist.

Dementsprechend wird definiert: Eine *Konjugation* auf einem JB^* -Tripel X über $\mathbb{C}$ ist ein konjugiert-linearer Tripel-Automorphismus auf X mit Periode 2. Mit dieser Definition ist also eine Abbildung auf einem JB^* -Tripel über $\mathbb{C}$ genau dann eine Konjugation, wenn sie eine Konjugation auf dem zugrunde liegenden Banachraum über $\mathbb{C}$ ist.

Ein JB^* -Tripel Y über $\mathbb{R}$ heißt eine *reelle Form eines JB^* -Tripels* X über $\mathbb{C}$, falls es auf X eine Konjugation $*$ gibt, so dass $Y = X_{(*)}$ ist. Man beachte, dass jede solche reelle Form auch eine reelle Form des Y zugrunde liegenden Banachraumes über $\mathbb{C}$ ist (siehe 3.1.11) und damit wiederum auch eine reelle Form des Y zugrunde liegenden Vektorraumes über $\mathbb{C}$ (Definition 2.9.23). Für ein wie hier zu Anfang definiertes Y ist somit $Y_{(*)}$, aufgefasst als ein JB^* -Tripel über $\mathbb{R}$, eine reelle Form des JB^* -Tripels Y über $\mathbb{C}$.

4.7.44 Hermiteifizierung von $*JB$ -Tripeln (ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seite 317). Sei nun X ein $*JB$ -Tripel. Nach Definition gibt es ein JB^* -Tripel Y über $\mathbb{C}$, so dass X als ein abgeschlossenes Untertripel des reellen JB^* -Tripels $Y_{\mathbb{R}}$ auffassbar ist. Dann ist die äußere direkte Summe $Y \oplus_{\infty} \overline{Y}$ ein JB^* -Tripel über $\mathbb{C}$. Sei ι die identische Abbildung von der Menge X in die Menge Y und κ wie in 4.7.43 die identische Abbildung von Y auf $\overline{Y}$. Sei $X_{(\mathbb{C})}$ die Komplexifizierung von X gemäß Definition 2.9.22. Dann ist

$$T: X_{(\mathbb{C})} \rightarrow Y \times \overline{Y}, \quad (x, y) \mapsto (\iota(x) + i\iota(y), \kappa(\iota(x) - i\iota(y)))$$

eine lineare, injektive Abbildung.

Beweis. Sei $(x, y) \in X_{(\mathbb{C})}$. Dann ist $T(i(x, y)) = T(-y, x) = \left(-\iota(y) + i\iota(x), \kappa(-\iota(y) - i\iota(x)) \right) = \left(i^2\iota(y) + i\iota(x), \kappa(i^2\iota(y) - i\iota(x)) \right) = \left(i(\iota(x) + i\iota(y)), i\kappa(-\iota(y) + \iota(x)) \right) = iT(x, y)$. $\square$

X kann also in Form von $T(X)_{\mathbb{R}}$ als ein abgeschlossenes reelles Untertripel von $(Y \oplus_{\infty} \bar{Y})_{\mathbb{R}}$ aufgefasst werden. Setze $U := T(X_{(\mathbb{C})})$. Ist $(u, v) \in U$, so überzeugt man sich leicht, dass mit $x := \frac{1}{2}\iota^{-1}(u + \kappa^{-1}v)$ und $y := \frac{1}{2}\iota^{-1}i(\kappa^{-1}v - u)$ die Gleichung $T(x, y) = (u, v)$ erfüllt ist. Als Norm auf $X_{(\mathbb{C})}$ nehme man die von $Y \times \bar{Y}$ per T auf $X_{(\mathbb{C})}$ induzierte Norm, soll heißen: $\|(x, y)\| := \|T(x, y)\|$ für alle $(x, y) \in X_{(\mathbb{C})}$. Somit also $U \cong X_{(\mathbb{C})}$ und konkret gilt $\|(x, y)\| = \max\{\|x + iy\|_Y, \|x - iy\|_Y\}$ für alle $(x, y) \in X_{(\mathbb{C})}$. U ist ein abgeschlossenes komplexes Untertripel des JB^* -Tripels $Y \oplus_{\infty} \bar{Y}$, also ein JB^* -Tripel über $\mathbb{C}$. Seien (a, x) , (b, y) und (c, z) drei Elemente von $X_{(\mathbb{C})}$. Dann gilt (das „ ι “ wird hier zur besseren Lesbarkeit weggelassen): $\{T(a, x)T(b, y)T(c, z)\} = (\{a + ix, b + iy, c + iz\}, \{\kappa(a - ix), \kappa(b - iy), \kappa(c - iz)\}) = (A + iB, \kappa(A - iB))$ mit $A = \{abc\} + \{ayz\} - \{xbz\} + \{xyc\}$ und $B = \{abz\} - \{ayc\} + \{xbc\} + \{xyz\}$. Da A und B Elemente von X sind, ist $\{T(a, x)T(b, y)T(c, z)\}$ ein Element von U . Somit ist auf $X_{(\mathbb{C})}$ das folgende Tripelprodukt erklärt:

$$\{xyz\} := T^{-1}\{Tx, Ty, Tz\} \quad \text{für alle } x, y, z \in X_{(\mathbb{C})},$$

womit T ein Tripelhomomorphismus ist. Nach Bemerkung 4.7.6 ist somit $X_{(\mathbb{C})}$ ein JB^* -Tripel über $\mathbb{C}$. Die kartesische Involution auf $X_{(\mathbb{C})}$ (siehe 3.1.8) stellt sich auf U wie folgt dar: $T(\overline{(x, y)}) = T((x, -y)) = (\iota(x) - i\iota(y), \kappa(\iota(x) + i\iota(y)))$. Für die somit naheliegende Austausch-Abbildung

$$*: Y \times \bar{Y} \rightarrow Y \times \bar{Y}, \quad (u, \kappa(v)) \mapsto (v, \kappa(u))$$

gilt also $T(\overline{(x, y)}) = T(x, y)^*$ für alle $(x, y) \in X_{(\mathbb{C})}$; des Weiteren gilt somit $* \upharpoonright U = T \circ \bar{} \circ T^{-1} \upharpoonright U$ und $\bar{} = T^{-1} \circ * \circ T$. Die Austausch-Abbildung hat die folgenden drei Eigenschaften:

(a) Invariant auf U : Klar wegen $* \upharpoonright U = T \circ \bar{} \circ T^{-1} \upharpoonright U$, explizit: $(\iota(x) + i\iota(y), \kappa(\iota(x) - i\iota(y)))^* = (\iota(x) - i\iota(y), \kappa(\iota(x) + i\iota(y))) = (\iota(x) + i\iota(-y), \kappa(\iota(x) - i\iota(-y)))$.

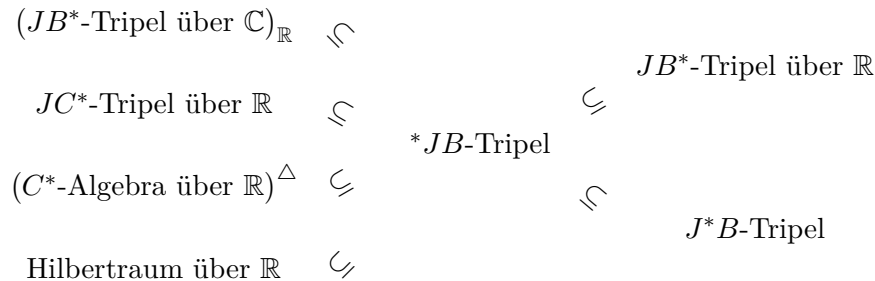
(b) Konjugiert-linear: Für alle $t \in \mathbb{C}$ und $u, v \in Y$ gilt $(t(u, \kappa(v)))^* = (tu, \kappa(\bar{t}v))^* = (\bar{t}v, \kappa(tu)) = \bar{t}(v, \kappa(u)) = \bar{t}(u, \kappa(v))^*$. (Dass sie es auf U ist, ist ja wieder wegen $* \upharpoonright U = T \circ \bar{} \circ T^{-1} \upharpoonright U$ klar.)

(c) Konjugiert-linearer Tripel-Homomorphismus: Klar, da das Tripelprodukt auf $Y \times \bar{Y}$ koordinatenweise definiert ist.

Die Abbildung $*$ auf $Y \times \bar{Y}$ ist also ein konjugiert-linearer Tripel-Automorphismus mit Periode 2 und nach 4.7.43 eine Konjugation auf dem JB^* -Tripel $Y \oplus_{\infty} \bar{Y}$ über $\mathbb{C}$. Da das mit der Komplexifizierung $X_{(\mathbb{C})}$ identifizierte Unter- JB^* -Tripel U über $\mathbb{C}$ nach (a) abgeschlossen unter der Abbildung $*$ ist, ist auf dem JB^* -Tripel U über $\mathbb{C}$ kanonisch eine Konjugation $*$ erklärt. Via der Abbildung T gilt für sie das Analogon der in 3.1.8 erwähnten Eigenschaft $X_{(\mathbb{C})}(\bar{})_{\mathbb{R}} = X$ der kartesischen Involution auf $X_{(\mathbb{C})}$, explizit: $U_{(*)} = \{(\iota(x) + i\iota(y), \kappa(\iota(x) - i\iota(y))) : x, y \in X, \iota(x) + i\iota(y) = \iota(x) - i\iota(y)\} = \{(\iota(x), (\kappa \circ \iota)(x)) : x \in X\} = T(X) = U$.

Zu jedem $*JB^*$ -Tripel X existiert somit auf $X_{(\mathbb{C})}$, versehen mit der kartesischen Involution $\bar{}$, eine eindeutige Struktur eines JB^* -Tripels über $\mathbb{C}$, mit der die kartesische Involution $\bar{}$ eine Konjugation ist, womit insbesondere $X_{(\mathbb{C})}$ eine Komplexifizierung nach LI ist. Dieses JB^* -Tripel $X_{(\mathbb{C})}$ über $\mathbb{C}$, versehen mit der kartesischen Konjugation, wird

4.7.45. Die Klasse der J^*B -Tripel umfasst alle *JB -Tripel (PERALTA [236] (2003), Seite 101, führt dazu ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147] (1995), Seite 318, Satz 2.5, an). Jeder Hilbertraum über $\mathbb{R}$ ist ein *JB -Tripel. Mit 4.7.11 und 4.7.41 hat man somit die folgenden Inklusionen:



PERALTA [236] (2003) zeigt für J^*B -Tripel X , in denen ein unitäres Element existiert: X ist genau dann ein *JB -Tripel, wenn in der Algebra $\mathcal{L}(X)$ der numerische Wertebereich die Inklusion $V(\text{Id}_X; D(X)) \subseteq \mathbb{R}^+$ erfüllt, und dies wiederum ist genau dann der Fall, wenn auf der Komplexifizierung des X zugrunde liegenden Vektorraumes die Struktur eines JB^* -Tripels über $\mathbb{C}$ existiert, deren Norm die von X fortsetzt.

4.7.47 Definition (ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seite 327). Ein $*JB$ -Tripel heißt ein $*JBW$ -Tripel, wenn es eine reelle Form eines JBW^* -Tripels über $\mathbb{C}$ (Definition 4.7.29) ist. (In ebd. wird so ein $*JBW$ -Tripel ein reelles JBW^* -Tripel genannt.)

4.7.49 Satz. Sei X ein *JB -Tripel. Dann sind äquivalent:

- (a) X ist ein *JBW -Tripel.
- (b) X ist ein w^* -abgeschlossenes Untertripel von der reellen Strukturierung eines JBW^* -Tripels über $\mathbb{C}$.
- (c) Der X zugrunde liegende Banachraum über $\mathbb{R}$ besitzt ein Prädual, bezüglich dessen das Tripel-Produkt auf X getrennt w^* -stetig ist.
- (d) Der X zugrunde liegende Banachraum über $\mathbb{R}$ besitzt ein Prädual.
- (e) Der X zugrunde liegende Banachraum über $\mathbb{R}$ besitzt ein Prädual, das ein stark eindeutiges Prädual ist.

Beweis. Die Äquivalenz von (a), (b) und (c) wird in ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seite 328, Theorem 4.4, gezeigt. MARTÍNEZ und PERALTA [206](2000), Seite 638, Satz 2.3, zeigen die Aussage (d) $\Rightarrow$ (e) und beweisen damit dann die Aussage (d) $\Rightarrow$ (c), ebd., Seite 643, Theorem 2.11; wie bei den JBW^* -Tripel über $\mathbb{C}$ kann also auf die getrennte w^* -Stetigkeit in (c) verzichtet werden. $\square$

4.7.50. Wie in 4.7.30 ist somit wieder wegen Satz 3.7.22 offenbar jede W^* -Algebra über $\mathbb{R}$ ein *JBW -Tripel.

Analog wie in 4.7.31 hat man das

4.7.51 Korollar (ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seite 328). *Jedes w^* -abgeschlossene Untertripel eines *JBW -Tripels ist ein *JBW -Tripel.*

4.7.52 Kartesische Involution. Sei X ein *JBW -Tripel. Betrachte die Hermitefizierung $X_{(\mathbb{C})}$ von X . Da dann die kartesische Involution von $X_{(\mathbb{C})}$ eine Isometrie ist, ist sie nach Korollar 2.10.9 schwach*-schwach*-stetig.

4.7.53 (ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seite 321). Sei X ein *JB -Tripel. Dann wird durch

$$\langle xyz \rangle := \frac{1}{3}(\{xyz\} + \{yzx\} + \{zxy\}) \quad \text{für alle } x, y, z \in X$$

ein Tripelprodukt auf X definiert; es wird das *symmetrisierte Jordan-Tripel-Produkt* genannt. Es ist eindeutig durch die Tatsache $\{xxx\} = \langle xxx \rangle$ für alle $x \in X$ bestimmt.

4.7.54 Satz (ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Seite 330, Theorem 4.8). *Seien X und Y zwei *JB -Tripel. Sei $T \in L(X, Y)$ bijektiv (es wird also T nicht als notwendig stetig verlangt). Bezeichne $\langle \rangle$ das symmetrisierte Jordan-Tripel-Produkt. Dann sind äquivalent:*

- (a) T ist eine Isometrie.
- (b) $T\{xxx\} = \{(Tx)(Tx)(Tx)\}$ für alle $x \in X$.
- (c) $T\langle xyz \rangle = \langle (Tx)(Ty)(Tz) \rangle$ für alle $x, y, z \in X$.

4.7.55 Bemerkung. Gemäß Satz 4.7.54 sind die *JB -Tripel genau die linearisometrischen Kopien von abgeschlossenen, reellen Untertripeln von JB^* -Tripeln über $\mathbb{C}$.

4.7.56 Ideale. Es sei an 4.7.34(d) erinnert. Sei X ein *JBW -Tripel und U ein schwach*-abgeschlossenes Ideal von X . Dann ist U ein M -Summand von X .

Beweis. (Nach einer Argumentation im Beweis eines Satzes in BECERRA GUERRERO, LÓPEZ PÉREZ, PERALTA und RODRÍGUEZ PALACIOS [16](2004), Seite 48). Betrachte das Annihilatorsystem der Orthogonalbildung, siehe 4.6.3(h). Nach ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147] (1995), Lemma 4.3, ist $U^\perp$ ebenfalls ein schwach*-abgeschlossenes Ideal von X und es gilt $X = U \dot{+} U^\perp$. Als abgeschlossene Untertripel von X sind U und $U^\perp$ jeweils *JB -Tripel (nach 4.7.51 sogar *JBW -Tripel). Somit ist auch $U \oplus_\infty U^\perp$ ein *JB -Tripel. Betrachte die bijektive Abbildung $S: U \oplus_\infty U^\perp \rightarrow X$, $(x, y) \mapsto x + y$. Sie ist ein Tripelisomorphismus: Seien $(a, x), (b, y), (c, z) \in U \oplus_\infty U^\perp$. Dann gilt einerseits $T\{(a, x)(b, y)(c, z)\} = T(\{abc\}, \{xyz\}) = \{abc\} + \{xyz\}$ und andererseits $\{T(a, x)T(b, y)T(c, z)\} = \{a + x, b + y, c + z\} = \{abc\} + \{abz\} + \{ayc\} + \{ayz\} + \{xbc\} + \{xbz\} + \{xyc\} + \{xyz\}$, was aber wegen $b \perp z, a \perp y, x \perp b$ und $y \perp c$ auch gleich $\{abc\} + \{xyz\}$ ist, das heißt, S ist ein Tripelisomorphismus. Nach Satz 4.7.54 ist somit S eine Isometrie und U damit ein M -Summand von X . $\square$

4.7.57 Prädual (BARTON und TIMONEY [12](1986), Seite 188, Satz 3.4; BECERRA GUERRERO, LÓPEZ PÉREZ, PERALTA und RODRÍGUEZ PALACIOS [16](2004), Seite 48, Satz 2.2).

Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Dann ist das Prädual von X ein L -Summand von X^* , mit anderen Worten, es ist L -eingebettet.

Der Beweis dieser Aussage baut im Wesentlichen auf den beiden Fakten auf, dass einerseits das Tripelprodukt auf X getrennt schwach*-stetig ist (4.7.28(b) und 4.7.49) und andererseits jedes schwach*-abgeschlossene Ideal von X ein M -Summand von X ist (4.7.34(d) und 4.7.56).

4.7.58 Spin-Faktoren. (a) (KAUP [174](1981), Seite 474; UPMEIER [280] (1987), Seite 6; KAUP [177](1997), Seite 199). Sei n eine beliebige Kardinalzahl und $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{C}$ mit Hilbertraumdimension n . Sei $*$ eine Konjugation auf H . H , versehen mit dem Tripelprodukt

$$\{xyz\} := \langle z, y \rangle x - \langle x, z^* \rangle y^* + \langle x, y \rangle z$$

(es sei an das Tripelprodukt (4.14) in 4.6.10 erinnert) und der (stets zu der gegebenen Hilbertraumnorm $\|\cdot\|_2$ äquivalenten) Norm

$$\|x\| := \sqrt{\langle x, x \rangle + \sqrt{\langle x, x \rangle^2 - |\langle x, x^* \rangle|^2}}$$

ist ein JBW^* -Tripel über $\mathbb{C}$; für $n \geq 3$ heißt es ein (*komplexer*) *Spin-Faktor* der Dimension n oder auch ein *Cartan-Faktor* IV_n . Auf $H_{(*)}$ gilt $\|\cdot\|_2 = \|\cdot\|$.

(b) (UPMEIER [279](1985), Seite 352). Sei $(H, \langle \cdot, \cdot \rangle)$ ein Hilbertraum über $\mathbb{R}$. Setze $X := \mathbb{R} \oplus H$. Dann ist auf dem Banachraum X über $\mathbb{R}$ per

$$(\alpha, x)(\beta, y) := \frac{1}{\sqrt{2}}(\alpha\beta + \langle x, y \rangle, \alpha y + \beta x)$$

ein kommutatives Produkt erklärt. Versehen mit diesem Produkt und der Norm

$$\|(\alpha, x)\| := |\alpha| + \|x\|$$

ist X eine Banach-Jordan-Algebra über $\mathbb{R}$ mit Eins $e := (\sqrt{2}, 0)$; sie heißt ein (*reeller*) *Spin-Faktor*. X , versehen mit dem Tripelprodukt

$$\{xyz\} := \frac{1}{2}(\langle x, y \rangle z + \langle z, y \rangle x - \langle x, z^* \rangle y^*)$$

ist ein JB^* -Tripel über $\mathbb{R}$.

4.7.59 Zerlegungen von Jordan-Tripeln. Sei X ein allgemeines Jordan-Tripel über $\mathbb{K}$. Sei $p \in \text{Tri}(X)$. Mögliche Eigenwerte des Operators $D(p)$ sind 0, $\frac{1}{2}$ und 1. Somit existiert stets die Zerlegung von X bezüglich des Tripotents p ,

$$X = X_0(p) \dot{+} X_{1/2}(p) \dot{+} X_1(p), \quad (4.19)$$

in die drei Eigenräume

$$X_\lambda(p) := \text{Eig}(D(p), \lambda), \quad \lambda \in \{0, \frac{1}{2}, 1\};$$

für einen Beweis siehe UPMEIER [279](1985), Seite 357, Satz 21.9, und — ähnlich wie bei SATAKE [254](1980), Seite 242 — STACHÓ [265](2007), Seite 3, Fußnote 2. Die kanonischen Projektionen $X \rightarrow X_\lambda(p)$ werden mit $P_\lambda(p)$ bezeichnet, $\lambda = 0, \frac{1}{2}, 1$; im Fall, dass X ein *JB -Tripel oder ein JB^* -Tripel über $\mathbb{C}$ ist, sind diese Projektionen

allesamt kontraktiv (PERALTA und STACHÓ [235](2001), Seite 83) und insbesondere ist somit dann die angegebene Zerlegung eine innere topologische Summe. Man beachte, dass, korrespondierend zu den Eigenwerten $\lambda = 0, \frac{1}{2}, 1$, manche Autoren anstatt des hier verwendeten Indexes λ den Index 2λ verwenden. Ist

$$X_0(p) = \{0\},$$

so heißt p *vollständig* (im Englischen *complete*). Versehen mit der p -Multiplikation als Multiplikation ist $X_1(p)$ eine Jordan-Algebra über $\mathbb{K}$ (es sei an 4.6.10 erinnert) mit Eins p ; im Fall, dass X ein Banach-Jordan-Tripel über $\mathbb{C}$ und $p \neq 0$ ist, ist diese Jordan-Algebra unital (UPMEIER, ebd., Seite 360, Lemma 21.11 und LOOS [200](1971), Seite 17); insbesondere ist also $X_1(p)$ eine JB^* -Algebra, falls X ein JB^* -Tripel über $\mathbb{C}$ und $p \neq 0$ ist.

(FRIEDMAN und RUSSO [100](1985), Seite 69; MARTÍNEZ und PERALTA [206] (2000); MEYBERG [212](1972), Seiten 163 und 164). Die Projektionen $P_\lambda(p)$, $\lambda = 0, \frac{1}{2}, 1$, erfüllen die folgenden Gleichungen:

$$\begin{aligned} P_0(p) &= B(p), \text{ also } P_0(p) = \text{Id}_X - 2D(p) + Q(p)^2, \\ P_{1/2}(p) &= 2(D(p) - Q(p)^2), \\ P_1(p) &= Q(p)^2. \end{aligned}$$

Nach zum Beispiel Gleichung (J.3.2) in FARAUT, KANEYUKI, KORÁNYI, LU und ROOS [92](2000), Seite 431, gilt für alle $x \in X$ die Gleichung $Q(x)^2 = 2D(x)^2 - D(x)$. Insbesondere lassen sich somit die Projektionen $P_\lambda(p)$, $\lambda = 0, \frac{1}{2}, 1$, auch wie folgt als ganzzahlige Polynome über $D(p)$ ausdrücken:

$$\begin{aligned} P_0(p) &= \text{Id}_X - 3D(p) + 2D(p)^2, \\ P_{1/2}(p) &= 4(D(p) - D(p)^2), \\ P_1(p) &= 2D(p)^2 - D(p). \end{aligned}$$

Ausgehend von diesen Gleichungen kann man wieder $P_\lambda(p)P_\mu(p) = \delta_{\lambda\mu}P_\mu(p)$ (Kronecker-Delta) für alle $\lambda, \mu \in \{0, \frac{1}{2}, 1\}$ zeigen, siehe FRIEDMAN und RUSSO, ebd., und MEYBERG, ebd.; des Weiteren erhält man mit diesen Projektionen genau wieder die anfangs angegebene Zerlegung (4.19) von X bezüglich des Tripotents p . Es gilt die Mengen-Gleichung

$$X_1(p) = Q(p)X. \quad (4.20)$$

Beweis. Sei $y \in Q(p)X$, also $y = Q(p)x$ für ein $x \in X$. Folglich ist $D(p)y = D(p)Q(p)x$ und dies ist nach der Homotopieformel (4.10) aus 4.6.7 gleich $Q(Q(p)p, p)x$, also $D(p)y = Q(p)x = y$ und $y \in X_1(p)$. Umgekehrt hat man $X_1(p) = Q(p)^2X \subseteq Q(p)X$. $\square$

Man bemerke $D(p) = \frac{1}{2}P_{1/2}(p) + P_1(p)$. Da p ein idempotentes Element in der Jordan-Algebra $X^{(p)}$ über $\mathbb{K}$ ist (siehe 4.6.10), ist die vorliegende Zerlegung offensichtlich gleich der Peirce-Zerlegung von $X^{(p)}$ bezüglich des Idempotents $p \in X^{(p)}$, siehe 2.3.17. Dementsprechend wird die vorliegende Zerlegung (4.19) die *Peirce-Zerlegung von X bezüglich des Tripotentes p* genannt, die $P_\lambda(p)$ heißen die *Peirce- λ -Projektionen* und die $X_\lambda(p)$ die *Peirce- λ -Räume* von X bezüglich des Tripotentes p , $\lambda = 0, \frac{1}{2}, 1$. Man bemerke, dass ein $x \in X$ genau dann unitär ist, wenn es tripotent ist und $X = X_1(x)$ gilt.

Nebenbei sei angemerkt, dass die Peirce-Zerlegung ein Beispiel ist einer allgemeineren von NEHER eingeführten *Peirce-Graduierung*; siehe EDWARDS und RÜTTIMANN [87](2003). Es gilt folgende „Peirce-Arithmetik“, siehe UPMEIER, ebd.:

$$\{X_0(p) \ X_1(p) \ X\} = \{X_1(p) \ X_0(p) \ X\} = \{0\} \quad (4.21)$$

(auf 4.7.61 vorgehend, in Worten: $X_0(p)$ und $X_1(p)$ sind zueinander orthogonal) und

$$\{X_k(p) X_\ell(p) X_m(p)\} \subseteq X_{k-\ell+m}(p),$$

wobei $X_\lambda(p) := 0$ für alle $\lambda \notin \{0, \frac{1}{2}, 1\}$ gesetzt ist. Insbesondere sind somit alle drei Eigenräume $X_\lambda(p)$, $\lambda = 0, \frac{1}{2}, 1$, Untertripel von X ; speziell sind $X_0(p)$ und $X_1(p)$ innere Ideale von X . Ist also X ein *JB -Tripel, so ist auch jeder dieser drei Eigenräume ebenfalls ein *JB -Tripel, und des Weiteren, falls X ein JB^* -Tripel über $\mathbb{C}$ ist, ist gemäß 4.7.11 jeder dieser drei Eigenräume ebenfalls ein JB^* -Tripel über $\mathbb{C}$. Falls X ein JBW^* -Tripel über $\mathbb{C}$ ist, so sind alle drei Peirce- λ -Projektionen wegen 4.7.28 schwach*-stetig, mithin sind dann damit alle drei Peirce- λ -Räume schwach*-abgeschlossen und mit 4.7.31 ebenfalls wieder JBW^* -Tripel über $\mathbb{C}$, $\lambda = 0, \frac{1}{2}, 1$; wegen Satz 4.7.49(c) und dessen Korollar 4.7.51 gilt das entsprechende auch für den Fall, dass X ein *JBW -Tripel ist (MARTÍNEZ und PERALTA [206](2000)).

Sei weiterhin wie gehabt X ein allgemeines Jordan-Tripel über $\mathbb{K}$. Sei $p \in \text{Tri}(X)$. Da p offensichtlich in dem Untertripel $P_1(p)$ ein unitäres Element ist, ist $X_1(p)$, versehen mit der p -Multiplikation als Multiplikation, gemäß 4.6.10 eine Jordan-Algebra über $\mathbb{K}$ mit Eins p ; des Weiteren ist dann im Fall, dass X ein Jordan-Tripel über $\mathbb{K}$ ist, $Q(p)$ eine Involution auf dieser nichtassoziativen Algebra über $\mathbb{K}$. Ist insbesondere X ein Banach-Jordan-Tripel über $\mathbb{K}$, so ist $X_1(p)$ eine involutive Banach-Jordan-Algebra über $\mathbb{K}$ (UPMEIER, ebd., Seite 360, Lemma 21.11). Ist also X ein JB^* -Tripel über $\mathbb{C}$ und das Tripotent p von Null verschieden, so ist $X_1(p)$ eine JB^* -Algebra.

Sei $p \in \text{Tri}(X)$. Gemäß der Fundamentalformel (4.12) von 4.6.7 gilt $Q(p)^3 = Q(p)$. Somit folgt aus der Eigenwertgleichung $Q(p)x = \lambda x$, dass ein Eigenwert $\lambda \in \mathbb{K}$ von $Q(p)$ stets die Gleichung $\lambda^3 = \lambda$ erfüllen muss. Mithin besitzt der Operator $Q(p)$ als mögliche Eigenwerte $-1, 0$ und 1 . Als Untervektorräume von $X_{\mathbb{R}}$ setzt man

$$X^\lambda(p) := \text{Eig}(Q(p), \lambda), \quad \lambda \in \{-1, 0, 1\}.$$

Die kanonischen Projektionen $X \rightarrow X^\lambda(p)$ werden mit $P^\lambda(p)$ bezeichnet, $\lambda = -1, 0, 1$. Ist X ein Jordan-Tripel über $\mathbb{K}$, so bemerke man, dass dann $X^1(p)$ der selbstadjungierte Teil der — gemäß des vorigen Absatzes mit der Involution $Q(p)$ ausgestatteten — Jordan-Algebra $X_1(p)$ über $\mathbb{K}$ ist, in Zeichen also

$$(X_1(p))_{(Q(p))} = X^1(p);$$

insbesondere ist dann also nach 3.1.10 die Projektion $P^1(p)$ gleich $\frac{1}{2}(\text{Id}_X + Q(p))P_1(p)$. Da nun also auch $Q(p)(1 - P_1(p)) = 0$ gilt, hat man für X wieder als ein allgemeines Jordan-Tripel über $\mathbb{K}$ die Zerlegungen

$$\begin{aligned} X_1(p) &= X^{-1}(p) \dot{+} X^1(p), \\ X^0(p) &= X_0(p) \dot{+} X_{\frac{1}{2}}(p), \\ X &= X^{-1}(p) \dot{+} X^0(p) \dot{+} X^1(p). \end{aligned}$$

Ist X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel, so gilt die folgende Rechenregel:

$$\{X^k(p) X^\ell(p) X^m(p)\} \subseteq X^{k\ell m}(p), \quad k, \ell, m \in \{-1, 1\};$$

insbesondere sind $X^{-1}(p)$ und $X^1(p)$ Untertripel von $X_{\mathbb{R}}$. Falls X ein Jordan-Tripel über $\mathbb{C}$ ist, gilt des Weiteren

$$X^{-1}(p) = iX^1(p),$$

denn: $x \in X^{-1}(p) \Leftrightarrow \{pxp\} = -x \Leftrightarrow \{p(-ix)p\} = -ix \Leftrightarrow -ix \in X^1(p) \Leftrightarrow x \in iX^1(p)$; falls X ein Jordan-Tripel über $\mathbb{R}$ ist, können $X^{-1}(p)$ und $X^1(p)$ verschiedene Dimensionen aufweisen (KAUP [177](1997), Seite 196).

4.7.60 M -Summand (FRIEDMAN und RUSSO [100](1985), Seite 72, Lemma 1.3). Sei X ein JB^* -Tripel über $\mathbb{K}$. Dann sind $X_0(p)$ und $X_1(p)$ komplementäre M -Summanden von $X_0(p) \dot{+} X_1(p)$. Mit anderen Zeichen gilt also:

$$(P_0(p) + P_1(p))X \cong X_0(p) \oplus_\infty X_1(p).$$

Das heißt, für alle $x \in X$ gilt die Gleichung $\|P_0(p)x + P_1(p)x\| = \max\{\|P_0(p)x\|, \|P_1(p)x\|\}$ und somit auch für alle $x^* \in X^*$ die Gleichung $\|x^* \circ P_0(p) + x^* \circ P_1(p)\| = \|x^* \circ P_0(p)\| + \|x^* \circ P_1(p)\|$.

Beweis. Da $P_0(p)$ und $P_1(p)$ kontraktiv sind, sind auch die Abbildungen $P_0(p) \upharpoonright X_0(p) \dot{+} X_1(p)$ und $P_1(p) \upharpoonright X_0(p) \dot{+} X_1(p)$ kontraktiv, das heißt, $\|P_0(p)x\| = \|P_0(p)(P_0(p) + P_1(p))x\| \leq \|P_0(p)x + P_1(p)x\|$ und analog auch $\|P_1(p)x\| \leq \|P_0(p)x + P_1(p)x\|$ für alle $x \in X$.

Sei $x \in X$. Setze $y_0 := P_0(p)x$ und $z_0 := P_1(p)x$. Setze $m := \max\{\|y_0\|, \|z_0\|\}$. Falls $m = 0$, ist nichts zu zeigen. Liege also der Fall $m \neq 0$ vor. Setze $y := y_0/m$ und $z := z_0/m$. Dann ist $\max\{\|y\|, \|z\|\} = 1$. Mit Satz 4.7.10 und Gleichung (4.21) in 4.7.59 hat man für alle $n \in \mathbb{N}$: $\|P_0(p)x + P_1(p)x\| = \|y_0 + z_0\| = m\|y + z\| = m\|\{(y+z)(y+z)(y+z)\}\|^{1/3} \equiv m\|(y+z)^3\|^{1/3} = \dots = m\|(y+z)^{(3^n)}\|^{(3^{-n})} = m\|y^{(3^n)} + z^{(3^n)}\|^{(3^{-n})} \leq m \cdot (\|y^{(3^n)}\| + \|z^{(3^n)}\|)^{(3^{-n})} = m \cdot (\|y\|^{(3^n)} + \|z\|^{(3^n)})^{(3^{-n})} \leq m \cdot (2)^{(3^{-n})}$ und dies konvergiert für $n \rightarrow \infty$ nach $m = \max\{\|P_0(p)x\|, \|P_1(p)x\|\}$. $\square$

4.7.61 Orthogonalität (LOOS [202](1977), Seite 3.8, Lemma 3.9). Sei X ein allgemeines Jordan-Tripel über $\mathbb{K}$. Gemäß 4.7.59 bezeichne $X_0(p)$ für eine beliebiges tripotentes Element p von X den Kern der Abbildung $D(p)$. Seien p und q zwei tripotente Elemente von X . Dann sind die folgenden Aussagen äquivalent:

- | | |
|----------------------|----------------------|
| (a) $p \in X_0(q)$, | (e) $q \in X_0(p)$, |
| (b) $D(q)p = 0$, | (f) $D(p)q = 0$, |
| (c) $\{qqp\} = 0$, | (g) $\{ppq\} = 0$, |
| (d) $D(q, p) = 0$, | (h) $D(p, q) = 0$. |

Beweis. Die Äquivalenzen (a) $\Leftrightarrow$ (b) $\Leftrightarrow$ (c) und (e) $\Leftrightarrow$ (f) $\Leftrightarrow$ (g) sind klar.

(c) $\Rightarrow$ (d): Gelte (c). Zuerst bemerke man $Q(q)p = \{qpq\} = \{qp\{qqq\}\} = \{q\{pqq\}q\}$ (Homotopieformel (4.10)) $= \{q\{qqp\}q\} = 0$. Damit gilt: $[D(q), D(q, p)] = D(\{qqq\}, p) - D(q, \{qqp\})$ (Bemerkung 4.6.8) $= D(q, p) = -[D(q, p), D(q)] = -D(\{qpq\}, q) + D(q, \{pqq\})$ (Bemerkung 4.6.8) $= D(q, \{qqp\})$ (ebige Vorbemerkung) $= 0$.

(d) $\Rightarrow$ (g): Gelte (d). Dann ist $\{ppq\} = \{qpp\} = D(q, p)p = 0$.

(g) $\Rightarrow$ (h): Vertauschen von p und q und dann (c) $\Rightarrow$ (d) anwenden.

(h) $\Rightarrow$ (c): $\{qqp\} = \{pqq\} = D(p, q)q$. $\square$

Insbesondere ist das System $(\text{Tri}(X), \text{Tri}(X); D \upharpoonright \text{Tri}(X)^2; L(X))$ orthosymmetrisch. Es sei an 4.6.20 erinnert. Sind p und q orthogonal, so kommutieren die beiden Operatoren $D(p)$ und $D(q)$, $p + q$ ist tripotent und es gilt $D(p + q) = D(p) + D(q)$.

Beweis. Seien p und q orthogonal zueinander. Dann gilt $[D(p), D(q)] = D(\{ppq\}, q) - D(q, \{ppq\})$ (Bemerkung 4.6.8) $= 0$ und $\{(p + q)(p + q)(p + q)\} = \{ppp\} + \{ppq\} + \{pqp\} + \{pqq\} + \{qpp\} + \{qpq\} + \{qqp\} + \{qqq\} = p + q$. $D(p + q) = D(p, p) + D(p, q) + D(q, p) + D(q, q) = D(p) + D(q)$. $\square$

Ist X ein Jordan-Tripel über $\mathbb{C}$ oder ein *JB -Tripel, so gilt: p und q sind genau dann orthogonal, wenn $p + q$ und $p - q$ tripotent sind (EDWARDS, MCCRIMMON und RÜTTIMANN [84](1996), Seite 200; ISIDRO, KAUP und RODRÍGUEZ PALACIOS [147](1995), Lemma 3.6).

4.7.62 (FERNÁNDEZ POLO, MARTÍNEZ MORENO, PERALTA [94](2004), Seite 441, Korollar 2.5). Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Sei $p \in X$ mit $\|p\| = 1$. Dann sind äquivalent:

- (a) p ist vollständig, sprich $p \in \text{Tri}(X)$ mit $B(p) = 0$. (Definition in 4.7.59.)
- (b) p ist ein Extrempunkt von B_X .
- (c) p ist ein reeller Extrempunkt von B_X .
- (d) Nur für $x = 0$ existiert $\lambda > 0$ mit $\|p + \lambda x\| = \|p - \lambda x\| = 1$.

Falls X ein JB^* -Tripel über $\mathbb{C}$ ist, ist die folgende Aussage äquivalent zu jeder der vier vorstehenden Aussagen (KAUP und UPMEIER [183](1977), Seite 190. Satz 3.5):

- (e) p ist ein komplexer Extrempunkt von B_X .

Ist X ein JB^* -Tripel über $\mathbb{C}$, so ist ein $p \in \text{Tri}(X)^\times$ genau dann vollständig, wenn es regulär ist (KAUP und UPMEIER [183](1977), Seite 189, Lemma 3.2). Für eine weitere äquivalente Formulierung für den Fall, dass X ein JBW^* -Tripel über $\mathbb{C}$ ist, siehe (4.24) in 4.7.66.

4.7.63 Hilbertraum (CHU und ISIDRO [47](2000), Beispiel 2.8). Sei H ein Hilbertraum über $\mathbb{C}$. Dann ist H , versehen mit dem Tripelprodukt

$$\{xyz\} := \frac{1}{2}(\langle x, y \rangle z + \langle z, y \rangle x),$$

ein JBW^* -Tripel über $\mathbb{C}$. Innerhalb eines jeden solchen JBW^* -Tripels über $\mathbb{C}$ gilt: Die Menge aller von Null verschiedenen Tripotenten ist genau gleich der Menge der Extrempunkte der abgeschlossenen Einheitskugel des gegebenen Hilbertraumes H . Ist $p \in \text{Tri}(H)$, so lauten die Peirce-Räume bezüglich p : $H_0(p) = \{0\}$, $H_{\frac{1}{2}}(p) = \{p\}^\perp$ und $H_1(p) = \mathbb{C}p$.

Speziell liefert der Hilbertraum $H = \mathbb{C}^2$ das folgende Gegenbeispiel für ein JB^* -Tripel über $\mathbb{R}$, für das sich die Aussage von 4.7.27 nicht überträgt: STACHÓ [267](2001) gibt für das *JB -Tripel $X = H_{\mathbb{R}}$ eine kontraktive Projektion p an, so dass $p(X)$, versehen mit dem projizierten Tripelprodukt (siehe 4.7.27), kein *JB -Tripel ist. Nichtsdestotrotz, und dabei von der in Satz 4.7.7(d) beschriebenen Situation für JB^* -Tripel über $\mathbb{C}$ abweichend, existieren auf $p(X)$ verschiedene Jordan-Tripelprodukte, so dass $p(X)$, jeweils weiterhin mit der kanonischen Hilbertraumnorm von H versehen (es sei an 3.4.10 erinnert), ein *JB -Tripel ist (ebd., Seite 227).

4.7.64 C^* -Algebra (BURGOS, FERNÁNDEZ POLO, GARCÉS, MARTÍNEZ MORENO und PERALTA [41](2008), Seite 229).

Sei A eine C^* -Algebra über $\mathbb{C}$. Bezeichne X das JB^* -Tripel A^Δ über $\mathbb{C}$. Existiere ein tripotentes Element $p \in \text{Tri}(X)$. Dann gilt $X_0(p) = (1 - pp^*)A(1 - p^*p)$, $X_{\frac{1}{2}}(p) = pp^*A(1 - p^*p) \dot{+} (1 - pp^*)Ap^*p$, $X_1(p) = pp^*Ap^*p$, wobei p jeweils auf der rechten Seite jeder Gleichung gemäß 4.7.3 als eine partielle Isometrie von A aufgefasst wird. Es sei an Gleichung (3.14) aus 3.5.33 erinnert. Des Weiteren vergleiche man auch mit Gleichung (2.19) von 2.1.33.

4.7.65 Stetige Funktionen (CHU und ISIDRO [47](2000), Beispiel 2.9). Sei $X := C(\Omega)$ die C^* -Algebra über $\mathbb{C}$ der komplexen, stetigen Funktionen auf einem kompakten Hausdorffraum Ω . Versehen mit dem Tripelprodukt $\{fgh\} := f\bar{g}h$ ist X ein JB^* -Tripel über $\mathbb{C}$. Es gilt $\text{Tri}(X) = \{f \in C(\Omega) : f \circ (1 - |f|^2) = 0\}$ und die Peirce-Räume bezüglich eines $g \in \text{Tri}(X)^\times$ lauten: $X_{\frac{1}{2}}(g) = \{0\}$, $X_1(g) = \{f \in C(\Omega) : g^{-1}(0) \subseteq f^{-1}(0)\}$.

4.7.66 Tripelordnung (BATTAGLIA [15](1991)). Für das Folgende sei als eine Analogie an die Äquivalenz (2.21) aus 2.1.45 erinnert. Sei X ein Jordan-Tripel über $\mathbb{K}$. Setze für alle $p, q \in \text{Tri}(X)$

$$p \leq q \Leftrightarrow (q - p \in \text{Tri}(X) \text{ und } p \perp q - p);$$

mit anderen Worten:

$$p \leq q \Leftrightarrow \exists r \in p^\perp \cap \text{Tri}(X) : q = p + r.$$

Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Seien $p, q \in \text{Tri}(X)$. Dann sind die folgenden Aussagen äquivalent:

- | | |
|----------------------|-------------------------|
| (i) $p \leq q$, | (iv) $D(p)q = p$, |
| (ii) $P_1(p)q = p$, | (v) $D(p, q) = D(p)$, |
| (iii) $Q(p)q = p$, | (vi) $D(q, p) = D(p)$. |

Beweis. Für die folgenden beiden Implikationen wird 4.7.61 verwendet:

$$(i) \Rightarrow (v): p \leq q \Rightarrow p \perp q - p \Rightarrow D(p, q - p) = 0;$$

$$(i) \Rightarrow (vi): p \leq q \Rightarrow D(q - p, p) = 0.$$

$$(v) \Rightarrow (iii): D(p, q) = D(p) \Rightarrow \{pqp\} = \{ppp\} = p; \text{ analog:}$$

$$(vi) \Rightarrow (iv): D(q, p) = D(p) \Rightarrow \{ppq\} = \{qpp\} = \{ppp\} = p.$$

$$(iii) \Rightarrow (ii): \text{Klar wegen } P_1(p) = Q(p)^2.$$

$$(iv) \Rightarrow (ii): \text{Schlichtes Anwenden der Jordan-Gleichung: } P_1(p)q = Q(p)^2q = \{p\{pqp\}p\} = -\{qp\{ppp\}\} + \{pp\{qpp\}\} + \{\{qpp\}pp\} = -p + p + p = p.$$

Während die bisher gezeigten Implikationen auch gelten, wenn X nur als ein Jordan-Tripel über $\mathbb{K}$ angenommen wird, wird die Äquivalenz von (i) und (ii) von FRIEDMAN und RUSSO [100](1985), Korollar 1.7, unter Verwendung der Struktur der JB^* -Tripel über $\mathbb{C}$ bewiesen. Man überlegt sich leicht, dass die Äquivalenz von (i) und (ii) dann auch für X als ein *JB -Tripel gilt. $\square$

Sei X ein Jordan-Tripel über $\mathbb{K}$ und seien $p, q \in \text{Tri}(X)$ mit $p \leq q$. Dann gilt die Gleichungskette

$$P_1(q)p = Q(q)p = D(q)p = p, \quad (4.22)$$

mit anderen Zeichen also

$$\{q\{qpq\}q\} = \{qpq\} = \{qqp\} = p.$$

Des Weiteren gilt

$$X^1(p) \subseteq X^1(q). \quad (4.23)$$

Beweis. $p \leq q \Rightarrow D(q, p) = D(p)$ wegen (vi). $\Rightarrow \{qpq\} = D(p)q = p$ wegen (iv). $\Rightarrow Q(q)^2p = \{q\{qpq\}q\} = \{qpq\} = p \Leftrightarrow p \in X_1(q) \Leftrightarrow D(q)p = p$.

Sei $x \in X^1(p)$, also $x = Q(p)x$. Unter Berücksichtigung der Fundamentalformel und der eben bewiesenen Gleichungskette gilt somit: $Q(q)x = Q(q)Q(p)x = Q(q)Q(Q(p)q)x = Q(q)Q(p)Q(q)Q(p)x = Q(Q(q)p)x = Q(p)x = x$, also $x \in X^1(q)$. $\square$

Hat man anfangs nur $Q(q)p = p$ vorliegen, so folgt zumindest $p = \{pqq\} = \{qpq\} = \{qqp\}$ und daraus wiederum $\{ppq\} = \{pqp\} = \{qpp\}$.

Beweis. $p = \{qpq\} = \{qp\{qqq\}\} = \{\{qpq\}qq\} = \{pqq\}$ und $\{ppq\} = \{p\{qpq\}q\} = \{pq\{pqq\}\} = \{pqp\}$. $\square$

Gilt für zwei Tripotente $p, q \in X$ sowohl die Ungleichung $X^1(p) \subseteq X^1(q)$ als auch die Gleichung $Q(q)p = Q(p)q$, so folgt $p \leq q$.

Beweis. Nach Voraussetzung gilt die Implikation $Q(p)x = x \Rightarrow Q(q)x = x$. Da $p = Q(p)p$, also auch $Q(q)p = p$. $\square$

Mit den nun vorhandenen Aussagen lässt sich jetzt zeigen, dass $\leq$ auf $\text{Tri}(X)$ eine partielle Ordnung mit 0 als kleinstem Element ist.

Beweis. $0 \leq x$ für alle $x \in \text{Tri}(X)$ ist klar. Ebenso die Reflexivität. Zur Antisymmetrie: Sei $x \leq y$ und $y \leq x$. Dann gilt $x - y = \{x - y, x - y, x - y\} = \{x - y, x - y, x\} - \{x - y, x - y, y\} = 0 - 0 = 0$. Es bleibt die Transitivität zu zeigen: Sei $x \leq y$ und $y \leq z$. Dann gilt die Gleichungskette (4.22), insbesondere also $Q(y)x = x$. Somit gilt $\{xzx\} = \{Q(y)x, z, Q(y)x\} = Q(Q(y)x)z$; dies ist nach der in 4.6.7 aufgeführten Fundamentalformel (4.12) gleich $Q(y)Q(x)Q(y)z$; dies wiederum ist nach (iii) gleich $Q(y)Q(x)y = Q(y)x = x$. Also gilt $Q(x)z = x$, das heißt, $x \leq z$. $\square$

Jeder lineare Tripel-Homomorphismus zwischen zwei Jordan-Tripeln über $\mathbb{K}$ erhält die Ordnung $\leq$.

Beweis. Seien $p, q \in \text{Tri}(X)$ mit $p \leq q$. Dann gilt $D(Tp)(Tq - Tp) = \{Tp, Tp, T(q - p)\} = T(D(p)(q - p)) = 0$. $\square$

Sei X ein JBW^* -Tripel über $\mathbb{C}$. Dann besitzt $\text{Tri}(X)$ genau dann ein größtes Element, wenn $X = \{0\}$ gilt. Für nicht triviales X gibt es also stets eine nicht leere Teilmenge von $\text{Tri}(X)$, die kein Supremum besitzt. Im Gegensatz dazu, existiert aber für jede nicht leere Teilmenge von $\text{Tri}(X)$ stets ihr Infimum. Für eine beliebige Teilmenge M von $\text{Tri}(X)$ existiert genau dann das Supremum von M , wenn M eine obere Schranke hat. Ist $(x_\nu)_{\nu \in I}$ ein aufsteigendes Netz in $\text{Tri}(X)$, so existiert das Supremum $\sup(x_\nu)_{\nu \in I}$ und ist gleich dem schwach*-Grenzwert von $(x_\nu)_{\nu \in I}$. Für jedes $p \in \text{Tri}(X)$ sind die beiden folgenden Aussagen äquivalent:

- (a) p ist vollständig (in 4.7.59 definiert)
 - (b) p ist ein maximales Element in $\text{Tri}(X)$.
- (4.24)

In diesem Zusammenhang sei an 4.7.62 erinnert.

(PERALTA und STACHÓ [235](2001), Seite 81). Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Dann gelten für alle $p, q \in \text{Tri}(X)$ mit $p \leq q$ die beiden Inklusionen $X_1(p) \subseteq X_1(q)$ und $X_0(q) \subseteq X_0(p)$.

Sei X ein Jordan-Tripel über $\mathbb{K}$. In Analogie zu den in 2.1.47 definierten minimal idempotenten Elementen eines Ringes heißt ein $p \in \text{Tri}(X)$ *minimal tripotent*, wenn $p \neq 0$ und p ein minimales Element in $(\text{Tri}(X)^\times, \leq)$ ist. Offensichtlich ist ein $p \in \text{Tri}(X)^\times$ genau dann minimal tripotent, wenn es keine zwei zueinander orthogonale $q, r \in \text{Tri}(X)^\times$ mit $p = q + r$ gibt. In Analogie zu der bei den minimal Idempotenten vorliegenden gleichen Situation, kann man daher minimal Tripotente auch als *irreduzibel tripotent* bezeichnen.

Die Abbildung $p \mapsto X^1(p)$ von $\text{Tri}(X)^\times$ in die Menge aller Untervektorräume von $X_{\mathbb{R}}$ ist eine streng ordnungserhaltende Abbildung.

Beweis. Dass die Abbildung ordnungserhaltend ist, ist wegen der Inklusion (4.23) klar. Sei nun $p < q$, also $p \leq q$, $p \neq q$. Angenommen, $X^1(p) \supseteq X^1(q)$. Nach Definition von $X^1(q)$ ist $q \in X^1(q)$. Nach Annahme also auch $q \in X^1(p)$, das heißt, $q = Q(p)q = \{pqp\}$, aber wegen $p \leq q$ ist $\{pqp\} = p$, somit $p = q$, *Widerspruch*. $\square$

(BECERRA GUERRERO, LÓPEZ PÉREZ, PERALTA und RODRÍGUEZ PALACIOS [16](2004)). Sei X ein allgemeines Jordan-Tripel über $\mathbb{K}$. Liege der Fall vor, dass p ein tripotentes Element von X ist. Das Tripotent p heißt *divisionsminimal tripotent* (im Englischen *division tripotent*), wenn es von Null verschieden und die Jordan-Algebra $X_1(p)$ über $\mathbb{K}$ eine Divisions-Jordan-Algebra über $\mathbb{K}$ ist (siehe Definition 2.3.15). (Man bemerke die wegen Gleichung (4.20) $X_1(p) = \{pXp\}$ aus 4.7.59 bestehende Analogie zu der Definition der schiefminimalen Idempotente in 2.1.47.)

Liege zusätzlich der Fall vor, dass das Tripotent p von Null verschieden sei. Ist X ein *JB -Tripel (also insbesondere zum Beispiel die reelle Strukturierung eines JB^* -Tripels über $\mathbb{C}$), so ist p genau dann divisionsminimal tripotent, wenn der selbstadjungierte Teil der Jordan-Algebra $X_1(p)$ über $\mathbb{R}$, also $X^1(p)$, gleich $\mathbb{R}p$ ist. Da für JB^* -Tripel X über $\mathbb{C}$ die Gleichung $X_1(p) = X^1(p) + iX^1(p)$ gilt, ist in einem JB^* -Tripel X über $\mathbb{C}$ ein $p \in \text{Tri}(X)^\times$ genau dann divisionsminimal tripotent, wenn die Jordan-Algebra $X_1(p)$ über $\mathbb{C}$ gleich $\mathbb{C}p$ ist; nach Gleichung (4.20) aus 4.7.59 ist dies für solche X wiederum äquivalent zu der Gleichung $Q(p)X = \mathbb{C}p$.

Wegen der vorhin erwähnten strengen Monotonie der Abbildung $p \mapsto X^1(p)$ ist offensichtlich jedes divisionsminimal tripotente Element minimal tripotent; im Fall, dass X ein Prädual besitzt, gilt auch die Umkehrung (PERALTA und STACHÓ [235](2001), Seite 81, Satz 2.2, und Seite 85), also:

Ist X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel,
so ist ein $p \in X$ genau dann minimal tripotent, wenn (4.25)
es divisionsminimal tripotent ist.

Als ein Beispiel, wo diese Umkehrung nicht gilt, hat man den Funktionenraum und das *JB -Tripel $C([0, 1], \mathbb{R})$ (ebd., Seite 81). Es sei an die in 2.4.18 dargestellte ähnliche Situation bei Ringen erinnert. Sei X ein *JBW -Tripel und p ein minimal tripotentes Element von X . Bezüglich der Eigenschaften, die dann p als ein Element von $X_{(\mathbb{C})}$ aufweist, siehe PERALTA und STACHÓ [235](2001), Seite 85. Das ebd. zu findende Lemma 3.2 wurde von BECERRA GUERRERO, LÓPEZ PÉREZ, PERALTA und RODRÍGUEZ PALACIOS [16](2004), Seite 55, Korollar 3.5, auf den Fall verallgemeinert, dass X ein *JB -Tripel und p divisionsminimal tripotent in X ist.

4.7.67. Sei X ein allgemeines Jordan-Tripel über $\mathbb{K}$. Sei U ein inneres Ideal von X . Dann ist ein $p \in \text{Tri}(U)$ genau dann ein divisionsminimal tripotentes Element von U , wenn es ein divisionsminimal tripotentes Element von X ist.

Beweis. Sei U erst einmal nur ein Untertripel von X und $p \in \text{Tri}(U)$. Dann gilt zunächst einmal (siehe 4.7.59) $X^1(p) \subseteq X_1(p) \stackrel{(4.20)}{=} Q(p)X$ und $X^1(p) \cap U = U^1(p)$. Liege nun zusätzlich der Fall vor, dass U ein inneres Ideal von X sei. Dann gilt $Q(p)X \subseteq U$, also $X^1(p) \subseteq U$. Es folgt $X^1(p) = U^1(p)$. $\square$

4.7.68. Sei X eine C^* -Algebra über $\mathbb{C}$. Dann ist eine idempotente partielle Isometrie $p \in X$ genau dann schiefminimal idempotent, wenn $p \in X^\Delta$ divisionsminimal tripotent ist.

Beweis. In X^Δ gilt für jedes $x \in X$: $\{p xp\} = \frac{1}{2}(px^*p + px^*p) = px^*p$. Wie in Bemerkung 3.1.4 erwähnt, gilt mit der dortigen Schreibweise hier $X = X^*$; daher ist $\{pXp\} = pX^*p = pXp$. $\square$

Nach 2.4.18 hat somit das JB^* -Tripel $C([0, 1], \mathbb{C})$ über $\mathbb{C}$ kein divisionsminimal tripotentes Element.

4.7.69 Atome in C^* -Algebren. Sei A eine C^* -Algebra über $\mathbb{K}$ und $p \in A^\times$ eine Projektion. Es sei an die Definition 3.5.57 eines Atoms erinnert. Es gilt: p ist genau dann ein Atom von A , wenn $p \in A^\Delta$ divisionsminimal tripotent ist.

Beweis. Da in A^Δ für jedes $x \in A$ die Gleichung $px^*p = \{pxp\}$ gilt, folgt mit Lemma 3.2.5: $(pAp)_{(*)\mathbb{R}} = \{x \in A : px^*p = x\}_{\mathbb{R}} = \{x \in A : \{pxp\} = x\}_{\mathbb{R}} = \{x \in A : Q(p)x = x\}_{\mathbb{R}} = (A^\Delta)^1(p)$. $\square$

Ist ein $x \in A^\Delta$ divisionsminimal tripotent, so ist x^*x ein Atom von A .

Beweis. Falls x eine Projektion von A ist, ist das nach der eben gemachten Aussage klar. Falls aber x keine Projektion von A ist, so ist x nach 4.7.3 zumindest eine partielle Isometrie von A . Ganz gemäß BECERRA GUERRERO und MARTÍN [18](2005), Seite 328, kann man nun wie folgt argumentieren. Sei p die Anfangsprojektion von x , also $p := x^*x$. Sei $y \in (pAp)_{(*)\mathbb{R}}$, also $py^*p = y$. Dann gilt $Q(x)(xy) = \{x, xy, x\} = x(xy)^*x = (xx^*x)(y^*x^*)x = x(x^*xy^*x^*x) = xpy^*p = xy$, das heißt, $xy \in (A^\Delta)^1(x)$. Nach Voraussetzung ist $(A^\Delta)^1(x) = \mathbb{R}x$, das heißt, $xy = \lambda x$ für ein $\lambda \in \mathbb{R}$. Somit gilt $\lambda p = \lambda x^*x = x^*(\lambda x) = x^*(xy) = x^*((xx^*x)(y^*x^*)x) = ppy^*p = py^*p = y$, das heißt, $y \in \mathbb{R}p$ und p ist ein Atom von A . $\square$

Somit enthält eine C^* -Algebra über $\mathbb{K}$ genau dann mindestens ein Atom, wenn A^Δ mindestens ein divisionsminimal tripotentes Element enthält.

4.7.70 Minimale Tripotente in $i_X(X)$. Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Sei $p \in X$. Es sei an die Äquivalenzaussage (4.25) von 4.7.66 erinnert. Die folgenden beiden Aussagen sind äquivalent:

- (a) p ist ein divisionsminimal tripotentes Element von X .
- (b) $i_X(p)$ ist ein minimal tripotentes Element von X^{**} .

Beweis. Zunächst sei daran erinnert, dass $i_X(X)$ schwach*-dicht in X^{**} ist, siehe 2.5.11(b). Zu (a) $\Rightarrow$ (b): Gelte (a), das heißt, $X^1(p) = \mathbb{R}p$. Da nach 4.7.28 bzw. 4.7.49 das Tripelprodukt auf X^{**} getrennt schwach*-stetig ist, ist die Projektion $P^1(\hat{p})$ von X^{**} auf $(X^{**})^1(\hat{p})$ schwach*-stetig. Somit ist $P^1(\hat{p})\hat{X}$ schwach*-dicht in $P^1(\hat{p})X^{**}$. Es ist $P^1(\hat{p})\hat{X} = \frac{1}{2}(\text{Id}_{X^{**}} + Q(\hat{p}))Q(\hat{p})^2\hat{X} = i_X\left(\frac{1}{2}(\text{Id}_X + Q(p))Q(p)^2X\right) = i_X(\mathbb{R}p) = \mathbb{R}\hat{p}$. Da $\mathbb{R}\hat{p}$ in X^{**} schwach*-abgeschlossen ist, folgt $P^1(\hat{p})X^{**} = \mathbb{R}\hat{p}$, das heißt, $\hat{p}$ ist minimal tripotent in X^{**} .

Zu (b) $\Rightarrow$ (a): Gelte (b). In $(X^{**})_{\mathbb{R}}$ ist $P^1(\hat{p})\hat{X}$ ein Untervektorraum, der schwach*-dicht in dem Unterraum $P^1(\hat{p})X^{**}$ enthalten ist. Nach Voraussetzung ist $P^1(\hat{p})X^{**} = \mathbb{R}\hat{p}$, also $P^1(\hat{p})\hat{X} = \mathbb{R}\hat{p}$. Wegen $P^1(\hat{p})\hat{X} = i_X(P^1(p)X)$, also $P^1(p)X = \mathbb{R}p$ und p ist divisionsminimal tripotent in X . $\square$

4.7.71 Bemerkung. Der folgende Satz wird in EDWARDS, FERNÁNDEZ POLO, HOSKIN und PERALTA [83](2010) und in FERNÁNDEZ POLO und PERALTA [95](2010) bewiesen und war eine bis dahin über 20 Jahre lang offene Vermutung.

Sei X ein JB^* -Tripel über $\mathbb{C}$. Dann sind sowohl die abgeschlossenen Seiten von B_X als auch die schwach*-abgeschlossenen Seiten von B_{X^*} kugelabgeschlossen.

Mit 2.11.7 hat man also für beliebige JB^* -Tripel über $\mathbb{C}$ als Folgerungen:

- (a) Eine Seite von B_X ist genau dann abgeschlossen, wenn sie kugelabgeschlossen ist.
- (b) Eine Seite von B_{X^*} ist genau dann schwach*-abgeschlossen, wenn sie kugelabgeschlossen ist.

4.7.72 Rang (KAUP [177](1997), Seite 196). Der *Rang* (im Englischen *rank*) eines JB^* -Tripels über $\mathbb{C}$ oder eines *JB -Tripels X ist die kleinste Kardinalzahl, die größer oder gleich der Mächtigkeit jeder orthogonalen Teilmenge von X ist.

Jeder Hilbertraum über $\mathbb{C}$, gemäß 4.7.63 aufgefasst als ein JB^* -Tripel über $\mathbb{C}$, ist vom Rang eins (BURGOS, FERNÁNDEZ POLO, GARCÉS, MARTÍNEZ MORENO und PERALTA [41](2008), Seite 228, Bemerkung 15).

4.7.73 Tripelordnung und Kugelsystem (EDWARDS und RÜTTIMANN [86] (2001); EDWARDS und RÜTTIMANN [81](1988)). Sei X ein Jordan-Tripel über $\mathbb{K}$. Bezeichne $\text{Tri}(X)^\sim$ die disjunkte Vereinigung von $\text{Tri}(X)$ mit einer einelementigen Menge $\{\omega\}$. $\text{Tri}(X)^\sim$ wird überdies als eine punktierte Menge mit Basispunkt 0 aufgefasst. Die Relation $\leq$ auf $\text{Tri}(X)$ wird zu der Relation

$$\leq \cup (\text{Tri}(X)^\sim \times \{\omega\})$$

auf $\text{Tri}(X)^\sim$ erweitert; diese wird dann ebenfalls mit $\leq$ bezeichnet und ist eine partielle Ordnung auf $\text{Tri}(X)^\sim$.

Sei X JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Betrachte das Kugelsystem des Dualsystems (X_*, X) . Setze $\omega_{\perp} := B_{X_*}$. (Also $\omega_{\perp}^{-1} = \emptyset$.) Die Abbildung

$$\text{Tri}(X)^\sim \rightarrow \text{Faces}(\|\cdot\|; B_{X_*}), \quad p \mapsto p_{\perp} \quad (4.26)$$

ist ein Ordnungsisomorphismus. Da $\text{Faces}(\|\cdot\|; B_{X_*})$ ein vollständiger Verband ist, ist somit auch $\text{Tri}(X)^\sim$ ein vollständiger Verband. Ein *Atom* von X ist ein Extrempunkt der abgeschlossenen Einheitskugel B_{X_*} des Präduals X_* . Wegen 2.11.12 sind somit genau die minimalen Elemente von $\text{Faces}(\|\cdot\|; B_{X_*})^\times$ die Atome von X . Insbesondere bildet somit die Abbildung $p \mapsto p_{\perp}$ die Menge der minimalen Elemente von $\text{Tri}(X)^\sim{}^\times$ bijektiv auf die Menge der Atome von X ab. Siehe hierzu auch 4.7.74(b).

Die Abbildung

$$\text{Tri}(X)^\sim \rightarrow \text{Faces}(w^*; B_X), \quad p \mapsto p_{\perp}^{-1} \quad (4.27)$$

ist ein Antiordnungsisomorphismus mit

$$p_{\perp}^{-1} = p + B_{X_0(p)} \quad \text{für alle } p \in \text{Tri}(X).$$

Insbesondere ist jedes tripotente Element p von X nach 4.7.61(a) jeweils das kleinste Element von $p_{\perp}^{-1} \cap \text{Tri}(X)$. Mit 2.11.12 hat man als eine Folgerung des Weiteren insbesondere die in 4.7.62 und (4.24) von 4.7.66 bereits für JBW^* -Tripel über $\mathbb{C}$ formulierte Gleichheit der Menge der maximalen Elemente von $\text{Tri}(X)$ und der Menge der Extrempunkte von B_X auch für *JBW -Tripel.

Es gilt:

$$\begin{aligned} \text{Expfaces}(\|\cdot\|; B_{X_*}) &= \text{Semiexpfaces}(\|\cdot\|; B_{X_*}) = \text{Faces}(\|\cdot\|; B_{X_*}) \\ \text{und} &\quad \text{Semiexpfaces}(w^*; B_X) = \text{Faces}(w^*; B_X). \end{aligned}$$

Beweis. Siehe [86], Seite 171, Lemma 2.1, Theorem 3.7 und Lemma 3.8. Bleibt zu zeigen, dass für *JBW -Tripel die Abbildung $p \mapsto p_{\perp}^{-1}$ ein Antiordnungsisomorphismus ist. Dazu betrachte ebd., Lemma 3.8. Unter Mitverwendung der dortigen Terminologie ergibt sich die Behauptung dann wie folgt:

Sei A die Hermitefizierung von X und σ die nach ebd., Lemma 3.1, existierende Involution mit $X = X_{(\mathbb{C})(\sigma^*)}$, wobei σ^* so wie in ebd., Lemma 3.1, aufzufassen ist. Bezeichne $\top$ das Kugelannihilatorsymbol von (A_*, A) und $\perp$ (wie gehabt) das Kugelannihilatorsymbol von $((A^{\sigma^*})_*, A^{\sigma^*}) = (A_*^{\sigma}, A^{\sigma^*}) = (X_*, X)$. Für jedes $p \in \text{Tri}(A^{\sigma^*})$ ist die Mengengleichung $p_{\top}^{\top} \cap A^{\sigma^*} = p_{\perp}^{\perp}$ zu zeigen. Die dazu entscheidende Aussage findet sich in ebd., Lemma 3.6(i), wonach $p_{\top} = p_{\perp}$ gilt.

„ $\subseteq$ “: Sei $x \in p_{\top}^{\top} \cap A^{\sigma^*}$ und $\ell \in p_{\perp}$. Nach der Vorbemerkung also $\ell \in p_{\top}$ und somit $\ell(x) = 1$, das heißt, $x \in p_{\perp}^{\perp}$.

„ $\supseteq$ “: Sei $x \in (A^{\sigma^*})_1$ mit: Für alle $\ell \in p_{\perp} : \ell(x) = 1$. Wieder nach der Vorbemerkung gilt also: Für alle $\ell \in p_{\top} : \ell(x) = 1$. Also $x \in p_{\perp}^{\perp}$. $\square$

4.7.74 Trägerpunkte. (a) Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Betrachte das Kugelsystem des Dualsystems (X_*, X) . Sei $x \in S_X$. Nach 2.11.7 gilt $x_{\perp} \in \text{Faces}(\|\cdot\|; B_{X_*}) \setminus \{B_{X_*}\}$. Das damit gemäß (4.26) aus 4.7.73 eindeutig bestimmte Element p aus $\text{Tri}(X)$ mit $p_{\perp} = x_{\perp}$ wird der *Träger von x (in X)* (im Englischen *support*) genannt, in Zeichen $p = \text{spt}_X(x)$. In EDWARDS und RÜTTIMANN [81](1988), Seite 324, Lemma 3.4, findet man für JBW^* -Tripel über $\mathbb{C}$ drei Aussagen, die den Träger von x in X charakterisieren.

(b) Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Betrachte das Kugelsystem des Dualsystems (X_*, X) . Sei $\ell \in S_{X_*}$. Nach 2.11.7 ist $\ell^{\perp} \in \text{Faces}(w^*; B_X)$. Das dann gemäß (4.27) aus 4.7.73 eindeutig bestimmte Element p aus $\text{Tri}(X)^{\times}$ mit $\ell^{\perp} = p_{\perp}^{\perp}$ wird der *Träger von ℓ (in X)* genannt, in Zeichen $p = \text{spt}_X(\ell)$; nach (4.27) gilt $\ell^{\perp} = \text{spt}_X(\ell) + B_{X_0(\text{spt}_X(\ell))}$. (Nach dem Satz von Hahn-Banach ist hier stets $\ell^{\perp} \neq \emptyset$, also $p \neq \omega$. p kann auch nicht Null sein, denn: $p = 0 \Rightarrow p_{\perp}^{\perp} = \emptyset^{\perp} = B_X \Rightarrow \ell^{\perp} = B_X$, was aber für kein $\ell \in S_{X_*}$ sein kann.) Man bemerke, dass wegen 1.1.19(d) und der Anmerkung nach (4.27) in 4.7.73 der Träger von ℓ das kleinste Element von $\ell^{\perp} \cap \text{Tri}(X)$ ist; also gilt insbesondere $\ell(p) = 1$, $p \perp \ell$ bezüglich des Kugelsystems des Dualsystems (X, X^*) , das heißt, mit 4.7.13 und Bemerkung 2.11.3 ist ℓ ein Stützfunctional für B_X bei p .

Für jedes $\ell \in S_{X_*}$ gilt die Äquivalenz der folgenden Aussagen:

- (a) ℓ ist ein Atom von X .
- (b) $\ell \perp p$ für ein minimal tripotentes p .
- (c) ℓ ist ein Stützfunctional für B_X bei einem minimal tripotenten p .
- (d) $\ell \perp p$ für $p = \text{spt}_X(\ell)$ und p ist minimal tripotent.
- (e) $\text{spt}_X(\ell)$ ist minimal tripotent.

Beweis. (a) $\Rightarrow$ (b): Wähle gemäß der Anmerkung zu (4.26) dasjenige $p \in \text{Tri}(X)$ mit $p_{\perp} = \{\ell\}$.

(b) $\Rightarrow$ (a): $\ell \perp p \Leftrightarrow p \in \ell^{\perp} \Rightarrow \ell \in \ell^{\perp\perp} \subseteq p_{\perp}$. Da p minimal tripotent ist, ist $p_{\perp}$ nach der Anmerkung zu (4.26) einelementig, also $p_{\perp} = \{\ell\}$, das heißt, ℓ ist ein Atom von X .

(b) $\Leftrightarrow$ (c): Siehe Bemerkung 2.11.3.

(b) $\Rightarrow$ (d): Wie vorhin festgestellt, ist $\text{spt}_X(\ell)$ das kleinste Element von $\ell^{\perp} \cap \text{Tri}(X)$. Es gilt also insbesondere $\text{spt}_X(\ell) \leq p$ und da p ein minimales Element in $\text{Tri}(X)^{\times}$ ist, folgt $\text{spt}_X(\ell) = p$.

Die Implikation (d) $\Rightarrow$ (b) und die Äquivalenz (d) $\Leftrightarrow$ (e) sind klar. $\square$

Für eine weitere äquivalente Aussage siehe Korollar 4.7.83.

Für den Fall, dass X ein JBW^* -Tripel über $\mathbb{C}$ ist, zeigt EDWARDS und RÜTTIMANN [81](1988), Seite 326, Lemma 3.6, dass der Träger von ℓ in X genau das in FRIEDMAN und RUSSO [100](1985), Seite 75, mit $e(\ell)$ bezeichnete eindeutige Tripotent p in X ist, für das sowohl die Gleichung $\ell \circ P_1(p) = \ell$ gilt, als auch $\ell \upharpoonright X_1(p)$ ein treuer, positiver, normaler Zustand auf der JBW^* -Algebra $X_1(p)$ ist (siehe 4.7.15).

Zwei Elemente $x^*, y^* \in S_{X^*}$ heißen *orthogonal* (zueinander), wenn ihre Träger in X orthogonal sind.

(c) Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Sei $x \in S_X$. Man betrachte die Kugelsysteme der Dualsysteme (X, X^*) , (X^*, X^{**}) und (X^{**}, X^{***}) . Es sei an die Gleichung (2.37) in 2.11.13 erinnert. Nach 2.11.7 gilt $x^\perp \in \text{Faces}(w^*; B_{X^*})$, also auch $x^\perp \in \text{Faces}(\|\cdot\|; B_{X^*}) = \text{Faces}(\|\cdot\|; B_{X^{**}})$. Es sei daran erinnert, dass wegen 4.7.28 und 4.7.48 das Bidual X^{**} ein JBW^* -Tripel über $\mathbb{C}$ bzw. ein *JBW -Tripel ist. Wegen des Ordnungsisomorphismus (4.26) aus 4.7.73 existiert genau ein $p \in \text{Tri}(X^{**})^\times$ mit $p_\perp = x^\perp$; es wird der *Träger von x in X^{**}* genannt, in Zeichen $p = \text{spt}_{X^{**}}(x)$. Speziell bemerke man, dass X genau dann glatt bei x ist, wenn $\text{spt}_{X^{**}}(x)$ minimal tripotent in X^{**} ist. Des Weiteren gilt wegen (b) im Fall, dass X bei x glatt ist, für das $\ell \in S_{X^*}$ mit $x^\perp = \{\ell\}$ die Gleichung $\text{spt}_{X^{**}}(\ell) = \text{spt}_{X^{**}}(x)$. Da für jeden normierten Vektorraum Y über $\mathbb{K}$ bezüglich der Dualsysteme (Y, Y^*) und (Y^*, Y^{**}) für alle $y \in Y$ stets die Gleichung $y^\perp = i_Y(y)_\perp$ gilt, hat man im speziellen Fall, dass x selbst bereits ein tripotentes Element ist, die Gleichung $i_X(x) = \text{spt}_{X^{**}}(x)$ vorliegen.

4.7.75. Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Sei $x^* \in X^*$ und p ein tripotentes Element von X mit $\|x^* \circ P_1(p)\| = \|x^*\|$. Dann gilt $x^* = x^* \circ P_1(p)$.

Beweis. Dies wird für JB^* -Tripel über $\mathbb{C}$ in FRIEDMAN und RUSSO [100](1985), Seite 73, Satz 1(a), bewiesen. Nach MARTÍNEZ und PERALTA [206](2000), Seite 642, Lemma 2.9, funktioniert der Beweis von [100] auch für *JB -Tripel. $\square$

In PERALTA und STACHÓ [235](2001), Seite 83, Lemma 2.7, wird für X ein *JB -Tripel allgemeiner gezeigt: Ist $x^* \in S_{X^*}$ und $p \in x^{*-1} \cap \text{Tri}(X)$, so gilt $x^* = x^* \circ P^1(p)$. (Dies ist allgemeiner, da $1 = x^*(p) = x^* \circ P_1(p)p \leq \|x^* \circ P_1(p)\| \leq \|x^*\| \|P_1(p)\| \leq 1$.)

4.7.76 Atomare Projektionen. Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Sei ℓ ein Atom von X und p der Träger von ℓ in X . Betrachte das Dualsystem (X, X_*) . Es sei an 2.2.26 erinnert. Dann ist der Rang-Eins-Operator $p \otimes \ell$ im Fall, dass X ein JBW^* -Tripel über $\mathbb{C}$ ist, gleich der Peirce-Projektion $P_1(p)$, und im Fall, dass X ein *JBW -Tripel ist, gleich der Projektion $P^1(p)$.

Beweis. Der JBW^* -Tripel über $\mathbb{C}$ -Fall findet sich in FRIEDMAN und RUSSO [100](1985), Seite 79, Satz 4(a). In diesem Zusammenhang sei an 3.3.3 und die Formulierung $X_1(p) = \mathbb{C}p$ für die Aussage, dass p divisionsminimal tripotent ist, erinnert.

Betrachte nun den *JBW -Tripel-Fall. Da p nach 4.7.74(b) divisionsminimal tripotent ist, gilt $X^1(p) = \mathbb{R}p$, das heißt, es gibt ein $x^* \in X^*$ mit $P^1(p) = p \otimes x^*$. Nach 4.7.75 gilt nun $\ell = \ell \circ P^1(p)$, womit klar ist, dass man für x^* das Funktional ℓ einsetzen darf. $\square$

4.7.77 Atomare Zerlegung. (a) Sei X ein JBW^* -Tripel über $\mathbb{C}$. Bezeichne X_* sein Prädual. Setze

$$A := \text{cl span ex } B_{X_*},$$

das heißt, A ist die abgeschlossene, lineare Hülle der Atome von X . Nach FRIEDMAN und RUSSO [100](1985), Seite 84, Theorem 1, ist A ein L -Summand von X_* . Bezeichne N den nach 2.5.17 eindeutig bestimmten, zu A komplementären L -Summanden. Falls N von $\{0\}$ verschieden ist, hat die abgeschlossene Einheitskugel B_N keinen Extrempunkt, da dann entweder $X = N$ vorliegt oder andernfalls 2.4.21 greift. Nach 2.4.19 gilt $X \cong A^* \oplus_\infty N^*$. Nach 2.4.20 (dort mit $p = 1$) gilt $A \cong X_*/N$ und $N \cong X_*/A$. Betrachte das Dualsystem (X_*, X) . Da nach 2.4.36(a) $(X_*/N)^* \cong N^\perp$ und $(X_*/A)^* \cong A^\perp$ gilt, ist somit $A^* \cong N^\perp$ und $N^* \cong A^\perp$. Setzt man $\mathcal{A} := N^\perp$ und $\mathcal{N} := A^\perp$, so hat man also

$$X = \mathcal{A} \oplus_\infty \mathcal{N}.$$

Ebd. wird die Gleichung $N = \text{span}\{\ell \in S_{X_*} : \ell \text{ ist orthogonal zu allen Atomen von } X\}$ gezeigt (siehe die Definition in 4.7.74(b)). Setze

$$M := \text{span}\{p \in X : p \text{ minimal tripotent}\}.$$

In ebd., Seite 85, Korollar 2.9, wird die Gleichung

$$N = M_{\perp}$$

gezeigt (man beachte dazu die Gleichung (2.26) in 2.2.27 und, dass $\text{spt}_X \text{ ex } B_{X_*}$ genau gleich der Menge der minimal tripotenten Elemente von X ist). Somit gilt wegen $M \subseteq M_{\perp}^{\perp} = N^{\perp}$ und des Bipolarensatzes die Gleichung

$$\mathcal{A} = \text{cl}(w^*; M).$$

Man bemerke, dass somit $\mathcal{N}$ kein minimal tripotentes Element enthalten kann. In ebd., Seite 86, Theorem 2, wird gezeigt, dass $\mathcal{A}$ ein Tripelideal von X ist. Somit ist $\mathcal{A}$ ein JBW^* -Tripel über $\mathbb{C}$, welches gleich der schwach*-abgeschlossenen Hülle seiner minimal tripotenten Elemente ist. In demselben Theorem von ebd. wird die Gleichung

$$\mathcal{N} = \bigcap \{X_0(p) : p \text{ minimal tripotent in } X\} \quad (4.28)$$

gezeigt. Somit ist $\mathcal{N}$ ein JBW^* -Tripel über $\mathbb{C}$, welches kein minimal tripotentes Element enthält. Mit 4.7.33 und der getrennt schwach*-Stetigkeit des Tripelproduktes von X folgt aus dieser Gleichung unmittelbar, dass $\mathcal{A}$ und $\mathcal{N}$ zueinander orthogonal sind, also $\mathcal{N} = \{x \in X : D(x, y) = 0 \text{ für alle } y \in \mathcal{A}\} = \{x \in X : D(x, p) = 0 \text{ für alle minimal tripotenten Elemente } p \text{ von } X\}$; insbesondere ist somit auch $\mathcal{N}$ ein Tripelideal von X .

(b) Sei X ein JBW^* -Tripel. Setze wie in (a) wieder $M := \text{span}\{p \in X : p \text{ minimal tripotent}\}$ (wobei diesmal natürlich nur reelle Linearkombinationen zu bilden sind). Des Weiteren setze dann wieder $\mathcal{A} = \text{cl}(w^*; M)$. PERALTA und STACHÓ [235](2001), Seite 86, Theorem 3.6, zeigen, dass $\mathcal{A}$ ein M -Summand von X ist. Bezeichne $\mathcal{N}$ den nach 2.5.17 eindeutig bestimmten, zu $\mathcal{A}$ komplementären M -Summanden. Dann wird in ebd. gezeigt, dass $\mathcal{N}$ gleich der zu M orthogonalen Menge ist; insbesondere sind also $\mathcal{A}$ und $\mathcal{N}$ zueinander orthogonal. Als ein M -Summand in einem Dualraum (siehe 2.5.17), aber auch da die Abbildung $D(x) \in L(X)$ für jedes $x \in X$ nach Satz 4.7.49(c) schwach*-stetig ist, ist $\mathcal{N}$ schwach*-abgeschlossen. Betrachte das Dualsystem (X_*, X) . Nach dem Bipolarensatz gilt $\mathcal{A} = \mathcal{A}_{\perp}^{\perp}$ und $\mathcal{N} = \mathcal{N}_{\perp}^{\perp}$. Mit 2.5.12 gilt $X_* = \mathcal{A}_{\perp} \dot{+} \mathcal{N}_{\perp}$ und aus der in 2.5.17 genannten Dualität von $\mathbb{P}$ -Abbildungen folgt $X_* = \mathcal{A}_{\perp} \oplus_1 \mathcal{N}_{\perp}$. Setze $A := \mathcal{N}_{\perp}$ und $N := \mathcal{A}_{\perp}$, also $X_* = A \oplus_1 N$. Wie in (a) sieht man $A^* \cong N^{\perp}$ und $N^* \cong A^{\perp}$; also wieder $\mathcal{A} \cong A^*$ und $\mathcal{N} \cong N^*$.

Da wegen $X_* = (X, \sigma(X, X_*))^*$ alle Elemente von X_* auf X w^* -stetig sind, gilt $\mathcal{A}_{\perp} = M_{\perp}$, also $N = M_{\perp}$.

Sei $\ell \in \text{ex } B_{X_*}$. Dann ist $p := \text{spt}_X(\ell) \in M$; insbesondere ist $\ell(p) = 1$, und daher $\ell \notin M_{\perp}$. Wegen 2.4.21 gilt somit $\text{cl span ex } B_{X_*} \subseteq \mathcal{A}$. Dass auch die umgekehrte Inklusion gilt, sieht man wie folgt: Bezeichne $\bar{}$ die kartesische Involution auf $X_{(\mathbb{C})}$. Nach 3.1.15 und 4.7.52 gilt $X_* = X_{(\mathbb{C})*}(\bar{})$ und $\mathcal{A}_* = \mathcal{A}_{(\mathbb{C})*}(\bar{})$. Wendet man (a) auf $X_{(\mathbb{C})}$ an, so ist die dort betrachtete schwach*-abgeschlossene lineare Hülle der minimal tripotenten Elemente von $X_{(\mathbb{C})}$ (also das dortige „ $\mathcal{A}$ “) nach ebd. gleich $\mathcal{A}_{(\mathbb{C})}$. Somit gilt nach (a): $A = \mathcal{A}_* = (\text{cl span ex } B_{X_{(\mathbb{C})*}})_{(\bar{})} \subseteq (\text{cl span ex } B_{X_*})_{(\bar{})} = \text{cl span ex } B_{X_*}$.

Da der atomare Teil von $X_{(\mathbb{C})}$, soll heißen, die in $X_{(\mathbb{C})}$ schwach*-abgeschlossene Hülle der minimal tripotenten Elemente von $X_{(\mathbb{C})}$, gleich $\mathcal{A}_{(\mathbb{C})}$ ist (siehe ebd.), und $\mathcal{A} = \mathcal{A}_{(\mathbb{C})}(\bar{})$ gilt, ist $\mathcal{A}$ ein Tripelideal von X . Ganz entsprechend sieht man, dass auch $\mathcal{N}$ ein Tripelideal von X ist.

Falls N von $\{0\}$ verschieden ist, also entweder $N = X$ oder 2.4.21 anwendbar ist, gilt $\text{ex } B_N = \emptyset$.

Sei $q \in \text{Tri}(X)$. Wegen $\left((X_{(\mathbb{C})})_0(q)\right)_{(-)} = (X_{(\mathbb{C})(-)})_0(q) = X_0(q)$ und $\mathcal{N} = \left(\bigcap \{(X_{(\mathbb{C})})_0(p) : p \text{ minimal tripotent in } X_{(\mathbb{C})}\}\right)_{(-)}$ gilt auch hier die Gleichung (4.28) von (a).

Insgesamt liegt also eine vollkommen zu (a) analoge Situation vor.

(c) Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Die dann gemäß (a) beziehungsweise (b) vorliegende direkte Summe $X_* = A \oplus_1 N$ heißt die *atomare Zerlegung* (im Englischen *atomic decomposition*) von X_* in den atomaren Teil A und den nicht atomaren Teil N . Man sagt auch, $X = \mathcal{A} \oplus_\infty \mathcal{N}$ sei die *atomare Zerlegung* von X in den atomaren Teil $\mathcal{A}$ und den nicht atomaren Teil $\mathcal{N}$. X heißt *atomar* (im Englischen *atomic*), falls $X = \mathcal{A}$ vorliegt.

4.7.78 Geometrische Charakterisierung. An dieser Stelle sei auf NEAL und RUSSO [223](2004), Seite 586, Theorem 3.1, hingewiesen, wonach ein Banachraum X über $\mathbb{C}$ genau dann isometrisch zu einem JB^* -Tripel über $\mathbb{C}$ ist, wenn sein Dualraum X^* gewisse rein geometrische Eigenschaften aufweist.

4.7.79 Atomar und Radon-Nikodým. (a) Ein JBW^* -Tripel über $\mathbb{C}$ heißt ein *Faktor*, wenn es nicht als eine ℓ_∞ -Summe von zwei von $\{0\}$ verschiedenen Tripelidealen geschrieben werden kann. Solch ein Faktor wiederum heißt *vom Typ I*, wenn er mindestens ein minimal tripotentes Element enthält. Nach FRIEDMAN und RUSSO [101](1986), Seite 144, Satz 2, ist jedes atomare JBW^* -Tripel über $\mathbb{C}$ eine ℓ_∞ -direkte Summe von JBW^* -Tripel-Faktoren vom Typ I. Nach HORN [139](1984), Seite 81, Korollar 9.1.2, sind die JBW^* -Tripel-Faktoren vom Typ I genau die sogenannten Cartan-Faktoren. (Für eine straffe Übersicht über die Cartan-Faktoren siehe FERNÁNDEZ POLO, MARTÍNEZ und PERALTA [93](2004).)

(b) Nach CHU und IOCHUM [49](1987), Seite 463, Theorem 2, besitzt das Prädual eines JBW^* -Tripels über $\mathbb{C}$ genau dann die Radon-Nikodým-Eigenschaft, wenn es eine ℓ_∞ -direkte Summe von Cartan-Faktoren ist.

(c) Sei X ein *JBW -Tripel. Das Prädual des atomaren Anteils der Hermitefizierung $X_{(\mathbb{C})}$ von X , also $\text{cl span ex } B_{X_{(\mathbb{C})}^*}$, hat nach (a) und (b) die Radon-Nikodým-Eigenschaft, und da die Radon-Nikodým-Eigenschaft sich auf Unterräume überträgt, besitzt auch $\mathcal{A}_* = \left(\text{cl span ex } B_{X_{(\mathbb{C})}^*}\right)_{(-)}$ die Radon-Nikodým-Eigenschaft, wobei $-$ die kartesische Involution von $X_{(\mathbb{C})}$ bezeichnet und $\mathcal{A}$ der atomare Teil von X ist.

(d) Zusammenfassend gilt also: Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Dann besitzt das Prädual des atomaren Teils von X die Radon-Nikodým-Eigenschaft

4.7.80 Prädual und Radon-Nikodým. Sei X ein JBW^* -Tripel X über $\mathbb{C}$. BARTON und GODEFROY [13](1986), Seite 300, Theorem 1, haben gezeigt: Das Prädual X_* von X besitzt die Radon-Nikodým-Eigenschaft genau dann, wenn $B_{X_*} \subseteq \text{cl co ex } B_{X_*}$ gilt. Als eine Folgerung aus ihren Beweis ebd., merken sie an, dass das Prädual X_* von X die Radon-Nikodým-Eigenschaft besitzt, wenn die Gleichung $B_{X_*} = \text{cl co ex } B_{X_*}$ gilt. Somit hat man: Das Prädual X_* von X besitzt genau dann die Radon-Nikodým-Eigenschaft, wenn es die Krein-Milman-Eigenschaft aufweist.

4.7.81 Satz. Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Sei $a \in S_X$. Die drei folgenden Aussagen sind äquivalent:

- (a) Die Norm von X ist bei a stark subdifferenzierbar.
- (b) $\text{spt}_{X^{**}}(a) \in i_X(X)$.

(c) Φ_X ist $(n - \overline{w})$ -oberhalbstetig bei a .

Des Weiteren ist die Norm bei allen Tripotenten von X stark subdifferenzierbar.

Beweis. Für X ein JB^* -Tripel über $\mathbb{C}$ siehe BECERRA GUERRERO und RODRÍGUEZ PALACIOS [20](2003), Seite 386, Theorem 2.7.

Für X ein *JB -Tripel siehe BECERRA GUERRERO und PERALTA [19](2004), Seite 506, Korollar 2.2 und Seite 507, Korollar 2.5; für (c) kann man wie im Beweis des JB^* -Tripel über $\mathbb{C}$ Falles argumentieren, da das dazu verwendete Theorem 2.5 aus [20] per Korollar 2.3 aus [19] analog auch für *JB -Tripel gilt. $\square$

4.7.82 Satz (BECERRA GUERRERO und MARTÍN [18](2005), Seite 323, Lemma 3.1). Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Sei $a \in S_X$. Die folgenden beiden Aussagen sind äquivalent:

(a) X ist Fréchet-glatt bei a . (Definition in 4.1.11)

(b) In X existiert ein divisionsminimal tripotentes Element p mit $i_X(p) = \text{spt}_{X^{**}}(a)$.

Beweis. Nach 4.1.7 ist X genau dann Fréchet-glatt bei a , wenn die Norm von X bei a sowohl stark subdifferenzierbar als auch glatt ist. Nach Satz 4.7.81 ist die starke Subdifferenzierbarkeit der Norm bei a äquivalent zu $\text{spt}_{X^{**}}(a) \in i_X(X)$. Nach der Bemerkung in 4.7.74(c) ist die Glattheit der Norm bei a äquivalent dazu, dass $\text{spt}_{X^{**}}(a)$ minimal tripotent in X^{**} ist. Also ist die Fréchet-Glattheit bei a äquivalent dazu, dass in X ein p existiert, für das sowohl $i_X(p) = \text{spt}_{X^{**}}(a)$ gilt als auch $i_X(p)$ minimal tripotent in X^{**} ist. Mit 4.7.70 folgt unmittelbar die Behauptung. $\square$

4.7.83 Korollar. Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Bezeichne X_* ein Prädual von X . Zusätzlich zu den in 4.7.74(b) aufgeführten äquivalenten Aussagen (a) bis (e) hat man: Sei $\ell \in S_{X_*}$. Dann sind äquivalent:

(a) ℓ ist ein Atom von X .

(b) ℓ ist per einem minimal tripotenten Element von X ein stark exponierter Punkt von B_{X_*} .

Beweis. Dass (a) aus (b) folgt ist klar, siehe das Diagramm in 2.11.11. Gelte nun (a). Nach 4.7.74(b) existiert ein minimal tripotentes Element p in X mit $p \in \ell^\perp$. Wie in 4.7.74(c) bemerkt, gilt $i_X(p) = \text{spt}_{X^{**}}(p)$. Nach Satz 4.7.82 heißt dies, dass X bei p Fréchet-glatt ist, was mit 4.1.11 sofort (b) liefert. $\square$

4.7.84 Normierende atomare Zerlegung. Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Nach 4.7.28 bzw. 4.7.48 ist X^{**} ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Mit den Bezeichnungen von 4.7.77 sei $X^{**} = \mathcal{A} \oplus_\infty \mathcal{N}$ die atomare Zerlegung von X^{**} und $X^* = A \oplus_1 N$ die atomare Zerlegung des Präduals von X^{**} . Sei $p_{\mathcal{A}}$ die zu dem M -Summand $\mathcal{A}$ korrespondierende M -Projektion und sei $p_{\mathcal{N}}$ die zu dem M -Summand $\mathcal{N}$ korrespondierende M -Projektion.

(a) (FRIEDMAN und RUSSO [101](1986), Seite 144, Satz 1). Die Abbildung $p_{\mathcal{A}} \circ i_X$ ist ein isometrischer Tripelhomomorphismus, bettet also X isometrisch in das atomare JBW^* -Tripel über $\mathbb{C}$ bzw. *JBW -Tripel ein. Folglich ist A ein normierender Unterraum von X^* .

Beweis. (Ganz nach ebd.) Offensichtlich ist die Abbildung $p_{\mathcal{A}} \circ i_X$ ein kontraktiver Tripelhomomorphismus. Bleibt zu zeigen, dass er isometrisch ist. Sei dazu $x \in S_X$ und $\ell \in \text{ex}(x^\perp)$. Da $x^\perp$ nach 2.11.7 eine extremale Teilmenge von B_{X^*} ist, ist wegen der Kompatibilitätseigenschaft extremer Mengen (siehe 2.2.44) das Funktional ℓ auch ein Extrempunkt von B_{X^*} , das heißt, $\ell \in A$. Wegen $A = \mathcal{N}_\perp$ gilt $\langle \ell, \mathcal{N} \rangle = \{0\}$. Folglich hat

man: $1 = \|x\| = \|i_X(x)\| \geq \|(p_{\mathcal{A}} \circ i_X)(x)\| \geq \ell((p_{\mathcal{A}} \circ i_X)(x)) = \ell(i_X(x)) = \ell(x) = 1$. Da wegen $A^* = \mathcal{A}$ die Abbildung $i_{X,A}$ (siehe 2.4.9) isometrisch ist, ist A ein normierender Unterraum von X^* . Dass A normierend ist, folgt aber bereits auch direkt mittels des Satzes von Krein-Milman aus der Tatsache, dass $\text{ex } B_{X^*} \subseteq B_A$ gilt. $\square$

(b) (BECERRA GUERRERO und MARTÍN [18](2005), Seite 325, Theorem 3.8). Falls X kein divisionsminimal tripotentes Element enthält, ist die Abbildung $p_{\mathcal{N}} \circ i_X$ ein isometrischer Tripelhomomorphismus, bettet also X isometrisch in den nicht atomaren Teil $\mathcal{N}$ von X^{**} ein. Folglich ist dann N ein normierender Unterraum von X^* .

Beweis. Setze $U := \ker(p_{\mathcal{N}} \circ i_X)$, also $U = i_X^{-1}(\mathcal{A})$. Da $\mathcal{A}$ ein abgeschlossenes Tripelideal von X^{**} ist, ist U ein abgeschlossenes Tripelideal von X . Betrachte die beiden Dualsysteme (X, X^*) und (X^*, X^{**}) . Da $\mathcal{A}$ schwach*-abgeschlossen ist, ist der schwach*-Abschluss von $i_X(U)$ in X^{**} , sprich $U^{\perp\perp}$, ein Untertripel von $\mathcal{A}$. Bezeichne $\mathcal{N}_U$ den nicht atomaren Teil von $U^{\perp\perp}$. Da $\mathcal{N}_U = \bigcap \{(U^{\perp\perp})_0(p) : p \text{ minimal tripotent in } U^{\perp\perp}\}$ gilt, gilt wegen 4.7.67 die Gleichung $\mathcal{N}_U = \bigcap \{(U^{\perp\perp})_0(p) : p \text{ minimal tripotent in } X^{**}\}$, also $\mathcal{N}_U = U^{\perp\perp} \cap \bigcap \{X_0(p) : p \text{ minimal tripotent in } X^{**}\}$, das heißt, $\mathcal{N}_U = U^{\perp\perp} \cap \mathcal{N}$, und somit $\mathcal{N}_U = \{0\}$, das heißt, $U^{\perp\perp}$ ist ebenfalls atomar. Wegen $U^{**} \cong (X^*/U^{\perp})^* \cong U^{\perp\perp}$ (konkret: Bezeichnet j die kanonische Einbettung von U in X und j^{**} die duale Abbildung der dualen Abbildung von j , so gilt $j^{**}(i_X(U)) = U^{\perp\perp}$.) ist U^* ein Prädual von $U^{\perp\perp}$; dieses besitzt nach 4.7.79 die Radon-Nikodým-Eigenschaft. Nach 4.1.8 ist U ein Asplundraum. Wäre nun $U \neq \{0\}$, so gäbe es eine in U dichte Menge von Punkten, wo die Norm jeweils Fréchet-differenzierbar wäre. Sei x solch ein, aber von 0 verschiedener, Punkt. Dann wäre die Norm bei $x/\|x\|$ Fréchet-differenzierbar, mit anderen Worten, bei $x/\|x\|$ Fréchet-glatt. Nach Satz 4.7.82 würde somit in U ein divisionsminimal tripotentes Element vorhanden sein und nach 4.7.67 müsste dann aber dieses Element auch ein divisionsminimal tripotentes Element von X sein, im Widerspruch zur anfangs gemachten Voraussetzung. Somit gilt $U = \{0\}$ und die Abbildung $p_{\mathcal{N}} \circ i_X$ ist injektiv.

Da $\mathcal{A}$ und $\mathcal{N}$ zueinander orthogonal sind, ist $p_{\mathcal{N}}$, und somit auch $p_{\mathcal{N}} \circ i_X$, ein Tripelhomomorphismus. Somit ist $(p_{\mathcal{N}} \circ i_X)(X)$ ein Untertripel von $\mathcal{N}$ und da die Abbildung $p_{\mathcal{N}} \circ i_X$ als stetige Abbildung beschränkt ist, ist dieses Untertripel auch abgeschlossen in $\mathcal{N}$, das heißt, es ist ein JB^* -Tripel über $\mathbb{C}$ bzw. ein *JB -Tripel. Nach Satz 4.7.17(a) und 4.7.54 ist die Abbildung $p_{\mathcal{N}} \circ i_X$ eine Isometrie. Ganz analog wie bei (a) sieht man, dass N normierend ist. $\square$

Kapitel 5

Die Daugavet-Eigenschaft

5.1 Definition, Beispiele und einige Eigenschaften von Daugavet-Objekten

5.1.1 Definition (KADETS, SHVIDKOY, SIROTKIN und WERNER [163](2000)). Sei Y ein normierter Vektorraum über $\mathbb{K}$ und X ein abgeschlossener Untervektorraum von Y . Sei $J: X \rightarrow Y$ die kanonische Einbettung. Sei $\mathcal{M}$ eine Teilmenge von $\mathcal{L}(X, Y)$. Man sagt, das Paar (X, Y) besitze die *Daugavet-Eigenschaft bezüglich $\mathcal{M}$* (abgekürzt $DP(\mathcal{M})$, im Englischen *Daugavet property for $\mathcal{M}$*), falls für alle $T \in \mathcal{M}$ die sogenannte *Daugavet-Gleichung*

$$\|J + T\| = 1 + \|T\|$$

erfüllt ist; falls dabei der Fall $X = Y$ vorliegt, sagt man einfach, X besitze *bezüglich $\mathcal{M}$ die Daugavet-Eigenschaft*. Es sei an die in 2.2.26 eingeführte Schreibweise erinnert. Ist U eine Teilmenge von X^* und besitzt dann das Paar (X, Y) die Daugavet-Eigenschaft bezüglich der Menge $\{y \otimes x^* \in L(X, Y) : x^* \in U, y \in Y\}$, so sagt man auch, das Paar (X, Y) (bzw. X , falls $X = Y$) habe die *Daugavet-Eigenschaft bezüglich U* (abgekürzt $DP(U)$); falls dabei $U = X^*$ vorliegt, also das Paar (X, Y) die Daugavet-Eigenschaft bezüglich der Menge aller Rang-Eins Elemente von $\mathcal{L}(X, Y)$ besitzt, lässt man den Bezug auf U in der Regel weg. (Üblicherweise wird die Teilmenge U von X^* ein normierender Unterraum von X^* sein.) X heißt ein *Daugavetraum*, falls X die Daugavet-Eigenschaft besitzt. Setze

$$\mathcal{L}_D(X, Y) := \{T \in \mathcal{L}(X, Y) : \|J + T\| = 1 + \|T\|\} \quad \text{und} \quad \mathcal{L}_D(X) := \mathcal{L}_D(X, X).$$

Die Elemente von $\mathcal{L}_D(X, Y)$ heißen *Daugavet-Operatoren* (für das Paar (X, Y)); $\mathcal{L}_D(X)$ entsprechend.

5.1.2. (a) DAUGAVET [59](1963) zeigte, dass alle kompakten Operatoren aus $\mathcal{L}(C([0, 1], \mathbb{R}))$ die Daugavet-Gleichung erfüllen; siehe WERNER [292] für eine Übertragung von [59] ins Englische. Sei M ein kompakter Hausdorffraum. Setze $Z := C(M, \mathbb{K})$. Sei $T \in \mathcal{L}(Z)$. Dann ist $\|\text{Id}_Z + T\|$ oder $\|\text{Id}_Z - T\|$ gleich $1 + \|T\|$. Genau für den Fall, dass M keinen isolierten Punkt aufweist, ist $C(M, \mathbb{K})$ ein Daugavetraum. Ist Y ein Banachraum über $\mathbb{K}$, so ist $C(M, Y)$ ein Daugavetraum, wenn $C(M, \mathbb{K})$ oder Y ein Daugavetraum ist. Ist (X, Σ, μ) ein Maßraum, so sind die Funktionenräume $L_1(X, \Sigma, \mu)$, $L_{1(\mathbb{C})}(X, \Sigma, \mu)$, $L_\infty(X, \Sigma, \mu)$ und $L_{\infty(\mathbb{C})}(X, \Sigma, \mu)$ jeweils genau dann Daugaveträume, wenn der Maßraum (X, Σ, μ) atomlos ist. Die Folgenräume c_0 und c sind keine Daugaveträume. Da für jedes $p \in \mathbb{N}$ mit $1 \leq p \leq \infty$ der Folgenraum ℓ_p gleich $L_p(\mathbb{N}, \mathcal{P}\mathbb{N}, \mu)$, μ das *Zählmaß* (im Englischen *counting measure*; für $A \subseteq \mathbb{N}$ ist $\mu(A)$ gleich der Anzahl der Elemente von A , wenn A endlich ist, sonst unendlich), und jedes $n \in \mathbb{N}$ ein Atom ist (μ ist also sogar rein atomar), sind die beiden Folgenräume ℓ_1 und ℓ_∞ keine Daugaveträume.

(b) Sei X ein normierter Vektorraum über $\mathbb{K}$. Dann gilt wegen 2.2.4 für alle $\lambda \geq 0$ die Inklusion $\lambda \cdot \mathcal{L}_D(X) \subseteq \mathcal{L}_D(X)$.

(c) Sei X ein Banachraum über $\mathbb{K}$, der die Daugavet-Eigenschaft besitzt. Dann hat die Einheitskugel B_X keinen stark exponierten Punkt. Insbesondere besitzt X also nicht die Radon-Nikodým-Eigenschaft und ist somit auch nicht reflexiv, geschweige endlichdimensional.

(d) Sei (X, Y) ein Paar von Banachräumen über $\mathbb{K}$. Setze für diese Nummer als eine Vereinfachung der Schreibweise

$$A := \{T \in \mathcal{L}(X, Y) : T \text{ ist stark Radon-Nikodým}\} \quad \text{und} \\ B := \{T \in \mathcal{L}(X, Y) : T \text{ fixiert nicht } \ell_1\}.$$

Für jedes $\lambda \in \mathbb{K}$ gilt zwar $\lambda \cdot A \subseteq A$, aber im Allgemeinen ist A kein Vektorraum. B hingegen ist ein Vektorraum über $\mathbb{K}$. Für den Fall $X = Y$ ist bekannt, dass $K_w(X) \subseteq A \cap B$ gilt. Liege nun der Fall vor, dass das Paar (X, Y) die Daugavet-Eigenschaft besitze. Dann gilt $A \cup B \subseteq \mathcal{L}_D(X, Y)$. Im Fall von $X = Y$ gilt $\text{span}(A + B) \subseteq \mathcal{L}_D(X)$.

Die beiden entsprechenden Aussagen für den Fall $X \neq Y$ der beiden eben gemachten Aussagen für den Fall $X = Y$ gelten im Allgemeinen nicht.

(e) Für jeden Banachraum X über $\mathbb{C}$ ist jedes $T \in \mathcal{L}(X)$ mit $\|T\| \in \sigma(T)$ ein Daugavet-Operator.

(f) Sei X ein Banachraum über $\mathbb{K}$. Falls X^* ein Daugavetraum ist, so ist auch X ein Daugavetraum. Für den Funktionenraum $C[0, 1]$ gilt die Umkehrung dieser Aussage nicht.

(g) Ein Banachraum X über $\mathbb{C}$ ist genau dann ein Daugavetraum, wenn für alle $x^* \in X^*$ und $x \in X_{\mathbb{R}}$ der Operator $x \otimes \text{Re } x^*$ aus $L(X_{\mathbb{R}})$ ein Element von $\mathcal{L}_D(X_{\mathbb{R}})$ ist. Das heißt, X ist genau dann ein Daugavetraum, wenn $X_{\mathbb{R}}$ einer ist.

(h) Sei X ein Banachraum über $\mathbb{K}$, der die Daugavet-Eigenschaft besitzt. Dann ist jedes M -Ideal von X ein Daugavetraum.

(i) Sei X ein Banachraum über $\mathbb{K}$ und $T \in \mathcal{L}(X)$. Dann sind die folgenden beiden Aussagen äquivalent:

(α) $T \in \mathcal{L}_D(X)$.

(β) $\sup \text{Re } V_{\text{spatial}}(T) = \|T\|$. (Definition 2.11.15)

(j) Sei X ein Banachraum über $\mathbb{K}$, der die Daugavet-Eigenschaft besitzt. Dann enthält X einen zu dem Folgenraum ℓ_1 isomorphen Untervektorraum, und X^* sogar einen zu ℓ_1 isometrisch isomorphen Untervektorraum.

Beweis. (a) HOLUB [138](1986), Seite 398. $C(M, Y)$: KADETS, KALTON und WERNER [160](2003), Seite 202. $L_1(\mu)$: WERNER [289](1996), Seite 453, Korollar 6. $L_\infty(\mu)$: KADETS, SHVIDKOY und WERNER [164](2001), Seite 269. Bezüglich $L_1(\mu)$ und $L_\infty(\mu)$ siehe auch die Einleitung in ABRAMOVICH [1](1991). c_0, c : MARTÍN [204](2003), Seite 635, Beispiel 4. (b) ABRAMOVICH, ALIPRANTIS und BURKINSHAW [3](1991), Seite 217, Lemma 2.1. (c) WOJTASZCZYK [293](1992), Seite 1048, Satz 1. (d) Bezüglich stark Radon-Nikodým siehe KADETS, SHVIDKOY, SIROTKIN und WERNER [163](2000), Seite 857. Bezüglich ℓ_1 siehe SHVIDKOY [260](2000), Seite 202, Theorem 4. (e) ABRAMOVICH und ALIPRANTIS [2](2002), Seite 456, Lemma 11.3. (f) KADETS, SHVIDKOY, SIROTKIN und WERNER [163](2000), Seite 857. (g) KADETS, MARTÍN und MERÍ [161](2007), Seite 2388, Bemerkung 2.1. (h) Für $\mathbb{K} = \mathbb{R}$ siehe KADETS, SHVIDKOY, SIROTKIN und WERNER [163](2000), Seite 861, Satz 2.10. Der $\mathbb{K} = \mathbb{C}$ -Fall folgt dann aus (g) und der Bemerkung über die $\mathbb{K}$ -Unabhängigkeit von $\mathbb{P}$ -Idealen in 2.5.17. (i) MARTÍN und OIKHBERG [205](2004), Seite 161, Lemma 2.3(a). (j) KADETS, SHVIDKOY, SIROTKIN und WERNER [163](2000), Seite 861, Theorem 2.9 und Seite 864, Korollar 2.13. $\square$

5.1.3. Die Daugavet-Eigenschaft lässt sich mit den in 2.6.13 definierten Scheiben und sogar allgemeiner mit schwach offenen Teilmengen formulieren. Sei (X, Y) ein Paar von Banachräumen über $\mathbb{K}$. Dann sind die folgenden Aussagen äquivalent:

- (a) (X, Y) besitzt die Daugavet-Eigenschaft.
- (b) Für jedes $y \in S_Y$, jedes $x^* \in S_{X^*}$ und jedes $\varepsilon > 0$ existiert ein $z^* \in S_{X^*}$ und ein $\eta > 0$, so dass $S(B_X, z^*, \eta) \subseteq S(B_X, x^*, \varepsilon)$ gilt und für alle $x \in S(B_X, z^*, \eta)$ die Ungleichung $\|x + y\| \geq 2 - \varepsilon$ erfüllt ist.
- (c) Für jedes $x^* \in S_{X^*}$, jedes $y \in S_Y$ und jedes $\varepsilon > 0$ existiert ein $z \in S_Y$ und ein $\eta > 0$, so dass $S(B_{Y^*}, z, \eta) \subseteq S(B_{Y^*}, y, \varepsilon)$ gilt und für alle $y^* \in S(B_{Y^*}, z, \eta)$ die Ungleichung $\|x^* + y^* \upharpoonright X\| \geq 2 - \varepsilon$ erfüllt ist.
- (d) Für jedes $y \in S_Y$, jedes $x^* \in S_{X^*}$ und jedes $\varepsilon > 0$ existiert ein $x \in S_X \cap S(B_X, x^*, \varepsilon)$, so dass $\|x + y\| \geq 2 - \varepsilon$ gilt.
- (e) Für jedes $y \in S_Y$, jedes $x^* \in S_{X^*}$ und jedes $\varepsilon > 0$ existiert ein $y^* \in S_{Y^*} \cap S(B_{Y^*}, y, \varepsilon)$, so dass $\|x^* + y^* \upharpoonright X\| \geq 2 - \varepsilon$ gilt.
- (f) Für jedes $y \in S_Y$, jedes $\varepsilon > 0$ und jede relativ zu X schwach offene, nicht leere Teilmenge U von B_X existiert eine relativ zu X schwach offene, nicht leere Teilmenge V von B_X mit $V \subseteq U$ und $\|x + y\| \geq 2 - \varepsilon$ für alle $x \in V$.
- (g) Für jedes $x^* \in S_{X^*}$, jedes $\varepsilon > 0$ und jede relativ schwach* offene, nicht leere Teilmenge U von B_{Y^*} existiert eine relativ schwach* offene, nicht leere Teilmenge V von B_{Y^*} mit $V \subseteq U$ und $\|x^* + y^* \upharpoonright X\| \geq 2 - \varepsilon$ für alle $y^* \in V$.

Beweis. (a) bis (e): KADETS, SHVIDKOY, SIROTKIN und WERNER [163](2000), Seite 857, Lemmata 2.1 und 2.2. (a), (f) und (g): SHVIDKOY [260](2000), Lemma 3. $\square$

5.1.4 Bemerkungen. Im Folgenden werden einige Möglichkeiten genannt, wie man weitere zu 5.1.3(a) äquivalente Aussagen — die fortan meist ohne explizite Erwähnung verwendet werden — erhalten kann.

(a) In 5.1.3 kann man statt $y^* \upharpoonright X$ auch $J^*(y^*)$ schreiben.

(b) In 5.1.3(b) und (c) kann man statt $\exists \eta > 0$ auch $\exists \eta \in (0, 1)$ schreiben, wenn man statt $\forall \varepsilon > 0$ den Ausdruck $\forall \varepsilon \in (0, 1)$ verwendet.

(c) In 5.1.3(b), (d) und (f) kann man statt $\|x + y\| \geq 2 - \varepsilon$ auch $\|x - y\| \geq 2 - \varepsilon$ schreiben.

(d) Analog kann man in 5.1.3(c), (e) und (g) statt $\|x^* + y^* \upharpoonright X\| \geq 2 - \varepsilon$ auch $\|x^* - y^* \upharpoonright X\| \geq 2 - \varepsilon$ schreiben.

Beweis. (a) Klar. (b) ABRAMOVICH und ALIPRANTIS [2](2002), Seite 488, Lemma 11.46. (c) und (d) Klar. $\square$

5.1.5 Bemerkung (KADETS, SHVIDKOY, SIROTKIN und WERNER [163](2000), Seite 3). Sei (X, Y) ein Paar von Banachräumen über $\mathbb{K}$, welches die Daugavet-Eigenschaft besitze. Als eine Folgerung von 5.1.4(c) und (d) ist der Durchmesser von sowohl jeder Scheibe von B_X als auch von jeder schwach*-Scheibe von B_{X^*} gleich 2. Damit hat man einerseits wieder die Bemerkung von 5.1.2(c), dass X nicht die Radon-Nikodým-Eigenschaft aufweist und andererseits, dass auch X^* nicht die Radon-Nikodým-Eigenschaft aufweist. Wegen der genannten Durchmesser-Eigenschaft von B_{X^*} ist nach 4.1.11 die Norm von X extrem rau.

Beweis. (Nach ABRAMOVICH und ALIPRANTIS [2](2002), Seite 490, Korollar 11.47.) Sei $x^* \in S_{X^*}$ und $\varepsilon > 0$. Sei $0 < \mu < \varepsilon$ beliebig und dann $y \in S_Y \cap S(B_X, x^*, \mu)$. Nach 5.1.3(b) existiert ein $z^* \in S_{X^*}$ und ein $\eta > 0$ mit $S(B_X, z^*, \eta) \subseteq S(B_X, x^*, \mu)$ und für alle $x \in S(B_X, z^*, \eta)$ gilt $\|x - y\| \geq 2 - \mu$. Da nun x und y beide Elemente von $S(B_X, x^*, \varepsilon)$ sind, folgt, dass $S(B_X, x^*, \varepsilon)$ Durchmesser 2 hat.

Die Aussage für B_{X^*} geht analog: Sei $x \in S_X$ und $\varepsilon > 0$ und dann $0 < \mu < \varepsilon$ und $x^* \in S_{X^*} \cap S(B_{X^*}, x, \mu)$. Nach 5.1.3(c) existiert ein $z \in S_Y$ und ein $\eta > 0$ mit $S(B_{Y^*}, z, \eta) \subseteq S(B_{Y^*}, x, \mu)$ und für alle $u^* \in S(B_{Y^*}, z, \eta)$ gilt $\|x^* - J^*(u^*)\| \geq 2 - \mu$. Da x^* und $J^*(u^*)$ beide Elemente von $S(B_{X^*}, x, \varepsilon)$ sind, folgt, dass $S(B_{X^*}, x, \varepsilon)$ Durchmesser 2 hat. $\square$

Mit einer zu dem ebenen Beweis ganz analogen Argumentation, angewandt auf (f) und (g) von 5.1.3, erhält man offenbar die allgemeinere Aussage, dass der Durchmesser von sowohl jeder schwach offenen, nicht leeren Teilmenge von B_X als auch von jeder schwach* offenen, nicht leeren Teilmenge von B_{X^*} gleich 2 ist.

5.1.6 Schmale Operatoren.

(KADETS, KALTON und WERNER [160](2003)). Seien X und Y zwei Banachräume über $\mathbb{K}$. Dann heißt ein $T \in \mathcal{L}(X, Y)$ *schmal* (im Englischen *narrow*), wenn für jedes $x \in S_X$, jedes $y \in S_X$, jede Scheibe S von B_X , die y enthält, und jedes $\varepsilon > 0$ ein $v \in S$ existiert, so dass $\|x + v\| \geq 2 - \varepsilon$ und $\|T(y - v)\| \leq \varepsilon$ gilt. Die schmalen Operatoren wurden in KADETS, SHVIDKOY und WERNER [164](2001) eingeführt.

(BILIK, KADETS, SHVIDKOY und WERNER [24](2005)). Ein $T \in \mathcal{L}(X, Y)$ ist genau dann schmal, wenn für jedes $x \in S_X$, jedes $y \in S_X$, jedes $x^* \in X^*$ und jedes $\varepsilon > 0$ ein $z \in S_X$ existiert, so dass $\|x + z\| \geq 2 - \varepsilon$ und $\|T(y - z)\| + |x^*(y - z)| \leq \varepsilon$ gilt.

Sei X ein Banachraum über $\mathbb{K}$, der die Daugavet-Eigenschaft besitzt. Dann gilt mit den Bezeichnungen von 5.1.2(d), dass $\text{span}(A + B)$ eine Teilmenge der Menge aller schmalen Elemente von $\mathcal{L}(X)$ ist.

5.1.7. Sei X ein Banachraum über $\mathbb{K}$. Dann sind die folgenden Aussagen äquivalent:

- (a) X ist ein Daugavetraum.
- (b) Für jedes $x \in S_X$, jede relativ schwach offene Teilmenge U von B_X und jedes $\varepsilon > 0$ existiert eine relativ schwach offene Teilmenge V von B_X mit $V \subseteq U$ und $\|x + y\| \geq 2 - \varepsilon$ für alle $y \in V$.
- (c) Für jedes $x \in S_X$, jede relativ schwach offene, nicht leere Teilmenge U von B_X und jedes $\varepsilon > 0$ existiert ein $y \in U$ mit $\|x + y\| \geq 2 - \varepsilon$.
- (d) Für jedes $x \in S_X$ und $\varepsilon > 0$ gilt $B_X = \text{cl co} \{y \in B_X : \|x + y\| \geq 2 - \varepsilon\}$.
- (e) Für jedes $x \in S_X$ und $\varepsilon > 0$ gilt $B_X = \text{cl co} (B_X \setminus (x + (2 - \varepsilon) \cdot B_X))$.
- (f) Für jedes $x^* \in S_{X^*}$, jeder schwach*-abgeschlossenen Scheibe S von B_{X^*} und jedes $\varepsilon > 0$ existiert eine Scheibe T von B_{X^*} mit $T \subseteq S$ und $\|x^* + y^*\| \geq 2 - \varepsilon$ für alle $y^* \in T$.
- (g) $0 \in \mathcal{L}(X)$ ist schmal.
- (h) $\mathcal{L}(X)$ enthält mindestens ein schmales Element.

Beweis. (a) $\Leftrightarrow$ (b): WERNER [290](2001), Lemma 2.2. (a) $\Leftrightarrow$ (c): KADETS, SHVIDKOY und WERNER [164](2001), Lemma 1.1. (a) $\Leftrightarrow$ (d): WERNER [290](2001), Korollar 2.3. (d) $\Leftrightarrow$ (e): Sei $x \in S_X, \varepsilon > 0$ und $y \in B_X$. Dann gilt: $\|y - x\| > 2 - \varepsilon \Leftrightarrow$ Für jedes $z \in B_X$

gilt $y - x \neq (2 - \varepsilon)z \Leftrightarrow$ Für jedes $z \in B_X$ gilt $y \neq x + (2 - \varepsilon)z \Leftrightarrow y \notin x + (2 - \varepsilon) \cdot B_X$.
 (a) $\Leftrightarrow$ (f): WERNER [290](2001), Lemma 2.4. (a), (g), (h): KADETS, KALTON und WERNER [160](2003). $\square$

5.1.8. Sei X ein Banachraum über $\mathbb{K}$ und V ein normierender Unterraum von X^* . Dann sind äquivalent:

- (a) X hat die Daugavet-Eigenschaft bezüglich V .
- (b) Für jedes $x \in S_X$, jedes $x^* \in S_V$ und jedes $\varepsilon > 0$ existiert ein $y \in S(B_X, x^*, \varepsilon)$ für das $\|x + y\| \geq 2 - \varepsilon$ gilt.
- (c) Für jedes $x \in S_X$, jedes $x^* \in S_V$ und jedes $\varepsilon > 0$ existiert ein $y^* \in S_V$ und ein $\eta > 0$, so dass $S(B_X, y^*, \eta) \subseteq S(B_X, x^*, \varepsilon)$ gilt und für alle $y \in S(B_X, y^*, \eta)$ die Ungleichung $\|x + y\| \geq 2 - \varepsilon$ erfüllt ist.

Beweis. KADETS, SHEPELSKA und WERNER [162](2008), Theorem 2.2. $\square$

5.1.9 Satz (BECERRA GUERRERO und MARTÍN [18](2005), Seite 320, Theorem 2.2). *Sei X ein von $\{0\}$ verschiedener Banachraum über $\mathbb{K}$ für den zwei normierende Untervektorräume Y und Z von X^* mit $X^* = Y \oplus_1 Z$ existieren. Dann ist X ein Daugavetraum.*

Beweis. (Ganz nach ebd.) Sei $x_0 \in S_X$, $f_0 \in S_{X^*}$ und $\varepsilon > 0$. Schreibe $f_0 = y_0 + z_0$ mit $y_0 \in Y$ und $z_0 \in Z$. Also $\|f_0\| = \|y_0\| + \|z_0\|$. Setze $U := S(B_{X^*}, x_0, \varepsilon)$. Nach 2.5.8 ist B_Z schwach*-dicht in B_{X^*} . Und da U in B_{X^*} eine schwach*-offene, nicht leere Teilmenge in B_{X^*} ist, existiert ein $z \in B_Z \cap U$. Wegen $\|z\| \geq \sup\{|z(x)| : x \in S_X\} = \sup\{\operatorname{Re} z(x) : x \in S_X\}$ gilt $\|z\| \geq 1 - \varepsilon$. Wieder nach 2.5.8 ist auch B_Y schwach*-dicht in B_{X^*} . Daher existiert ein in B_Y liegendes Netz $(y_\nu)_{\nu \in I}$, das schwach* gegen z konvergiert. Da z ein Element der in B_{X^*} schwach*-offenen Scheibe U ist, kann man ohne Einschränkung annehmen, dass das Netz $(y_\nu)_{\nu \in I}$ ganz in U liegt. Nach 4.1.12(a) ist die Norm von X^* schwach*-unterhalbstetig. Da nun das Netz $(y_\nu + y_0)_{\nu \in I}$ schwach* gegen $z + y_0$ konvergiert, gilt gemäß 1.2.9 die Ungleichung $\liminf_\nu \|y_\nu + y_0\| \geq \|z + y_0\|$. Wegen $\|z + y_0\| = \|z\| + \|y_0\| > 1 + \|y_0\| - \varepsilon$, also $\liminf_\nu \|y_\nu + y_0\| > 1 + \|y_0\| - \varepsilon$. Setze $\Delta := \|z + y_0\| - (1 + \|y_0\| - \varepsilon)$, also $\Delta = \|z\| - (1 - \varepsilon)$ und $0 < \Delta \leq \varepsilon$. Nach Definition des Limes inferior existiert ein $\nu_0 \in I$ mit $\inf\{\|y_\mu + y_0\| : \nu_0 \leq \mu\} \geq \liminf_\nu \|y_\nu + y_0\| - \Delta$, das heißt, es existiert ein $\mu \in I$ mit $\|y_\mu + y_0\| \geq \|z + y_0\| - \Delta = 1 + \|y_0\| - \varepsilon$. Somit hat man einerseits $\|f_0 + y_\mu\| = \|y_0 + z_0 + y_\mu\| = \|(y_0 + y_\mu) + z_0\| = \|y_0 + y_\mu\| + \|z_0\| \geq 1 + \|y_0\| + \|z_0\| - \varepsilon = 1 + \|y_0 + z_0\| - \varepsilon = 1 + \|f_0\| - \varepsilon = 2 - \varepsilon$ und andererseits, da $y_\mu \in U$ ist, $\operatorname{Re} y_\mu(x_0) > 1 - \varepsilon$. Mit der Besetzung von y als x_0 , x^* als f_0 und y^* als y_μ gibt 5.1.3(e) die Behauptung. $\square$

5.1.10 Korollar (BECERRA GUERRERO und MARTÍN [18](2005), Seite 320, Korollar 2.3). *Sei X ein L -eingebetteter Banachraum über $\mathbb{K}$ mit $\operatorname{ex} B_X = \emptyset$. Dann ist X^* (und somit auch X) ein Daugavetraum.*

Beweis. (Ganz nach ebd.) Sei $\operatorname{ex} B_X = \emptyset$. Folglich ist X nicht reflexiv. Mithin sei Z ein von $\{0\}$ verschiedener Untervektorraum von X^{**} mit $X^{**} = i_X(X) \oplus_1 Z$. Mit 2.4.21 hat man somit $\operatorname{ex} B_{X^{**}} = \operatorname{ex} B_Z$. Nach dem Satz von Krein-Milman 2.11.11 gilt somit $B_{X^{**}} = \operatorname{cl}(w^*; \operatorname{co} \operatorname{ex} B_{X^{**}}) = \operatorname{cl}(w^*; \operatorname{co} \operatorname{ex} B_Z) = \operatorname{cl}(w^*; B_Z)$, das heißt, B_Z ist in $B_{X^{**}}$ schwach*-dicht. Nach dem Satz von Goldstine 2.5.11(b) ist $i_X(X)$ ebenfalls schwach*-dicht in $B_{X^{**}}$. Nach 2.5.8 sind also $i_X(X)$ und Z beide für X^* normierende Untervektorräume von X^{**} und mit Satz 5.1.9 ist X^* ein Daugavetraum. Dass dann auch X ein Daugavetraum ist, ist in 5.1.2(f) bemerkt worden. $\square$

5.2 Die Daugavet-Eigenschaft in JB^* -Tripeln

5.2.1 Satz (BECERRA GUERRERO und MARTÍN [18](2005), Satz 3.9). *Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Falls X von $\{0\}$ verschieden ist, ist X genau dann ein Daugavetraum, wenn X kein divisionsminimal tripotentes Element besitzt.*

Beweis. Ist X von $\{0\}$ verschieden und besitzt es keine divisionsminimal tripotenten Elemente, so folgt mit 4.7.84 und Satz 5.1.9, dass X ein Daugavetraum ist. Besitzt andererseits X ein divisionsminimal tripotentes Element, etwa p , so ist die Norm von X nach Satz 4.7.82 bei p Fréchet-glatt, das heißt, die Rauheit von X bei p ist gleich 0. Da aber nach Bemerkung 5.1.5 die Norm eines Daugavetraumes extrem rau ist, also die Rauheit von X bei jedem $x \in S_X$ gleich 2 ist, kann X dann kein Daugavetraum sein. $\square$

Mit 4.7.69 hat man somit unmittelbar das

5.2.2 Korollar. *Eine von $\{0\}$ verschiedene C^* -Algebra über $\mathbb{K}$ ist genau dann ein Daugavetraum, wenn sie nicht atomar ist.*

5.2.3 Von Neumann Typen. Sei H ein Hilbertraum über $\mathbb{K}$ und A eine von Neumann-Algebra auf H über $\mathbb{K}$. Ist A vom Typ II, Typ III oder eine direkte Summe von einer Typ II- und einer Typ III-von Neumann-Algebra, so weist A keine Atome auf, siehe 3.7.58, 3.7.59, 3.7.60, 4.7.68 und 4.7.69; somit ist also jedes solch ein A ein Daugavetraum.

5.2.4 Beispiele von Faktoren vom Typ II und III (FILLMORE [96](1996), Seiten 61, 62 und 81; FREMLIN [99](2006), 255M und Kapitel 44; JONES und SUNDER [159](1997), Kapitel 1; LARSEN [196] (1973), Seite 21, Beispiel 1.2.9; LI [197](1992), Kapitel 7.3; SUNDER [271](1987), Seite 121, Beispiel 4.1.14, und Seite 146; TAKESAKI [274](2001), Abschnitt V.7).

Nach 5.2.3 ist jede von Neumann-Algebra über $\mathbb{C}$, die ein Faktor vom Typ II_1 , II_∞ oder III ist, ein Daugavetraum. Im Folgenden werden einige konkrete Beispiele solcher Faktoren angeführt.

(a) Sei M eine von Neumann-Algebra über $\mathbb{C}$ auf einem Hilbertraum H , also $M \subseteq \mathcal{L}(H)$. Sei G eine Gruppe, die auf M operiert, soll heißen, es gibt einen Gruppenhomomorphismus r von G in die Gruppe aller * -Automorphismen von M ; als Operation von G auf M bezeichne dann α die mit der Darstellung r assoziierte Auswertungsabbildung, also $\alpha: G \times M \rightarrow M$, $\alpha(g, m) := r(g)(m)$. Betrachte das Hilbertraum-Tensorprodukt $H \otimes \ell_2(G)$ von H und $\ell_2(G)$, wobei $\ell_2(G)$ der Hilbertraum aller quadratsummierbaren, komplexwertigen Funktionen auf G ist; dazu unitär äquivalent wird $H \otimes \ell_2(G)$ hier aufgefasst als der Hilbertraum $\mathcal{H} := \ell_2(G, H)$ aller quadratsummierbaren H -wertigen Funktionen auf G , also

$$\mathcal{H} = \left\{ \varphi: G \rightarrow H : \sum_{g \in G} \|\varphi(g)\|^2 < \infty \right\}.$$

(Für manche Betrachtungen ist es oftmals praktisch die Elemente von $\mathcal{H}$ aufzufassen als Spaltenvektoren mit normquadratsummierbaren Einträgen von H , in der in 2.4.19 eingeführten Schreibweise also als $\ell_2\left(\bigoplus_{g \in G} H_g\right)$, wobei $H_g := H$ für alle $g \in G$ gesetzt worden ist.) Durch

$$(\pi(x)\varphi)(g) := \alpha(g^{-1}, x)\varphi(g) \quad \text{für alle } x \in M, \varphi \in \mathcal{H}, g \in G$$

ist $\pi: M \rightarrow \mathcal{L}(\mathcal{H})$ eine treue W^* -Darstellung von M auf $\mathcal{H}$, soll heißen, ein $w^* - w^*$ -stetiger, injektiver * -Homomorphismus von M in $\mathcal{L}(\mathcal{H})$. Durch

$$(\lambda(g)\varphi)(h) := \varphi(g^{-1}h) \quad \text{für alle } g, h \in G, \varphi \in \mathcal{H}$$

ist $\lambda: G \rightarrow \mathcal{L}(\mathcal{H})$ eine treue, unitäre Darstellung von G auf $\mathcal{H}$, soll heißen, ein injektiver Gruppenhomomorphismus von G in die Gruppe der unitären Elemente von $\mathcal{L}(\mathcal{H})$. Dann wird die kleinste von Neumann-Algebra auf $\mathcal{H}$, die sowohl $\pi(M)$ als auch $\lambda(G)$ umfasst, also $(\pi(M) \cup \lambda(G))''$, mit $M \times_\alpha G$ bezeichnet (andere auch übliche Bezeichnungen sind $M \rtimes_\alpha G$, $M \otimes_\alpha G$ und $\mathcal{R}(M, G, \alpha)$); sie heißt das (*per* α) *verschränkte Produkt* (im Englischen *crossed product*) von M und G (bei TAKESAKI, ebd., auch die *Kovarianz-von Neumann Algebra von* (M, G, α)).

(b) Sei $M := \{m_z \in \mathcal{L}(\mathbb{C}) : z \in \mathbb{C}\}$, wobei $m_z(w) := zw$ für alle $z, w \in \mathbb{C}$ gesetzt sei, die von Neumann-Algebra über $\mathbb{C}$ auf dem Hilbertraum $H := \mathbb{C}$, die durch die kanonische Einbettung von $\mathbb{C}$ in $\mathcal{L}(\mathbb{C})$ entsteht. Sei G eine abzählbare ICC-Gruppe (siehe 4.3.9) mit $G \neq \{e\}$. Sei α die triviale Operation von G auf M , also

$$\alpha(g, m_z) = m_z \quad \text{für alle } g \in G, z \in \mathbb{C}.$$

Setze $\mathcal{H}$, π und λ wie in (a). Dann ist $(\pi(m_z)\varphi)(g) = m_z(\varphi(g)) = z \cdot \varphi(g)$ für alle $z \in \mathbb{C}$, $\varphi \in \mathcal{H}$, $g \in G$, also

$$\pi(m_z)\varphi = z\varphi \quad \text{für alle } z \in \mathbb{C}, \varphi \in \mathcal{H},$$

das heißt, π entspricht der kanonischen Einbettung von $\mathbb{C}$ in $\mathcal{L}(\mathcal{H})$. Dann ist die von Neumann-Algebra $M \times_\alpha G = (\pi(M) \cup \lambda(G))'' = (\mathbb{C} \cdot \text{Id}_{\mathcal{H}} \cup \lambda(G))'' = \lambda(G)'' \subseteq \mathcal{L}(\mathcal{H})$ ein Faktor vom **Typ II**₁, die sogenannte *Linksgruppen-von Neumann-Algebra der Gruppe* G ; man symbolisiert sie auch per $\mathbb{C} \rtimes_\alpha G$, $\mathcal{R}(G)$ oder $W^*(G)$.

(c) Sei Ω eine lokal kompakte, Hausdorff'sche topologische Gruppe, die das zweite Abzählbarkeitsaxiom erfüllt (sprich, deren Topologie eine abzählbare Basis besitzt), versehen mit einem linken Haarmaß μ auf der Borel- σ -Algebra Σ von Ω . Dann ist Ω separabel und σ -finit. Sei G eine abzählbare, dichte Untergruppe von Ω . Setze

$$\beta_g(h) := \beta(g, h) := gh \quad \text{für alle } g \in G, h \in \Omega.$$

Die Gruppe G operiert also auf der Gruppe Ω per Linkstranslationen β_g , $g \in G$. Bezeichne für zwei Elemente $A, B \in \Sigma$ mit $A \triangle B$ die symmetrische Differenz $(A \cup B) \setminus (A \cap B)$ von A und B . Die Linkstranslationen β_g , $g \in G$, operieren ergodisch — soll heißen, dass für jedes $A \in \Sigma$ die Gültigkeit der Gleichung $\mu(\beta_g(A) \triangle A) = 0$ für alle $g \in G$ stets $\mu(A) = 0$ oder $\mu(\Omega \setminus A) = 0$ impliziert — und frei — soll heißen, dass für jedes $g \in G \setminus \{e\}$ die Gleichung $\mu(\{h \in \Omega : \beta_g(h) = h\}) = 0$ gilt — von G auf dem Maßraum (Ω, Σ, μ) per maßerhaltenden Automorphismen — soll heißen, per Bijektionen $T: \Omega \rightarrow \Omega$ mit (i) $A \in \Sigma$ impliziert $T(A), T^{-1}(A) \in \Sigma$ und (ii) für alle $A \in \Sigma$ gilt genau dann $\mu(A) = 0$, wenn $\mu(T^{-1}(A)) = 0$ gilt —. Setze

$$H := L_2(\mathbb{C})(\Omega, \Sigma, \mu)$$

und dann

$$M := \left\{ m_f \in \mathcal{L}(H) : f \in L_\infty(\mathbb{C})(\Omega, \Sigma, \mu) \right\},$$

wobei $m_f(g) := f \cdot g$ für alle $f \in L_\infty(\mathbb{C})(\Omega, \Sigma, \mu)$, $g \in H$ gesetzt worden sei. M ist eine maximal abelsche (heißt, $M = M'$) von Neumann-Algebra (also eine so genannte *masa*) auf dem Hilbertraum H . Sei α die von β induzierte (wegen der Separabilität und σ -Finitheit von Ω eindeutig bestimmte) Operation von G auf M , das heißt,

$$\alpha(g, m_f) = m_{f \circ \beta(g^{-1}, \cdot)} \quad \text{für alle } g \in G, m_f \in M.$$

Dann ist die von Neumann-Algebra $M \times_\alpha G$ ein Faktor vom Typ I oder vom Typ II. (Statt $M \times_\alpha G$ schreibt man auch $L_\infty(\mathbb{C})(\Omega, \Sigma, \mu) \rtimes_\alpha G$.) Genauer gilt: $M \times_\alpha G$ ist vom **Typ I**, wenn G diskret ist und $G = \Omega$ gilt; das Zählmaß auf $\mathcal{P}(\Omega)$ ist dann ein Haarmaß.

Zum Beispiel ist $M \times_\alpha G$ vom **Typ I_n**, wenn $G = \Omega = \mathbb{Z}_n$ die zyklische Gruppe der Ordnung n , $n \in \mathbb{N}^\times$, ist, und vom **Typ I_∞**, wenn G und Ω die kompakte, abelsche topologische Gruppe $S_\mathbb{C}$ ist (mit Gruppenoperation die Multiplikation von komplexen Zahlen). Wenn G nicht diskret, aber kompakt ist, ist die von Neumann-Algebra $M \times_\alpha G$ ein Faktor vom **Typ II₁**; dies ist zum Beispiel der Fall, wenn man $\Omega := S_\mathbb{C}$ setzt, als Haarmaß μ das durch 2π dividierte Lebesguemaß des mit $S_\mathbb{C}$ identifizierten Intervalls $(-\pi, \pi]$, also $\mu(\Omega) = 1$, und als G eine abzählbare, unendliche, dichte Untergruppe von Ω wählt, zum Beispiel $G := \{\exp(i2\pi n\omega) : n \in \mathbb{Z}\}$ für eine beliebig, aber fest gewählte irrationale Zahl ω . Die Operation α von G auf M stellt sich dann dar als

$$\begin{aligned} (\alpha(z, m_f)g)(w) &= \left(\left(f \circ \beta\left(\frac{1}{z}, \cdot\right) \right) \cdot g \right)(w) \\ &= f\left(\frac{w}{z}\right) \cdot g(w) \quad \text{für alle } z \in G, w \in \Omega, m_f \in M, g \in H. \end{aligned}$$

Mit den Bezeichnungen von (a) gilt

$$\begin{aligned} (\pi(m_f)\varphi)(z)(w) &= (\alpha\left(\frac{1}{z}, m_f\right)\varphi(z))(w) \\ &= (m_{f \circ \beta(z, \cdot)}(\varphi(z)))(w) \\ &= f(zw) \cdot \varphi(z)(w) \end{aligned}$$

für alle $m_f \in M$, $z \in G$, $w \in \Omega$, $\varphi \in \mathcal{H}$; des Weiteren ist

$$(\lambda(z)\varphi)(w) = \varphi\left(\frac{w}{z}\right) \quad \text{für alle } w, z \in G, \varphi \in \mathcal{H}.$$

Wenn G weder diskret noch kompakt ist (beachte, dass μ genau dann endlich ist, wenn G kompakt ist), ist die von Neumann-Algebra $M \times_\alpha G$ ein Faktor vom **Typ II_∞**; diese Situation liegt zum Beispiel vor, wenn man Ω als die lokal kompakte, abelsche Gruppe $(\mathbb{R}, +)$ wählt, wobei sich als Haarmaß unmittelbar das Lebesguemaß anbietet, und als G die rationalen Zahlen. Wieder mit den Bezeichnungen von (a) gilt dann offensichtlich für alle $z \in \mathbb{Q}$, $m_f \in M$, $\varphi \in \mathcal{H}$

$$\begin{aligned} (\alpha(z, m_f)g)(w) &= f(w - z) \cdot g(w) && \text{für alle } w \in \mathbb{R}, g \in H, \\ (\pi(m_f)\varphi)(z)(w) &= f(w + z) \cdot \varphi(z)(w) && \text{für alle } w \in \mathbb{R}, \text{ und} \\ (\lambda(z)\varphi)(w) &= \varphi(w - z) && \text{für alle } w \in \mathbb{Q}. \end{aligned}$$

(d) Sei Ω die lokal kompakte, abelsche Gruppe $(\mathbb{R}, +)$, versehen mit dem Lebesguemaß auf der Borel- σ -Algebra Σ von $\mathbb{R}$. Sei $p > 1$ eine beliebig, aber fest gewählte rationale Zahl. Sei G die Gruppe $\mathbb{Z} \times \mathbb{Q}$. Setze

$$\beta((n, q), t) := p^n t + q \quad \text{für alle } (n, q) \in \mathbb{Z} \times \mathbb{Q}, t \in \mathbb{R}.$$

Die Abbildung β operiert also von G auf $\mathbb{R}$ per Homöomorphismen. Sei M wieder wie in (c) die von $L_\infty(\mathbb{C})(\Omega, \Sigma, \mu)$ auf $L_2(\mathbb{C})(\Omega, \Sigma, \mu)$ operierende von Neumann-Algebra und sei α ebenfalls wie in (c) definiert. Dann ist die von Neumann-Algebra $M \times_\alpha G$ ein Faktor vom **Typ III**.

5.2.5 Satz (BECERRA GUERRERO und MARTÍN [18](2005), Satz 3.10). *Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Falls X von $\{0\}$ verschieden ist, sind die folgenden Aussagen äquivalent:*

- (α) X ist ein Daugavetraum.
- (β) Die Norm von X ist extrem rau.
- (γ) Die Norm von X ist an keinem Punkt von S_X Fréchet-glatt.

Beweis. $(\alpha) \Rightarrow (\beta)$: Bemerkung 5.1.5.

$(\beta) \Rightarrow (\gamma)$: Klar.

$(\gamma) \Rightarrow (\alpha)$: Falls (γ) gilt, existiert in X gemäß Satz 4.7.82 kein divisionsminimal tripotentes Element von X , so dass X nach Satz 5.2.1 ein Daugavetraum sein muss. $\square$

5.2.6 Bemerkungen (BECERRA GUERRERO und MARTÍN [18](2005), Seiten 326, 327, Bemerkungen 3.11 und 3.12).

(a) Für den Folgenraum ℓ_1 gilt zwar die Aussage (β) (und somit auch (γ) von Satz 5.2.5), er ist aber, wie in 5.1.2(a) bemerkt, kein Daugavetraum.

(b) JOHN und ZIZLER [157](1978), Seite 341, Bemerkung 4, versehen $\ell_1(I)$, I eine nicht abzählbare Menge, mit einer äquivalenten Norm, die nirgendwo glatt ist — für die also insbesondere die Aussage (γ) von Satz 5.2.5 zutrifft —, aber nicht rau ist.

(c) Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Nach BECERRA GUERRERO, LÓPEZ PÉREZ, PERALTA und RODRÍGUEZ PALACIOS [16](2004), Seite 51, Korollar 2.5, gilt: Ist X nicht reflexiv, so ist die Norm des Präduals von X extrem rau.

5.2.7 Satz (BECERRA GUERRERO und MARTÍN [18](2005), Seiten 321, 322, Theorem 3.2). *Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Bezeichne X_* das Prädual von X . Falls X von $\{0\}$ verschieden ist, sind die folgenden Aussagen äquivalent:*

- (a) X ist ein Daugavetraum.
- (b) X_* ist ein Daugavetraum.
- (c) Der Durchmesser von jeder relativ schwach offenen, nicht leeren Teilmenge von B_{X_*} ist gleich 2.
- (d) Der Durchmesser von jeder Scheibe von B_{X_*} ist gleich 2.
- (e) B_{X_*} hat keine stark exponierten Punkte.
- (f) B_{X_*} hat keine Extrempunkte, mit anderen Worten, X hat keine Atome.
- (g) X hat keine minimal tripotenten Elemente.
- (h) X hat keine divisionsminimal tripotenten Elemente.

Beweis. (a) $\Rightarrow$ (b): Dies wurde in 5.1.2(f) bemerkt.

(b) $\Rightarrow$ (c): Bemerkung 5.1.5.

(c) $\Rightarrow$ (d): Klar.

(d) $\Rightarrow$ (e): Klar, da sonst Widerspruch zu der in 2.11.11 erstgenannten äquivalenten Formulierung eines stark exponierten Punktes.

(e) $\Rightarrow$ (f): Sonst Widerspruch zu Korollar 4.7.83.

(f) $\Leftrightarrow$ (g): Dies wurde nach (4.26) in 4.7.73 bemerkt.

(g) $\Leftrightarrow$ (h): Aussage (4.25) in 4.7.66.

Für die letzte zum Ringschluss benötigte Implikation hat man nun zwei Möglichkeiten: Die Äquivalenz (h) $\Leftrightarrow$ (a) gemäß Satz 5.2.1 oder die Implikation (f) $\Rightarrow$ (a), die wie folgt begründet werden kann: X_* ist nach 4.7.57 L -eingebettet. Gilt (f), so ist X_* nach dem Satz von Krein-Milman 2.11.11 nicht reflexiv und mit Korollar 5.1.10 folgt (a). $\square$

Wieder mit 4.7.69 hat man somit unmittelbar das

5.2.8 Korollar. *Das Prädual einer von $\{0\}$ verschiedenen W^* -Algebra A über $\mathbb{K}$ ist genau dann ein Daugavetraum, wenn A nicht atomar ist.*

5.2.9 Korollar (BECERRA GUERRERO und MARTÍN [18](2005), Seite 324, Korollar 3.5). *Sei X JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Dann ist X^* kein Daugavetraum. Falls X ein Bidual eines Raumes ist, ist X kein Daugavetraum.*

Beweis. (Ganz nach ebd.) X^{**} ist nach 4.7.28 bzw. nach 4.7.48 ein JBW^* -Tripel über $\mathbb{C}$ bzw. ein *JBW -Tripel. Nach dem Satz von Alaoglu-Bourbaki und dem Satz von Krein-Milman 2.11.11 besitzt $B_{X^*} = B_{X^{**}}$ mindestens einen Extrempunkt. Ist $X = Y^{**}$, so hat aus dem gleichen Grund $B_{X^*} = B_{Y^*}$ mindestens einen Extrempunkt. $\square$

5.2.10 Bemerkungen (BECERRA GUERRERO und MARTÍN [18](2005), Seite 324, Bemerkung 3.6 und Seite 327, Bemerkung 3.12). **(a)** In 5.1.2(a) wurde festgestellt, dass der Folgenraum c_0 kein Daugavetraum ist. Nichtsdestotrotz gelten die Aussagen (c) (siehe BECERRA GUERRERO, LÓPEZ PÉREZ und RODRÍGUEZ PALACIOS [17](2003), Seite 754, Lemma 2.2) und (f) des Satzes 5.2.7, wenn man dort für X_* den Raum c_0 einsetzt.

(b) Es existiert ein Banachraum Z , dessen abgeschlossene Einheitskugel beulbar ist — also Z in der Rolle von X_* nicht die Aussage (d) des Satzes 5.2.7 erfüllt —, aber keinen Extrempunkt aufweist. (In ebd. wird hierzu EDELSTEIN [80](1973), Satz 1, angeführt.)

(c) Obgleich jede Scheibe der abgeschlossenen Einheitskugel des Folgenraumes ℓ_∞ Durchmesser 2 hat (insbesondere also B_{ℓ_∞} keinen stark exponierten Punkt hat), besitzt B_{ℓ_∞} , da ℓ_∞ ein Dualraum ist, Extrempunkte.

(d) In Hinblick auf die in Satz 5.2.7(c) formulierte Eigenschaft des Prädualballs sei folgendes angemerkt: Sei X ein JB^* -Tripel über $\mathbb{C}$ oder ein *JB -Tripel. Betrachte die Eigenschaft: Der Durchmesser jeder relativ schwach offenen Teilmenge von B_X ist gleich 2. In BECERRA GUERRERO, LÓPEZ PÉREZ, PERALTA und RODRÍGUEZ PALACIOS [16](2004), Seite 49, Theorem 2.3, wurde gezeigt, dass X diese Eigenschaft aufweist, wenn X nicht reflexiv ist.

5.2.11 Satz. *Sei X ein JBW^* -Tripel über $\mathbb{C}$ oder ein *JBW -Tripel. Dann ist der nicht atomare Teil von X ein Daugavetraum. Somit ist auch der nicht atomare Teil von X_* ein Daugavetraum.*

Beweis. Da der nicht atomare Teil von X ein JBW^* -Tripel über $\mathbb{C}$ bzw. ein *JBW -Tripel ist, das keine minimal tripotenten Elemente enthält, ist er nach Satz 5.2.7 ein Daugavetraum. Der nicht atomare Teil von X_* ist das Prädual des nicht atomaren Teils von X . Nun siehe 5.1.2(f). $\square$

Literatur

- [1] ABRAMOVICH, YURI A.: *New classes of spaces on which compact operators satisfy the Daugavet equation*. J. Operator Theory, 25:331–345, 1991. [5.1.2](#)
- [2] ABRAMOVICH, YURI A. und CHARALAMBOS D. ALIPRANTIS: *An invitation to operator theory / Problems in operator theory*, Band 50 der Reihe *Graduate studies in mathematics*. American Mathematical Society, Providence, Rhode Island, 2002. [2.4.19](#), [5.1.2](#), [5.1.4](#), [5.1.5](#)
- [3] ABRAMOVICH, YURI A., CHARALAMBOS D. ALIPRANTIS und O. BURKINSHAW: *The Daugavet equation in uniformly convex Banach spaces*. Journal of Functional Analysis, 97:215–230, 1991. [5.1.2](#)
- [4] ADÁMEK, JIŘÍ, HORST HERRLICH und GEORGE E. STRECKER: *Abstract and concret categories. The joy of cats*. Universität Bremen, 2004. Newest online version: <http://katmat.math.uni-bremen.de/acc>. Download vom 19.9.2011. [1.1.1](#)
- [5] AKEMANN, CHARLES und NIK WEAVER: *Geometric characterizations of some classes of operators in C^* -algebras and von Neumann algebras*. Proc. Amer. Math. Soc., 130(10):3033–3037, 2002. [3.5.33](#), [3.5.36](#), [3.5.43](#)
- [6] ALBERT, ABRAHAM ADRIAN: *Structure of algebras*, Band XXIV der Reihe *American Mathematical Society, Colloquium Publications*. American Mathematical Society, Providence, Rhode Island, 1961. Vorwort von 1939; revised printing. [2.1.33](#)
- [7] ALFSEN, ERIK M. und EDWARD G. EFFROS: *Structure in real Banach spaces. Part I and II*. Annals of Mathematics. Second Series, 96:98–173, 1972. [2.5.17](#), [2.5.17](#)
- [8] ALFSEN, ERIK MAGNUS und FREDERIC W. SHULTZ: *State spaces of operator algebras. Basic theory, orientations, and C^* -products*. Birkhäuser Boston, Basel, Berlin, 2001. [2.4.22](#), [3.4.22](#), [3.5.35](#), [3.7.39](#), [3.7.45](#)
- [9] ARAZY, JONATHAN: *An application of infinite dimensional holomorphy to the geometry of Banach spaces*. In: *Geometrical aspects of Functional Analysis, edited by Joram Lindenstrauss and V. Milman. Israel seminar 1985–86*, Band 1267 der Reihe *Lecture Notes in Mathematics*. Springer, 1987. [4.5.18](#), [4.5.31](#), [4.7.24](#)
- [10] ARTIN, MICHAEL: *Algebra*. Birkhäuser Verlag, Basel, Boston, Berlin, 1993. Aus dem Englischen übersetzt von Annette A’Campo. [2.1.22](#), [2.1.43](#)
- [11] AUBIN, JEAN-PIERE und HÉLÈNE FRANKOWSKA: *Set-valued analysis*. Birkhäuser Boston, Basel, Berlin, 1990. [1.1.6](#)
- [12] BARTON, THOMAS JOSEPH und RICHARD M. TIMONEY: *Weak*-continuity of Jordan triple products and its applications*. Math. Scand., 59:177–191, 1986. [4.7.28](#), [4.7.34](#), [4.7.57](#)
- [13] BARTON, THOMAS JOSEPH (TOM) und GILLES GODEFROY: *Remarks on the predual of a JB^* -triple*. Journal of the London Mathematical Society, Second Series, 34:300–304, 1986. [4.7.80](#)
- [14] BARTSCH, RENÉ: *Allgemeine Topologie I*. Siehe <http://math.marvinus.net/>, 13.04.2005, Download vom 24.6.2010. [1.2.28](#)
- [15] BATTAGLIA, M.: *Order theoretic type decomposition of JBW^* -triples*. The Quarterly Journal of Mathematics, Oxford Second Series, 42(166):129–147, 1991. [4.7.66](#)

- [16] BECERRA GUERRERO, JULIO, GINÉS LÓPEZ PÉREZ, ANTONIO M. PERALTA und ÁNGEL RODRÍGUEZ PALACIOS: *Relatively weakly open sets in closed balls of Banach spaces, and real JB^* -triples of finite rank*. Mathematische Annalen, 330:45–58, 2004. [4.7.56](#), [4.7.57](#), [4.7.66](#), [4.7.66](#), [5.2.6](#), [5.2.10](#)
- [17] BECERRA GUERRERO, JULIO, GINÉS LÓPEZ PÉREZ und ÁNGEL RODRÍGUEZ PALACIOS: *Relatively weakly open sets in closed balls of C^* -algebras*. Journal of the London Mathematical Society, Second Series, 68:753–761, 2003. [3.5.56](#), [4.1.11](#), [5.2.10](#)
- [18] BECERRA GUERRERO, JULIO und MIGUEL MARTÍN: *The Daugavet property of C^* -algebras, JB^* -triples, and of their isometric preduals*. J. Funct. Anal., 224(2):316–337, 2005. [4.1.11](#), [4.7.69](#), [4.7.82](#), [4.7.84](#), [5.1.9](#), [5.1.10](#), [5.2.1](#), [5.2.5](#), [5.2.6](#), [5.2.7](#), [5.2.9](#), [5.2.10](#)
- [19] BECERRA GUERRERO, JULIO und A. PERALTA: *Subdifferentiability of the norm and the Banach-Stone theorem for real and complex JB^* -triples*. manuscripta mathematica, 114:503–516, 2004. [4.7.81](#)
- [20] BECERRA GUERRERO, JULIO und ÁNGEL RODRÍGUEZ PALACIOS: *Strong subdifferentiability of the norm on JB^* -triples*. The Quarterly Journal of Mathematics, 54(4):381–390, 2003. [4.7.81](#)
- [21] BEHREND, EHRHARD: *M-Structure and the Banach-Stone Theorem*, Band 736 der Reihe *Lecture Notes in Mathematics*. Springer-Verlag Berlin, Heidelberg, New York, 1979. [2.5.17](#), [2.5.17](#)
- [22] BEHRENS, ERNST-AUGUST: *Nichtassoziative Ringe*. Math. Annalen, 127:441–452, 1954. [2.1.42](#)
- [23] BEHRENS, ERNST-AUGUST: *Ringtheorie*. Bibliographisches Institut AG, Zürich, 1975. Erweiterte und korrigierte Ausgabe des zuerst 1972 in New York erschienenen Buches „Ring Theory“. [2.1.34](#), [2.1.34](#), [2.1.43](#), [2.1.59](#)
- [24] BILIK, DIMITRY, VLADIMIR KADETS, ROMAN SHVIDKOY und DIRK WERNER: *Narrow operators and the Daugavet property for ultraproducts*. Positivity, 9:45–62, 2005. [5.1.6](#)
- [25] BIRKHOFF, GARRETT: *Lattice Theory*, Band XXV der Reihe *American Mathematical Society Colloquium Publications*. American Mathematical Society, Providence, Rhode Island, 3. Auflage, 1973. [1.1.13](#)
- [26] BLACKADAR, BRUCE: *Operator algebras: Theory of C^* -algebras and von Neumann algebras*. Encyclopaedia of Mathematical Sciences, Volume 122. Operator Algebras and Non-Commutative Geometry III. Springer Berlin, Heidelberg, 2006. [3.5.48](#)
- [27] BLYTH, THOMAS SCOTT und MELVIN F. JANOWITZ: *Residuation theory*, Band 102 der Reihe *International Series of Monographs in Pure and Applied Mathematics*. Pergamon Press, Oxford, 1972. [1.1.18](#)
- [28] BONSALE, FRANK FEATHERSTONE: *A survey of Banach algebra theory*. Bull. London Math. Soc., 2:257–274, 1970. [3.5.32](#)
- [29] BONSALE, FRANK FEATHERSTONE und JOHN DUNCAN: *Numerical ranges of operators on normed spaces and of elements of normed algebras*, Band 2 der Reihe *London Mathematical Society Lecture Note Series*. Cambridge at the University Press, 1971. [2.9.27](#), [2.11.15](#), [2.11.17](#), [2.11.19](#), [2.11.21](#), [2.11.23](#), [2.11.27](#), [2.11.31](#), [2.12.8](#), [2.12.9](#), [2.12.10](#), [3.1.8](#), [3.2.22](#), [3.4.9](#)
- [30] BONSALE, FRANK FEATHERSTONE und JOHN DUNCAN: *Complete normed algebras*, Band 80 der Reihe *Ergebnisse der Mathematik und ihrer Grenzgebiete*. Springer-Verlag, 1973. [2.12.12](#), [2.12.15](#), [3.1.2](#), [3.1.7](#), [3.2.4](#), [3.2.17](#), [3.2.18](#), [3.2.19](#), [3.3.2](#), [3.3.3](#)
- [31] BONSALE, FRANK FEATHERSTONE und JOHN DUNCAN: *Numerical ranges II*, Band 10 der Reihe *London Mathematical Society Lecture Note Series*. Cambridge at the University Press, 1973. [3.5.37](#)
- [32] BOURBAKI, NICOLAS: *Elements of Mathematics: General Topology, Part 1*. Hermann, Publishers in Arts and Science, Paris, France; Addison-Wesley Publishing Company, Reading, Massachusetts, 1966. [1.1.16](#), [1.2.1](#), [1.2.31](#), [2.4.1](#), [2.4.7](#)

- [33] BOURBAKI, NICOLAS: *Variétés différentielles et analytiques. Fascicule de résultats; paragraphes 1 à 7*, Band 33 der Reihe *Éléments de mathématique*. Hermann, Paris, 2., korr. Auflage, 1971. [4.4.3](#)
- [34] BOURBAKI, NICOLAS: *Elements of Mathematics: Algebra, Part 1, Chapters 1-3*. Hermann, Publishers in Arts and Science, Paris, France; Addison-Wesley Publishing Company, Reading, Massachusetts, 1974 (1943, 1947, 1948, 1971). [2.1.1](#), [2.1.9](#), [2.1.34](#), [2.9.1](#), [2.9.2](#), [2.9.4](#), [3.4.1](#)
- [35] BOURBAKI, NICOLAS: *Elements of Mathematics: Topological vector spaces, Chapters 1-5 (Espaces vectoriels topologiques, Paris 1981)*. Springer-Verlag Berlin, Heidelberg, New York, London, Paris, Tokyo, 1987. [2.2.6](#), [2.2.8](#), [2.2.11](#), [2.2.25](#), [2.2.31](#), [2.2.34](#), [2.2.37](#), [2.4.8](#), [2.4.16](#), [2.4.29](#), [2.4.30](#), [2.4.37](#), [2.4.39](#), [2.5.5](#), [2.5.14](#), [2.6.3](#), [2.6.4](#), [2.6.8](#), [2.6.12](#)
- [36] BOURGIN, RICHARD D.: *Geometric aspects of convex sets with the Radon-Nikodým property*, Band 993 der Reihe *Lecture Notes in Mathematics*. Springer-Verlag, 1983. [2.6.13](#), [2.6.14](#), [2.6.17](#), [2.10.12](#), [2.11.11](#), [4.1.4](#), [4.1.8](#)
- [37] BRATTELI, OLA und DEREK W. ROBINSON: *Operator algebras and quantum statistical mechanics 1; C^* - and W^* -Algebras, symmetry groups, decomposition of states*. Springer-Verlag New York, Berlin, Heidelberg, London, Paris, Tokyo, 2. Auflage, 1987. [3.6.8](#), [3.6.14](#), [3.6.16](#), [3.7.39](#)
- [38] BRAUN, HEL und MAX KOECHER: *Jordan-Algebren*. Springer, 1966. [2.3.16](#)
- [39] BROWN, A. L.: *On the canonical projection of the third dual of a Banach space onto the first dual*. Bulletin of the Australian Mathematical Society, 15:351–354, 1976. [2.5.16](#)
- [40] BUNCE, L. J., C. H. CHU und B. ZALAR: *Structure spaces and decomposition in JB^* -triples*. Math. Scand., 86:17–35, 2000. [4.7.34](#)
- [41] BURGOS, MARÍA, FRANCISCO JOSÉ FERNÁNDEZ POLO, JORGE J. GARCÉS, JUAN MARTÍNEZ MORENO und ANTONIO M. PERALTA: *Orthogonality preservers in C^* -algebras and JB^* -triples*. Journal of Mathematical Analysis and Applications, 348:220–233, 2008. [4.7.64](#), [4.7.72](#)
- [42] CARADUS, SELWYN ROSS, WILLIAM ELMER PFAFFENBERGER und BERTRAM YOOD: *Calkin Algebras and Algebras of Operators on Banach Spaces*, Band 9 der Reihe *Lecture Notes in Pure and Applied Mathematics*. Marcel Dekker, Inc., New York, 1974. [3.7.59](#)
- [43] CAROTHERS, NEAL L.: *A short course on Banach space theory*, Band 64 der Reihe *London Mathematical Society, Student Texts*. Cambridge University Press, 2005. [1.2.22](#)
- [44] CARTAN, ÉLIE: *Sur les domaines bornés homogènes de l'espace de n variables complexes*. Abhandlungen aus dem mathematischen Seminar der Universität Hamburg, 11(1):116–162, 1935. [4.5.24](#)
- [45] CHOQUET, GUSTAVE: *Lectures on Analysis. Volume I: Integration and Topological Vector Spaces*. W. A. Benjamin, Inc., Reading, Massachusetts, 1969. Edited by Jerry Marsden, Tim Lance, and Steve Gelbart. [1.2.9](#), [2.9.32](#)
- [46] CHOQUET, GUSTAVE: *Lectures on Analysis. Volume II: Representation Theory*. W. A. Benjamin, Inc., Reading, Massachusetts, 1969. Edited by Jerry Marsden, Tim Lance, and Steve Gelbart. [2.4.37](#), [2.4.38](#), [2.6.6](#)
- [47] CHU, C.-H. und JOSÉ M. ISIDRO: *Manifolds of tripotents in JB^* -triples*. Mathematische Zeitschrift, 233:741–754, 2000. [4.7.63](#), [4.7.65](#)
- [48] CHU, CHO-HO, TRUONG DANG, BERNARD RUSSO und BELISARIO VENTURA: *Surjective Isometries of real C^* -algebras*. J. London Math. Soc. (2), 47:97–118, 1993. [3.5.10](#), [3.7.22](#), [3.7.25](#), [3.7.26](#)
- [49] CHU, CHO-HO und BRUNO IOCHUM: *On the Radon-Nikodým property in Jordan triples*. Proceedings of the American Mathematical Society, 99(3), March 1967. [4.7.79](#)
- [50] CHU, CHO-HO und PAULINE MELLON: *Jordan structures in Banach spaces and symmetric manifolds*. Expositiones Mathematicae, 16:157–180, 1998. [4.7.24](#)

- [51] CONNES, ALAIN: *Noncommutative Geometry*. Academic Press, Inc.; A division of Harcourt Brace & Company 525 B Street, Suite 1900, San Diego, California, 1994. English translation by Sterling Berberian from *Géométrie non commutative*, InterEditions, Paris, 1990. [3.7.59](#)
- [52] CONWAY, JOHN B.: *A course in Functional Analysis*, Band 96 der Reihe *Graduate texts in Mathematics*. Springer-Verlag, New York, 2. Auflage, 1990. [2.4.19](#), [3.4.11](#)
- [53] CONWAY, JOHN B.: *A course in operator theory*, Band 21 der Reihe *Graduate Studies in Mathematics*. American Mathematical Society, Providence, Rhode Island, 1999. [3.5.23](#), [3.5.53](#), [3.6.5](#), [3.6.14](#), [3.7.15](#), [3.7.20](#), [3.7.31](#), [3.7.45](#), [3.7.51](#), [3.7.53](#)
- [54] COX, DAVID, JOHN LITTLE und DONAL O'SHEA: *Ideals, Varieties, and Algorithms. An Introduction to Computational Algebraic Geometry and Commutative Algebra*. Undergraduate texts in mathematics. Springer-Verlag New York, Inc., 1997. 2nd edition. [2.1.57](#), [4.3.5](#)
- [55] CUDIA, DENNIS F.: *The geometry of Banach spaces. Smoothness*. Trans. Amer. Math. Soc., 110:284–314, 1964. [4.1.12](#)
- [56] DALES, H. GARTH: *Banach algebras and automatic continuity*, Band 24 der Reihe *London Mathematical Society Monographs New Series*. Clarendon Press, Oxford University Press, Great Britain, 2000. [2.1.13](#), [2.1.43](#), [2.4.10](#), [2.4.17](#), [2.7.2](#), [3.2.27](#)
- [57] DANG, T.: *Real isometries between JB^* -triples*. Proceedings of the American Mathematical Society, 114(4):971–980, April 1992. [4.7.19](#)
- [58] DANG, TRUONG C. und BERNARD RUSSO: *Real Banach Jordan triples*. Proceedings of the American Mathematical Society, 122(1), September 1994. [4.7.35](#), [4.7.37](#), [4.7.39](#), [4.7.42](#)
- [59] DAUGAVET, I. K.: *On a property of completely continuous operators in the space C* . Uspekhi Mat. Nauk, 18(5):157–158, 1963. In Russisch. [5.1.2](#)
- [60] DAVIDSON, KENNETH R.: *C^* -Algebras by example*. Fields Institute Monographs. American Mathematical Society, Providence, Rhode Island, 1996. [3.5.14](#)
- [61] DAY, MAHLON M.: *Normed linear spaces*, Band 21 der Reihe *Ergebnisse der Mathematik und ihrer Grenzgebiete*. Springer, 3. Auflage, 1973. [2.4.7](#), [2.6.14](#), [2.11.11](#)
- [62] DEVILLE, ROBERT, GILLES GODEFROY und VACLAV ZIZLER: *Smoothness and renormings in Banach spaces*, Band 64 der Reihe *Pitman monographs and surveys in Pure and Applied Mathematics*. Longman Scientific & Technical, Longman Group UK Limited, 1993. [4.1.11](#)
- [63] DICKSON, LEONARD EUGENE: *Algebras and their arithmetics*. The University of Chicago Press, Chicago, 1923. Unabridged and unaltered republication, ca. 1960, by Dover Publications, Inc., New York, of the first edition 1923. [2.1.33](#), [64](#)
- [64] DICKSON, LEONARD EUGENE und ANDREAS SPEISER: *Algebren und ihre Zahlentheorie. Mit einem Kapitel von Andreas Speiser (Zürich) über Idealtheorie*. Orell Füssli Verlag Zürich und Leipzig, 1927. Vom Verfasser vollständig überarbeitete und überwachte Übersetzung von [63]. [2.1.47](#)
- [65] DIESTEL, JOSEPH: *Geometry of Banach spaces - Selected topics*, Band 485 der Reihe *Lecture Notes in Mathematics*. Springer, 1975. [2.4.22](#)
- [66] DIESTEL, JOSEPH und JOHN JERRY UHL: *Vector measures*. Nummer 15 in *Mathematical surveys*. American Mathematical Society, Providence, Rhode Island, 1977. [2.6.17](#)
- [67] DIEUDONNÉ, JEAN: *Complex structures on real Banach spaces*. Proceedings of the American Mathematical Society, 3(1):162–164, February 1952. [2.9.19](#)
- [68] DIEUDONNÉ, JEAN ALEXANDRE EUGÈNE: *Grundzüge der modernen Analysis. Foundations of Modern Analysis. Treatise on Analysis. 9 Volumes. Aus der Reihe: Logik und Grundlagen der Mathematik*. vieweg Verlag, ab 1960. [3.2.24](#), [3.4.30](#), [4.1.5](#), [4.6.13](#)
- [69] DINEEN, SEÁN: *Complex analysis in locally convex spaces*, Band 57 (83) der Reihe *North-Holland Mathematics Studies (Notas de matematica)*. North-Holland Publishing Company - Amsterdam - New York - Oxford, 1981. [2.2.44](#), [4.5.23](#)

- [70] DINEEN, SEÁN: *Complete holomorphic vector fields on the second dual of a Banach space*. Math. Scand., 59:131–142, 1986. [4.7.24](#), [4.7.28](#)
- [71] DINEEN, SEÁN: *The second dual of a JB^* -triple system*, Band 125 der Reihe *North-Holland Mathematics studies*, Seiten 67–69. Elsevier Science Publishers, Amsterdam, Netherlands B. V., 1986. Complex analysis, functional analysis and approximation theory. Proceedings of the conference on complex analysis and approximation theory, Univ. Estadual de Campinas, Brazil, 23/27 July, 1984. Edited by Jorge Mujica. [4.7.2](#), [4.7.28](#)
- [72] DIVINSKY, NATHAN JOSEPH: *Rings and radicals*. Allen & Unwin, 1965. [2.1.19](#), [2.1.34](#), [2.1.43](#), [2.1.47](#), [2.1.48](#), [2.1.50](#), [2.1.50](#), [2.1.53](#), [2.1.54](#), [2.1.55](#), [2.1.56](#), [2.1.58](#), [2.1.59](#)
- [73] DIVINSKY, NATHAN JOSEPH und A. SULÍŃSKI: *Kurosh radicals of rings with operators*. Can. J. of Math., 17:278–280, 1965. [2.1.28](#), [2.1.49](#)
- [74] DIXMIER, JACQUES: *Sur un théorème de Banach*. Duke Math. J., 15:1057–1071, 1948. [2.5.16](#)
- [75] DIXMIER, JACQUES: *C^* -Algebras*. North-Holland Publishing Company, Amsterdam, New York, Oxford, 1977. [3.5.43](#)
- [76] DIXMIER, JACQUES: *Von Neumann algebras*. North-Holland Publishing Company, Amsterdam, New York, Oxford, 1981. [3.4.22](#), [3.4.28](#), [3.6.3](#), [3.6.5](#), [3.6.7](#), [3.6.12](#), [3.6.16](#), [3.6.18](#), [3.7.6](#), [3.7.59](#)
- [77] DORAN, ROBERT STUART und VICTOR ALLEN BELFI: *Characterizations of C^* -algebras. The Gelfand-Naimark Theorems*, Band 101 der Reihe *Monographs and textbooks in pure and applied mathematics*. Marcel Dekker, Inc., New York, Basel, 1986. [2.11.23](#), [2.11.24](#), [2.11.28](#), [3.2.33](#), [3.2.35](#), [3.3.2](#), [3.5.7](#), [3.5.8](#), [3.5.22](#), [3.5.22](#), [3.5.24](#), [3.5.31](#), [3.5.37](#)
- [78] DUNFORD, NELSON und JACOB T. SCHWARTZ: *Linear operators. Part I: General Theory*, Band VII der Reihe *Pure and applied mathematics*. Interscience Publishers, Inc., New York, 1970. [1.1.7](#), [1.1.9](#), [1.2.20](#), [2.1.25](#), [2.5.11](#), [2.10.11](#), [4.1.9](#)
- [79] DUNFORD, NELSON und JACOB T. SCHWARTZ: *Linear Operators. Part II: Spectral Theory, Self Adjoint Operators in Hilbert space*. Interscience Publishers (John Wiley & Sons), 1973. [3.6.1](#)
- [80] EDELSTEIN, MICHAEL: *Concerning dentability*. Pacific Journal of Mathematics, 46(1):111–114, 1973. [5.2.10](#)
- [81] EDWARDS, C. M. und GOTTFRIED TRAUGOTT RÜTTIMANN: *On the facial structure of the unit balls in a JBW^* -triple and its predual*. Journal of the London Mathematical Society, Second Series, 38(2):317–332, 1988. [4.7.73](#), [4.7.74](#)
- [82] EDWARDS, C. MARTIN: *On Jordan W^* -algebras*. Bulletin des Sciences Mathématiques, Deuxième Série, 104:393–403, 1980. [4.7.14](#), [4.7.15](#)
- [83] EDWARDS, C. MARTIN, FRANCISCO JOSÉ FERNÁNDEZ POLO, CHRISTOPHER S. HOSKIN und ANTONIO M. PERALTA: *On the facial structure of the unit ball in a JB^* -triple*. Journal für die reine und angewandte Mathematik, 641:123–144, 2010. [4.7.15](#), [4.7.71](#)
- [84] EDWARDS, C. MARTIN, KEVIN MCCRIMMON und GOTTFRIED TRAUGOTT RÜTTIMANN: *The range of a structural projection*. Journal of Functional Analysis, 139:196–224, 1996. [4.7.61](#)
- [85] EDWARDS, C. MARTIN und GOTTFRIED TRAUGOTT RÜTTIMANN: *A characterization of inner ideals in JB^* -triples*. Proc. Amer. Math. Soc., 116:1049–1057, 1992. [4.7.12](#)
- [86] EDWARDS, C. MARTIN und GOTTFRIED TRAUGOTT RÜTTIMANN: *The facial and inner ideal structure of a real JBW^* -triple*. Math. Nachr., 222:159–184, 2001. [2.11.8](#), [4.7.73](#), [4.7.73](#)
- [87] EDWARDS, C. MARTIN und GOTTFRIED TRAUGOTT RÜTTIMANN: *Involutive and Peirce gradings in JBW^* -triples*. Communications in Algebra, 31:2819–2848, 2003. [4.7.59](#)

- [88] EDWARDS, ROBERT EDMUND: *Functional Analysis. Theory and Applications*. Holt, Rinehart and Winston, Inc., New York, Chicago, San Francisco, Toronto, London, 1965. [2.2.11](#), [2.4.15](#), [2.4.16](#), [2.5.11](#), [3.4.7](#)
- [89] EL AMIN, KAIDI, ANTONIO MORALES CAMPOY und ÁNGEL RODRÍGUEZ PALACIOS: *A holomorphic characterization of C^* - and JB^* -algebras*. manuscripta mathematica, 104:467–478, 2001. [4.7.3](#)
- [90] ENDERTON, HERBERT B.: *Elements of set theory*. Academic Press, New York, San Francisco, London, 1977. [1.1.1](#), [1.1.1](#)
- [91] ERNÉ, M., J. KOSLOWSKI, A. MELTON und G. E. STRECKER: *A Primer on Galois Connections*. Annals of the New York Academy of Sciences, 704:103–125, 1993. [1.1.18](#)
- [92] FARAUT, JACQUES, SOJI KANEYUKI, ADAM KORÁNYI, QI-KENG LU und GUY ROOS: *Analysis and geometry on complex homogeneous domains*, Band 185 der Reihe *Progress in Mathematics*. Birkhäuser Boston, Basel, Berlin, 2000. [4.6.6](#), [4.7.59](#)
- [93] FERNÁNDEZ POLO, FRANCISCO JOSÉ, JUAN MARTÍNEZ und ANTONIO M. PERALTA: *Surjective isometries between real JB^* -triples*. Math. Proc. Camb. Phil. Soc., 137:709–723, 2004. [4.7.79](#)
- [94] FERNÁNDEZ POLO, FRANCISCO JOSÉ, JUAN MARTÍNEZ MORENO und ANTONIO M. PERALTA: *Geometric characterization of tripotents in real and complex JB^* -triples*. Journal of Mathematical Analysis and Applications, 295(2):435–443, 15 July 2004. [4.7.17](#), [4.7.62](#)
- [95] FERNÁNDEZ POLO, FRANCISCO JOSÉ und ANTONIO M. PERALTA: *On the facial structure of the unit ball in the dual space of a JB^* -triple*. Mathematische Annalen, 348:1019–1032, 2010. [4.7.71](#)
- [96] FILLMORE, PETER A.: *A user's guide to operator algebras*. Canadian Mathematical Society Series of Monographs and Advanced Texts. John Wiley & Sons, Inc., New York, Chichester, Bristane, Toronto, Singapore, 1996. [3.5.53](#), [5.2.4](#)
- [97] FRAENKEL, ABRAHAM A., YEHOSHUA BAR-HILLEL und AZRIEL LEVY: *Foundations of Set Theory*. North-Holland Publishing Company, 1973. Second revised edition. [1.1.1](#)
- [98] FRANCHETTI, CARLO und RAFAEL PAYÁ: *Banach spaces with strongly subdifferentiable norm*. Bolletino della Unione Matematica Italiana, Serie VII, VII-B:45–70, 1993. [4.1.8](#)
- [99] FREMLIN, D. H.: *Measure Theory*. Torres Fremlin, 25 Ireton Road, Calchester CO3 3AT, England, 2000-2008. [2.1.24](#), [2.2.7](#), [2.4.17](#), [2.6.16](#), [3.5.48](#), [5.2.4](#)
- [100] FRIEDMAN, YAAKOV und BERNARD RUSSO: *Structure of the predual of a JBW^* -triple*. Journal für die reine und angewandte Mathematik, 356:67–89, 1985. [4.7.59](#), [4.7.60](#), [4.7.66](#), [4.7.74](#), [4.7.75](#), [4.7.76](#), [4.7.77](#)
- [101] FRIEDMAN, YAAKOV und BERNARD RUSSO: *The Gelfand-Naimark theorem for JB^* -triples*. Duke Mathematical Journal, 53(1):139–148, March 1986. [4.6.9](#), [4.7.17](#), [4.7.79](#), [4.7.84](#)
- [102] FRIEDMAN, YAAKOV und BERNARD RUSSO: *Conditional expectation and bicontractive projections on Jordan C^* -algebras and their generalizations*. Mathematische Zeitschrift, 194:227–236, 1987. [3.1.11](#)
- [103] FRIEDRICHS DORF, ULF und ALEXANDER PRESTEL: *Mengenlehre für den Mathematiker*. Vieweg, Braunschweig, 1985. [1.1.1](#)
- [104] FUKAMIYA, M., Y. MISONOU und Z. TAKEDA: *On order and commutativity of B^* -algebras*. Tohoku Math. J. (2), 6(1):89–93, 1954. [3.5.48](#)
- [105] GÄHLER, SIEGFRIED und WERNER GÄHLER: *Beiträge zu einer streng komplexen Analysis*. In: NAAS, JOSEF (Herausgeber): *Beiträge zur komplexen Analysis und deren Anwendungen in der Differentialgeometrie*, Schriftenr. Zentralinst. Math. Mech. Akad. Wiss., DDR, Seiten 63–92. Akademie-Verlag, Berlin, DDR, 1974. [2.9.20](#)
- [106] GELFAND, ISRAEL MOISSEJEWITSCH: *Normierte Ringe*. Matematicheskii Sbornik, Recueil Mathématique, 9(51)(1):3–24, 1941. [2.9.11](#), [2.12.8](#)

- [107] GERHARDT, CLAUS: *Analysis II*. International Series in Analysis. Graduate Series in Analysis. IP International Press, Somerville, Massachusetts, 2006. Ruprecht-Karls-Universität, Institut für Angewandte Mathematik, Heidelberg. [4.1.5](#)
- [108] GILES, J. R., D. A. GREGORY und BRILLY SIMS: *Geometrical implications of upper semi-continuity of the duality mapping on a Banach space*. Pacific Journal of Mathematics, 79(1):99–109, 1978. [2.4.40](#), [2.11.13](#), [4.1.12](#)
- [109] GILES, JOHN ROBILLIARD: *Convex analysis with application in the differentiation of convex functions*, Band 58 der Reihe *Research Notes in Mathematics*. Pitman Books Limited, 1982. [2.2.44](#), [2.6.13](#), [2.6.14](#), [2.11.11](#), [2.11.11](#), [4.1.5](#), [4.1.8](#)
- [110] GLUBRECHT, JÜRGEN-MICHAEL, ARNOLD OBERSCHELP und GÜNTER TODT: *Klassenlogik*. Bibliographisches Institut Mannheim, Wien, Zürich, 1983. [1.1.1](#)
- [111] GODEFROY, GILLES: *Espaces de Banach: existence et unicité de certains préduaux*. Annales de l'institut Fourier, 28(3):87–105, 1978. [2.10.7](#), [2.10.8](#)
- [112] GODEFROY, GILLES: *Étude des projections de norme 1 de E'' sur E . Unicité de certains préduaux. Applications*. Annales de l'institut Fourier, 29(4):53–70, 1979. [2.10.7](#)
- [113] GODEFROY, GILLES: *Parties admissibles d'un espace de Banach. Applications*. Annales scientifiques de l'É.N.S. 4e série, 16(1):109–122, 1983. [2.10.8](#)
- [114] GODEFROY, GILLES: *Quelques remarques sur l'unicité des préduaux*. The Quarterly Journal of Mathematics. Oxford Second Series, 35:147–152, 1984. [2.10.8](#)
- [115] GOLDBERGER, JACOB K. und GERTRUDE EHRLICH: *Algebra*. The Macmillan Company, 1970. [2.1.1](#), [2.1.39](#), [2.9.25](#), [4.3.9](#)
- [116] GOODEARL, KENNETH RALPH: *Ring Theory. Nonsingular rings and modules*, Band 33 der Reihe *Monographs and textbooks in Pure and Applied Mathematics*. Marcel Dekker, Inc., New York, Basel, 1976. [2.1.34](#)
- [117] GOODEARL, KENNETH RALPH: *Notes on real and complex C^* -Algebras*. Shiva Publishing Limited, 1982. [2.12.17](#), [3.4.10](#), [3.5.11](#), [3.5.13](#), [3.5.17](#), [3.5.18](#), [3.5.39](#)
- [118] GOODEARL, KENNETH RALPH und ROBERT B. WARFIELD, JR.: *An introduction to noncommutative Noetherian rings*, Band 61 der Reihe *London Mathematical Society, Student Texts*. Cambridge University Press, 2. Auflage, 2004. [2.1.50](#)
- [119] GÖPFERT, ALFRED, THOMAS RIEDRICH und CHRISTIANE TAMMER: *Angewandte Funktionalanalysis. Motivationen und Methoden für Mathematiker und Wirtschaftswissenschaftler*. Studienbücher Wirtschaftsmathematik. Vieweg+Teubner, Wiesbaden, 2009. 1. Auflage. [2.11.4](#)
- [120] GRAY, MARY: *A radical approach to algebra*. Addison-Wesley Publishing Company, 1970. [2.1.43](#), [2.1.48](#), [2.1.50](#), [2.1.54](#), [2.1.56](#), [2.3.6](#)
- [121] GREGORY, DAVID A.: *Upper semi-continuity of subdifferential mappings*. Canadian Mathematical Bulletin. Bulletin Canadien de Mathématiques, 23(1):11–19, 1980. [4.1.12](#)
- [122] GROVE, LARRY CHARLES: *Algebra*. Pure and Applied Mathematics. A series of Monographs and Textbooks. Academic Press, New York, 1983. [2.1.11](#), [2.1.34](#)
- [123] HALL, BRIAN C.: *Lie groups, Lie algebras, and representations: an elementary introduction*. Springer New York, Berlin, Heidelberg, 2003. [2.3.10](#)
- [124] HALMOS, PAUL RICHARD: *Naïve Mengenlehre*, Band 6 der Reihe *Moderne Mathematik in elementarer Darstellung*. Vandenhoeck & Ruprecht in Göttingen, 1969. [1.1.1](#)
- [125] HANCHE-OLSEN, HARALD und ERLING STØRMER: *Jordan operator algebras*, Band 21 der Reihe *Monographs and studies in mathematics*. Pitman Publishing, 1984. [4.7.15](#)
- [126] HARMAND, PETER, DIRK WERNER und WEND WERNER: *M-Ideals in Banach spaces and Banach algebras*, Band 1547 der Reihe *Lecture Notes in Mathematics*. Springer-Verlag, Berlin, 1993. [2.5.16](#), [2.5.17](#), [2.5.17](#), [2.5.17](#), [2.10.7](#), [3.7.56](#)

- [127] HARRIS, LAWRENCE A.: *Schwarz's Lemma in normed linear spaces*. Proceedings of the National Academy of Sciences in the United States of America, 62:1014–1017, 1969. [4.5.5](#)
- [128] HARRIS, LAWRENCE A.: *Bounded symmetric homogeneous domains in infinite dimensional spaces*. In: *Proceedings on Infinite dimensional holomorphy, University of Kentucky, Lexington, May 28 - June 1, 1973*. Edited by T. L. Hayden and T.J. Suffridge, Band 364 der Reihe *Lecture Notes in Mathematics*, Seiten 13–40. Springer-Verlag Berlin, Heidelberg, New York, 1974. [4.5.6](#), [4.5.15](#), [4.7.11](#)
- [129] HAZEWINKEL, MICHIEL (Herausgeber): *Encyclopaedia of Mathematics. An updated and annotated translation of the Soviet „Mathematical Encyclopaedia“*. Reidel; Kluwer Academics Publishers, Dordrecht, Boston, Lancaster, Tokyo, 1988–2001. <http://eom.springer.de>. [2.1.28](#)
- [130] HEINRICH, STEFAN: *Ultraproducts in Banach space theory*. Journal für die Reine und Angewandte Mathematik (Crelle Journal), 313:72–104, 1980. [2.4.24](#), [4.7.28](#)
- [131] HELGASON, SIGURDUR: *Differential Geometry and Symmetric Spaces*. Academic Press, 1969. Fourth printing. [4.5.27](#)
- [132] HERRLICH, HORST und GEORGE E. STRECKER: *Category theory*. Allyn and Bacon Inc., Boston, 1973. [1.1.1](#)
- [133] HEUSER, HARRO: *Funktionalanalysis (Theorie und Anwendung)*. Teubner, 1. Auflage: 1975, 2. Auflage: 1986, 4. Auflage: 2006. [2.2.32](#), [2.4.29](#), [2.5.13](#), [2.6.1](#), [2.6.5](#), [2.6.7](#), [2.6.10](#), [2.6.19](#), [2.6.20](#)
- [134] HEUSER, HARRO: *Lehrbuch der Analysis. Teil 2*. B. G. Teubner, Stuttgart, Leipzig, Wiesbaden, 2002. 12. Auflage. [2.2.10](#)
- [135] HOLDGRÜN, H. S.: *Dualitätssätze; Harmonische Analyse, Teil 2. Nach einer Vorlesung aus dem Sommersemester 1983*. Mathematisches Institut der Universität Göttingen, 1983. [3.6.16](#), [3.7.15](#)
- [136] HOLDGRÜN, H. S.: *Klassifizierung von von Neumann-Algebren*. Mathematisches Institut der Universität Göttingen, 1987. [3.7.15](#)
- [137] HOLMES, RICHARD B.: *Geometric functional analysis and its applications*. Springer-Verlag, New York, Heidelberg, Berlin, 1975. [2.2.44](#), [2.4.8](#), [2.5.1](#), [2.5.9](#), [2.5.16](#), [2.10.2](#), [2.10.3](#), [2.10.4](#), [2.10.5](#), [3.6.4](#)
- [138] HOLUB, JAMES R.: *A property of weakly compact operators on $C[0, 1]$* . Proc. Amer. Math. Soc., 97(3):396–398, 1986. [5.1.2](#)
- [139] HORN, GÜNTHER: *Klassifikation der JBW^* -Tripel vom Typ I*. Doktorarbeit, Mathematische Fakultät der Eberhard-Karls-Universität Tübingen, 1984. [4.7.79](#)
- [140] HORN, GÜNTHER: *Characterization of the predual and ideal structure of a JBW^* -triple*. Math. Scand., 61:117–133, 1987. [2.10.8](#), [4.7.9](#), [4.7.28](#), [4.7.34](#)
- [141] HORNFECK, BERNHARD: *Algebra*. Walter de Gruyter, Berlin, New York, 1976. 3., verbesserte Auflage. [2.1.22](#), [2.1.43](#), [4.3.4](#), [4.3.4](#)
- [142] HORVÁTH, JOHN: *Topological vector spaces and distributions. Volume I*. Addison-Wesley Series in Mathematics. Addison-Wesley Publishing Company, Reading, Massachusetts, 1966. [2.4.5](#), [2.4.27](#), [2.6.2](#), [2.6.6](#), [2.6.12](#)
- [143] HUNGERFORD, THOMAS W.: *Algebra*. Springer-Verlag, New York, 1980. [1.1.2](#), [2.1.39](#), [2.9.1](#), [4.3.9](#)
- [144] INGELSTAM, LARS: *A vertex property for Banach algebras with identity*. Math. Scand., 11:22–32, 1962. [2.9.29](#), [2.9.30](#)
- [145] INGELSTAM, LARS: *Real Banach Algebras*. Arkiv för Matematik, 5(16):239–270, 1964. [2.9.11](#), [2.9.15](#), [2.9.16](#), [2.9.17](#), [2.9.19](#), [2.9.29](#), [2.12.4](#), [2.12.11](#), [3.2.8](#), [3.2.31](#), [3.5.11](#), [3.5.12](#), [3.5.16](#), [3.5.22](#)
- [146] ISIDRO, JOSÉ M.: *A glimpse at the theory of Jordan-Banach triple systems*. Revista Matematica Complutense, supplementary, 2:145–156, 1989. [4.5.35](#), [4.7.3](#), [4.7.11](#)

- [147] ISIDRO, JOSÉ M., WILHELM KAUP und ÁNGEL RODRÍGUEZ PALACIOS: *On real forms of JB^* -triples*. Manuscripta mathematica, 86:311–335, 1995. [3.2.18](#), [4.7.14](#), [4.7.16](#), [4.7.39](#), [4.7.40](#), [4.7.43](#), [4.7.44](#), [4.7.45](#), [4.7.46](#), [4.7.47](#), [4.7.48](#), [4.7.49](#), [4.7.51](#), [4.7.53](#), [4.7.54](#), [4.7.56](#), [4.7.61](#)
- [148] ISIDRO, JOSÉ M. und ÁNGEL RODRÍGUEZ PALACIOS: *On the definition of real W^* -algebras*. Proceedings of the American Mathematical Society, 124(11):3407–3410, November 1996. [3.7.22](#), [3.7.27](#), [3.7.42](#)
- [149] ISIDRO, JOSÉ M. und LÁZLÓ L. STACHÓ: *Holomorphic automorphism groups in Banach spaces: An elementary introduction*, Band 105 der Reihe *North-Holland Mathematics studies*. Elsevier Science Publishers B.V., Amsterdam, The Netherlands, 1984. [2.12.15](#), [4.5.7](#), [4.5.10](#), [4.5.11](#), [4.5.18](#), [4.5.25](#), [4.5.28](#), [4.5.30](#), [4.5.36](#), [4.6.10](#)
- [150] JACOBSON, NATHAN: *Structure of rings*, Band XXXVII. American Mathematical Society, Colloquium Publications, 1964. Revised. [2.1.33](#), [2.1.34](#), [2.1.43](#), [2.1.45](#), [2.1.46](#), [2.1.59](#)
- [151] JACOBSON, NATHAN: *Structure and representation of Jordan algebras*, Band 39 der Reihe *American Mathematical Society Colloquium Publications*. American Mathematical Society, Providence, Rhode Island, 1968. [2.3.17](#)
- [152] JACOBSON, NATHAN: *Basic algebra I + II*. W. H. Freeman and Company, San Francisco, 1974 + 1980. [2.1.9](#), [2.1.34](#), [2.1.34](#), [2.3.8](#), [2.3.30](#), [2.3.31](#), [3.3.4](#), [3.3.5](#), [3.3.6](#), [3.3.7](#)
- [153] JAMESON, GRAHAM: *Ordered Linear Spaces*, Band 141 der Reihe *Lecture Notes in Mathematics*. Springer-Verlag Berlin, Heidelberg, New York, 1970. [2.2.6](#)
- [154] JÄNICH, KLAUS: *Funktionentheorie. Eine Einführung*. Springer-Verlag Berlin, Heidelberg, 2004 (1977). 6. Auflage. [1.2.30](#)
- [155] JECH, THOMAS: *Set theory. The third Millenium edition, revised and expanded*. Springer-Verlag, Berlin, Heidelberg, 2003. [1.1.1](#), [1.1.2](#)
- [156] JEGGLE, HANSGEORG: *Nichtlineare Funktionalanalysis: Existenz von Lösungen nichtlinearer Gleichungen*. B. G. Teubner, Stuttgart, 1979. [4.1.6](#)
- [157] JOHN, K. und VACLAV ZIZLER: *On rough norms on Banach spaces*. Commentationes mathematicae universitatis Carolinae, 19(2):335–349, 1978. [5.2.6](#)
- [158] JOHNSON, WILLIAM B. und JORAM LINDENSTRAUSS (Herausgeber): *Handbook of the geometry of Banach spaces. Volume 1 and 2*. Elsevier, 2001 und 2003. [1.2.24](#), [2.4.24](#), [2.11.11](#)
- [159] JONES, V. und V. S. SUNDER: *Introduction to Subfactors*, Band 234 der Reihe *London Mathematical Society Lecture Notes Series*. Cambridge University Press, 1997. [5.2.4](#)
- [160] KADETS, VLADIMIR, NIGEL KALTON und DIRK WERNER: *Remarks on rich subspaces of Banach spaces*. Studia Mathematica, 159(2):195–206, 2003. [5.1.2](#), [5.1.6](#), [5.1.7](#)
- [161] KADETS, VLADIMIR, MIGUEL MARTÍN und JAVIER MERÍ: *Norm inequalities for operators on Banach spaces*. Indiana Univ. Math. J., 56(5):2385–2412, 2007. [5.1.2](#)
- [162] KADETS, VLADIMIR, VARVARA SHEPELSKA und DIRK WERNER: *Quotients of Banach spaces with the Daugavet property*. Bull. Pol. Acad. Sci., 56(2):131–147, 2008. [5.1.8](#)
- [163] KADETS, VLADIMIR M., ROMAN V. SHVIDKOY, GLEB G. SIROTKIN und DIRK WERNER: *Banach spaces with the Daugavet property*. Trans. Am. Math. Soc., 352(2):855–873, 2000. [5.1.1](#), [5.1.2](#), [5.1.3](#), [5.1.5](#)
- [164] KADETS, VLADIMIR M., ROMAN V. SHVIDKOY und DIRK WERNER: *Narrow operators and rich subspaces of Banach spaces with the Daugavet property*. Studia mathematica, 147(3):269–298, 2001. [5.1.2](#), [5.1.6](#), [5.1.7](#)
- [165] KADISON, RICHARD VINCENT: *Order properties of bounded self-adjoint operators*. Proceedings of the American Mathematical Society, 2(3):505–510, June 1951. [3.5.48](#)
- [166] KADISON, RICHARD VINCENT: *Operator algebras with a faithful weakly closed representation*. Annals of Mathematics, 64(1):175–181, July 1956. [3.7.39](#)

- [167] KADISON, RICHARD VINCENT und JOHN R. RINGROSE: *Fundamentals of the theory of operator algebras*, Band 100 der Reihe *Pure and applied mathematics*. Academic Press, 1983 + 1986 + 1991. [2.2.17](#), [2.2.20](#), [2.4.1](#), [2.4.29](#), [2.4.34](#), [3.4.22](#), [3.5.21](#), [3.5.33](#), [3.5.44](#), [3.5.47](#), [3.5.53](#), [3.5.55](#), [3.7.39](#), [3.7.45](#), [3.7.48](#), [3.7.51](#), [3.7.59](#)
- [168] KAMENSKII, MIKHAIL, VALERI OBUKHOVSKII und PIETRO ZECCA: *Condensing multivalued maps and semilinear differential inclusions in Banach spaces*, Band 7 der Reihe *De Gruyter series in nonlinear analysis and applications*. Walter de Gruyter, Berlin, 2001. [2.4.40](#)
- [169] KAMTHAN, P. K. und MANJUL GUPTA: *Sequence spaces and series*. Marcel Dekker, Inc. New York, Basel, 1981. [2.4.7](#)
- [170] KAPLANSKY, IRVING: *Normed algebras*. Duke Math. J., 16:399–418, 1949. [2.9.18](#), [2.9.27](#)
- [171] KAPLANSKY, IRVING: *Rings of operators*. W. A. Benjamin, Inc., New York, Amsterdam, 1968. [3.7.14](#)
- [172] KAUP, LUDGER und BURCHARD KAUP: *Holomorphic functions of several variables: An introduction to the fundamental theory*, Band 3 der Reihe *de Gruyter studies in mathematics*. de Gruyter, Berlin, New York, 1983. [4.5.21](#), [4.5.22](#)
- [173] KAUP, WILHELM: *Algebraic characterization of symmetric complex Banach manifolds*. Math. Ann., 228(1):39–64, 1977. [4.2.1](#), [4.5.2](#), [4.5.18](#), [4.5.19](#), [4.6](#), [4.6.21](#), [4.7.2](#), [4.7.26](#)
- [174] KAUP, WILHELM: *Über die Klassifikation der symmetrischen hermiteschen Mannigfaltigkeiten unendlicher Dimension I*. Math. Ann., 257:463–483, 1981. [4.6.11](#), [4.7.58](#)
- [175] KAUP, WILHELM: *A Riemann mapping theorem for bounded symmetric domains in complex Banach spaces*. Mathematische Zeitschrift, 183:503–529, 1983. [4.5.25](#), [4.5.38](#), [4.6.24](#), [4.6.25](#), [4.6.26](#), [4.6.27](#), [4.7.7](#), [4.7.16](#), [4.7.17](#), [4.7.22](#)
- [176] KAUP, WILHELM: *Contractive projections on Jordan C^* -algebras and generalizations*. Math. Scand., 54:95–100, 1984. [4.7.27](#)
- [177] KAUP, WILHELM: *On real Cartan factors*. manuscripta mathematica, 92:191–222, 1997. [4.7.58](#), [4.7.59](#), [4.7.72](#)
- [178] KAUP, WILHELM: *Bounded symmetric domains and derived geometric structures*. Rend. Mat. Acc. Lincei, Serie IX, 13:243–257, 2002. [4.5.17](#), [4.5.27](#), [4.5.38](#), [4.7.21](#), [4.7.22](#), [4.7.26](#)
- [179] KAUP, WILHELM: *Bounded symmetric domains and generalized operator algebras*. In: *The 46th Real Analysis and Functional Analysis Joint Symposium, August 7–9, 2007, Fukuoka, Kyushu University, Japan*, Seiten 46–56. 2007. [4.5.34](#), [4.6.12](#), [4.7.11](#), [4.7.14](#), [4.7.16](#), [4.7.22](#), [4.7.26](#), [4.7.27](#)
- [180] KAUP, WILHELM: *Komplexe Analysis. Vorlesung*, April 2007. <http://www.mathematik.uni-tuebingen.de/~kaup/skripte.html>, Download vom 12.2.2010. [4.1.6](#), [4.2.7](#), [4.4.1](#), [4.4.2](#), [4.4.3](#), [4.5.29](#)
- [181] KAUP, WILHELM und HARALD UPMEIER: *Banach spaces with biholomorphically equivalent unit balls are isomorphic*. Proc. Amer. Math. Soc., 58:129–133, July 1976. [4.5.29](#), [4.5.32](#)
- [182] KAUP, WILHELM und HARALD UPMEIER: *An infinitesimal version of Cartan's uniqueness theorem*. manuscripta math., 22:381–401, 1977. [4.2.11](#), [4.5.10](#)
- [183] KAUP, WILHELM und HARALD UPMEIER: *Jordan algebras and symmetric Siegel domains in Banach spaces*. Mathematische Zeitschrift, 157:179–200, 1977. [4.7.62](#), [4.7.62](#)
- [184] KELLEY, JOHN LEROY: *General Topology*. David van Nostrand Company, Inc., Princeton, New Jersey, Toronto, London, 1955. April 1967. [1.1.1](#), [1.1.10](#), [1.1.16](#), [1.2.15](#), [4.5.23](#)
- [185] KELLEY, JOHN LEROY, ISAAC NAMIOKA, WILLIAM F. DONOGHUE, JR., G. BALEY PRICE, KENNETH R. LUCAS, WENDY ROBERTSON, B. J. PETTIS, W. R. SCOTT, EBBE THUE POULSEN und KENNAN T. SMITH: *Linear Topological Spaces*. David van Nostrand Company, Inc., Princeton, New Jersey, 1963. [2.4.4](#), [2.4.26](#), [2.4.30](#), [2.4.36](#), [2.5.1](#), [2.5.7](#), [2.5.14](#)
- [186] KLINGENBERG, WILHELM und PETER KLEIN: *Lineare Algebra und analytische Geometrie. Band 2*. Bibliographisches Institut AG, Mannheim, 1972. [2.1.35](#), [2.9.2](#), [2.9.4](#), [2.9.21](#), [2.9.22](#), [2.9.25](#), [3.1.5](#), [3.3.1](#)

- [187] KÖTHE, GOTTFRIED: *Topologische lineare Räume*, Band 107 der Reihe *Die Grundlehren der mathematischen Wissenschaften in Einzeldarstellungen mit besonderer Berücksichtigung der Anwendungsgebiete*. Springer-Verlag Berlin, Heidelberg, New York, 1966. [2.2.8](#), [2.2.27](#), [2.2.32](#), [2.2.36](#), [2.4.3](#), [2.4.15](#), [2.4.16](#), [2.4.33](#), [2.5.1](#), [2.5.3](#), [2.5.5](#), [2.6.9](#), [2.6.12](#), [3.1.10](#)
- [188] KOWALSKY, HANS-JOACHIM und GERHARD O. MICHLER: *Lineare Algebra*. Walter de Gruyter, Berlin, New York, 2003. 12., überarbeitete Auflage. [2.9.25](#)
- [189] KRANTZ, STEVEN GEORGE: *Function theory of several complex variables*. John Wiley & Sons, New York, Chichester, Bristane, Toronto, Singapore, 1982. [4.5.23](#)
- [190] KULKARNI, S. H. und B. V. LIMAYE: *Gelfand-Naimark theorems for real Banach $*$ -algebras*. Math. Japonica, 25(5):545–558, 1980. [2.12.4](#), [3.5.11](#)
- [191] KUNEN, KENNETH: *Set Theory*, Band 102 der Reihe *Studies in Logic and the Foundation of Mathematics*. North-Holland Publishing Company, Amsterdam, New York, Oxford, 1980. [1.1.1](#), [1.1.1](#), [1.1.11](#)
- [192] KURATOWSKI, K.: *Topology I*. Academic Press, New York, London, 1966. [1.2.8](#)
- [193] KUROSCH, ALEXANDER GENNADJEWITSCH: *Vorlesungen über allgemeine Algebra*. Übersetzung: H.-H. Buchsteiner, Magdeburg; wissenschaftliche Betreuung: Prof. A. Kertész, Debrecen (z. Zt. in Halle). Verlag Harri Deutsch, Zürich und Frankfurt/Main, 1964. [1.1.13](#), [2.1.1](#), [2.1.28](#), [2.1.32](#)
- [194] LAM, TSIT-YUEN: *A first course in noncommutative rings*. Springer-Verlag New York, Inc., 2. Auflage, 2001. [2.1.33](#)
- [195] LANG, SERGE: *Differential manifolds*. Springer-Verlag, 1988. Second printing. [4.2.5](#)
- [196] LARSEN, RONALD: *Banach Algebras. An Introduction*. Marcel Dekker, Inc., New York, 1973. [2.3.2](#), [2.12.4](#), [3.3.2](#), [5.2.4](#)
- [197] LI, BINGREN: *Introduction to Operator Algebras*. World Scientific, Singapore, New Jersey, London, Hong Kong, 1992. [3.5.33](#), [3.7.9](#), [3.7.13](#), [3.7.20](#), [3.7.34](#), [3.7.57](#), [3.7.58](#), [5.2.4](#)
- [198] LI, BINGREN: *Real operator algebras*. World Scientific Publishing Co. Private Limited, River Edge, New Jersey, 2003. [2.9.9](#), [2.9.27](#), [2.10.10](#), [2.12.17](#), [3.1.12](#), [3.1.13](#), [3.1.15](#), [3.2.16](#), [3.2.19](#), [3.2.33](#), [3.2.37](#), [3.4.10](#), [3.5.11](#), [3.5.13](#), [3.5.15](#), [3.5.17](#), [3.5.21](#), [3.5.27](#), [3.5.33](#), [3.5.40](#), [3.5.52](#), [3.5.53](#), [3.5.54](#), [3.6.11](#), [3.6.13](#), [3.6.14](#), [3.6.17](#), [3.7.1](#), [3.7.5](#), [3.7.9](#), [3.7.12](#), [3.7.13](#), [3.7.20](#), [3.7.25](#), [3.7.26](#), [3.7.27](#), [3.7.29](#), [3.7.34](#), [3.7.44](#), [3.7.45](#), [3.7.47](#), [3.7.52](#), [3.7.57](#), [3.7.58](#), [3.7.60](#)
- [199] LIN, BOR-LUH, PEI-KEE LIN und S. L. TROYANSKI: *Characterizations of denting points*. Proceedings of the American Mathematical Society, 102(3):526–528, March 1988. [2.6.14](#), [2.6.15](#)
- [200] LOOS, OTTMAR: *Lectures on Jordan triples*. The University of British Columbia, Vancouver, Canada, September 1971. [4.6.10](#), [4.7.59](#)
- [201] LOOS, OTTMAR: *Jordan pairs*, Band 460 der Reihe *Lecture Notes in Mathematics*. Springer, 1975. [4.6.6](#)
- [202] LOOS, OTTMAR: *Bounded symmetric domains and Jordan pairs*. The University of California at Irvine, 1977. [4.6.6](#), [4.6.9](#), [4.7.61](#)
- [203] LUMER, G.: *Semi-inner-product spaces*. Trans. Amer. Math. Soc., 100:29–43, 1961. [2.11.20](#)
- [204] MARTÍN, MIGUEL: *The Daugavetian index of a Banach space*. Taiwanese Journal of Mathematics, 7(4):631–640, December 2003. [5.1.2](#)
- [205] MARTÍN, MIGUEL und TIMUR OIKHBERG: *An alternative Daugavet property*. J. Math. Appl., 294(1):158–180, 2004. [5.1.2](#)
- [206] MARTÍNEZ, JUAN und ANTONIO M. PERALTA: *Separate weak*-continuity of the triple product in dual real JB^* -triples*. Mathematische Zeitschrift, 234:635–646, 2000. [4.7.49](#), [4.7.59](#), [4.7.75](#)

- [207] MARTÍNEZ MORENO, J., J. F. MENA JURADA, R. PAYÁ ALBERT und ÁNGEL RODRÍGUEZ PALACIOS: *An approach to numerical ranges without Banach algebra theory*. Illinois Journal of Mathematics, 29(4):609–626, winter 1985. [2.11.16](#)
- [208] MATHIEU, MARTIN: *Funktionalanalysis*. Spektrum Akademischer Verlag, Heidelberg, Berlin, 1998. [2.3.26](#), [2.4.19](#), [3.4.11](#), [3.4.19](#), [3.4.20](#), [3.5.33](#), [3.5.33](#), [3.5.36](#), [3.5.47](#), [3.6.6](#), [3.7.23](#), [3.7.31](#), [3.7.44](#), [3.7.48](#), [3.7.48](#), [3.7.49](#), [3.7.53](#), [3.7.55](#)
- [209] MCCOY, NEAL HENRY: *The theory of rings*. The Macmillan Company, New York; Collier-Macmillan Limited, London, 1964. [2.1.13](#), [2.1.14](#), [2.1.33](#), [2.1.34](#), [2.1.43](#), [2.1.44](#), [2.1.50](#), [2.1.59](#)
- [210] MCCRIMMON, KEVIN: *A taste of Jordan algebras*. Springer, 2004. [2.3.9](#), [2.3.14](#), [2.3.17](#), [2.9.27](#), [2.11.26](#), [3.3.7](#), [4.5.17](#), [4.5.25](#), [4.6.13](#)
- [211] MEGGINSON, ROBERT E.: *An introduction to Banach space theory*, Band 183 der Reihe *Graduate texts in mathematics*. Springer, New York, 1998. [2.1.25](#), [2.2.6](#), [2.4.10](#), [2.4.15](#), [2.4.19](#), [2.4.31](#), [2.5.11](#), [2.5.16](#), [2.6.4](#), [2.7.2](#), [2.9.19](#), [2.10.9](#), [2.11.3](#), [2.11.5](#), [2.11.6](#), [3.5.55](#), [4.1.5](#), [4.1.11](#), [4.1.12](#)
- [212] MEYBERG, KURT: *Lectures on algebras and triple systems*. The University of Virginia, Charlottesville, 1972. [4.6.7](#), [4.7.25](#), [4.7.59](#)
- [213] MEYBERG, KURT: *Algebra. Teil 2*. Carl Hanser Verlag, München, Wien, 1976. 1. Auflage. [4.3.4](#)
- [214] MEYBERG, KURT: *Algebra. Teil 1*. Carl Hanser Verlag, München, Wien, 1980. 2. Auflage. [2.1.22](#), [2.1.43](#)
- [215] MOORE, ROBERT T.: *Hermitian functionals on B -algebras and duality characterizations of C^* -algebras*. Transactions of the American Mathematical Society, 162:253–265, 1971. [3.5.37](#)
- [216] MUKHERJEA, A. und K. POTHoven: *Real and Functional Analysis. Part B: Functional Analysis*, Band 27 der Reihe *Mathematical concepts and methods in science and engineering*. Plenum Press, New York and London, 2. Auflage, 1986. [2.2.44](#)
- [217] MURPHY, GERARD J.: *C^* -Algebras and operator theory*. Academic Press, Inc., 2004. [3.4.22](#), [3.5.21](#), [3.5.43](#), [3.5.53](#), [3.6.3](#), [3.7.6](#), [3.7.14](#), [3.7.16](#), [3.7.29](#)
- [218] NACHBIN, LEOPOLDO: *Topology on spaces of holomorphic mappings*. Springer-Verlag Berlin, Heidelberg, New York, 1969. [4.1.14](#)
- [219] NACHBIN, LEOPOLDO: *Introduction to functional analysis: Banach spaces, and differential calculus*, Band 60 der Reihe *Monographs and textbooks in pure and applied mathematics*. Marcel Dekker, Inc., New York and Basel, 1981. Translation from the Portuguese by Richard M. Aron, University of Dublin, Dublin, Ireland, of *Introdução à análise funcional: Espaços de Banach e Cálculo Diferencial*. Barra da Tijuca, Rio de Janeiro, Brazil and University of Rochester, Rochester, New York. [4.1.6](#)
- [220] NAGUMO, MITIO: *Einige analytische Untersuchungen in linearen, metrischen Ringen*. Japanese Journal of Mathematics, 13:61–80, 1936. [2.12.8](#)
- [221] NAIMARK, MARK ARONOVICH: *Normed algebras. 3. compl. rev. American ed.; translated from 2nd Russian edition published 1968, Moscow. (1st Russian edition: 1955, translated 1959)*. Wolters-Noordhoff series of monographs and textbooks on pure and applied mathematics. Wolters-Noordhoff, Groningen, 1972. [2.1.47](#), [2.3.2](#), [2.3.29](#), [2.7.2](#), [3.2.10](#), [3.2.28](#), [3.2.29](#), [3.6.15](#), [3.6.16](#)
- [222] NARASIMHAN, RAGHAVAN: *Several complex variables*. Chicago Lectures in Mathematics Series. The Chicago University Press, Chicago and London, 1971. [4.5.23](#)
- [223] NEAL, MATTHEW und BERNARD RUSSO: *State spaces of JB^* -triples*. Mathematische Annalen, 328:585–624, 2004. [2.4.22](#), [4.7.78](#)
- [224] NEUMANN, JOHN VON: *Zur Algebra der Funktionaloperationen und Theorie der normalen Operatoren*. Mathematische Annalen, 102:370–427, 1929. [3.7.2](#)

- [225] OBERSCHHELP, ARNOLD: *Set theory over classes*. Doktorarbeit, Westfälische Wilhelms-Universität, Münster, 1973. Dissertationes Mathematicae; CVI (106), Warszawa. [1.1.1](#)
- [226] OBERSCHHELP, ARNOLD: *Allgemeine Mengenlehre*. B.I.-Wissenschaftsverlag, Mannheim, Leipzig, Wien, Zürich; (Bibliographisches Institut & F.A. Brockhaus AG, Mannheim), 1994. [1.1.1](#)
- [227] ONO, TAMIO: *A real analogue of the Gelfand-Neumark theorem*. Proceedings of the American Mathematical Society, 25(1):159–160, May 1970. [3.5.11](#)
- [228] OUCHEVA, VERA: *Symmetric complex Banach manifolds and hermitian Jordan triple systems*. Diplomarbeit bei Wilhelm Kaup, Universität Tübingen, 2001. [4.5.18](#), [4.7.26](#)
- [229] PACHALE, HELMUT: *Nichtlineare Funktionalanalysis*. Vorlesungsskript SS 1974 Freie Universität Berlin, 1974. [4.1.6](#)
- [230] PALMER, THEODORE WINDLE: *Real C^* -algebras*. Pacific Journal of Mathematics, 35(1):195–204, 1970. [3.2.32](#), [3.5.11](#)
- [231] PALMER, THEODORE WINDLE: *Banach Algebras and the General Theory of $*$ -Algebras. Volume I: Algebras and Banach Algebras*. Nummer 49 in *Encyclopedia of Mathematics and its Applications*. Cambridge University Press, 1994. [2.1.9](#), [2.1.33](#), [2.1.42](#), [2.1.43](#), [2.1.44](#), [2.1.51](#), [2.1.52](#), [2.1.60](#), [2.2.26](#), [2.3.2](#), [2.3.18](#), [2.3.22](#), [2.3.23](#), [2.3.24](#), [2.3.26](#), [2.7.2](#), [2.7.4](#), [2.8.1](#), [2.11.10](#), [2.11.23](#), [2.12.2](#), [2.12.4](#), [2.12.13](#), [2.12.15](#), [2.12.16](#), [2.12.17](#), [3.2.24](#), [3.3.2](#), [3.3.3](#)
- [232] PALMER, THEODORE WINDLE: *Banach Algebras and the General Theory of $*$ -Algebras. Volume II: $*$ -algebras*. Nummer 79 in *Encyclopedia of Mathematics and its Applications*. Cambridge University Press, 2001. [2.10.8](#), [2.10.10](#), [3.2.19](#), [3.2.34](#), [3.4.9](#), [3.4.23](#), [3.4.25](#), [3.4.26](#), [3.4.27](#), [3.4.29](#), [3.5.1](#), [3.5.8](#), [3.5.10](#), [3.5.22](#), [3.5.25](#), [3.5.26](#), [3.5.29](#), [3.5.30](#), [3.5.33](#), [3.5.38](#), [3.5.41](#), [3.5.49](#), [3.5.50](#), [3.5.51](#), [3.5.53](#), [3.6.2](#), [3.6.3](#), [3.6.12](#), [3.6.15](#), [3.7.1](#), [3.7.4](#), [3.7.6](#), [3.7.11](#), [3.7.13](#), [3.7.19](#), [3.7.30](#), [3.7.37](#), [3.7.40](#), [3.7.41](#)
- [233] PEDERSEN, GERT KJÆRGÅRD: *C^* -Algebras and their automorphism groups*, Band 14 der Reihe *London Mathematical Society Monographs*. Academic Press Inc. (London) Ltd., 1979. [2.4.19](#), [3.4.21](#), [3.5.34](#), [3.5.53](#), [3.7.22](#), [3.7.39](#), [3.7.56](#), [3.7.57](#)
- [234] PEDERSEN, GERT KJÆRGÅRD: *Analysis Now*, Band 118 der Reihe *Graduate texts in Mathematics*. Springer-Verlag New York, Berlin, Heidelberg, London, Paris, Tokyo, 1989. [2.4.19](#), [3.2.12](#), [3.4.14](#), [3.7.29](#)
- [235] PERALTA, A. und L. L. STACHÓ: *Atomic decomposition of real JBW^* -triples*. The Quarterly Journal of Mathematics, 52:79–87, 2001. [4.7.59](#), [4.7.66](#), [4.7.66](#), [4.7.75](#), [4.7.77](#)
- [236] PERALTA, ANTONIO M.: *On the axiomatic definition on real JB^* -triples*. Math. Nachr., 256:100–106, 2003. [4.7.45](#)
- [237] PFLAUMANN, ERIKA und HEINZ UNGER: *Funktionalanalysis II*. Bibliographisches Institut AG, Zürich, 1974. [4.1.5](#)
- [238] PHELPS, ROBERT RALF: *Convex functions, monotone operators and differentiability*, Band 1364 der Reihe *Lecture Notes in Mathematics*. Springer-Verlag Berlin Heidelberg, 2. Auflage, 1993. (1. Edition war 1989). [2.11.4](#), [2.11.11](#), [4.1.8](#)
- [239] PIERCE, RICHARD S.: *Associative algebras*, Band 88 der Reihe *Graduate texts in mathematics*. Springer-Verlag New York, Heidelberg, Berlin, 1982. [2.1.43](#)
- [240] PINCHUK, S. I.: *Holomorphic mappings in C^n and the problem of holomorphic equivalence*. In: KHENKIN (Herausgeber): *Several complex variables III*, Band 9 der Reihe *Geometric function theory; Encyclopaedia of Mathematical Sciences*, Seiten 173–199. Springer, 1989. translation from Itogi Nauki Tekh., Ser. Sovrem. Probl. Mat., Fundam. Napravleniya 9, 195–223 (1986). [4.5.23](#)
- [241] PITTS, C. G. C.: *Introduction to metric spaces*. University Mathematical Texts. Oliver & Boyd, Edinburgh, 1972. [2.1.26](#)
- [242] POINCARÉ, HENRI: *Les fonctions analytiques de deux variables et la représentation conforme*. Rendiconti del Circolo Matematico di Palermo, 23(1):185–220, 1907. [4.5.23](#)

- [243] PREUSS, GERHARD: *Allgemeine Topologie*. Springer-Verlag Berlin, Heidelberg, New York, 1975 (1972). 2., korrigierte Auflage. [1.1.16](#), [1.2.16](#)
- [244] QUERENBURG, BOTO VON: *Mengentheoretische Topologie*. Springer-Verlag, Berlin, Heidelberg, New York, 1979. 2., neubearbeitete und erweiterte Auflage. [1.1.2](#), [1.1.7](#), [1.1.16](#), [1.2.1](#), [1.2.3](#), [1.2.30](#), [1.2.31](#)
- [245] REED, MICHAEL und BARRY SIMON: *Methods of modern mathematical physics. Volume 1: Functional analysis*. Academic press, 1980. Revised and enlarged. [2.2.11](#), [2.5.1](#), [3.6.2](#)
- [246] RICKART, CHARLES E.: *General theory of Banach algebras*. David van Nostrand Company, Inc., 1960. [2.1.9](#), [2.1.20](#), [2.3.20](#), [2.4.10](#), [2.8.7](#), [2.9.27](#), [2.12.1](#), [2.12.2](#), [2.12.4](#), [2.12.6](#), [2.12.14](#), [2.12.17](#), [3.2.7](#), [3.2.8](#), [3.2.9](#), [3.2.10](#), [3.2.23](#), [3.2.25](#), [3.2.28](#), [3.2.31](#), [3.2.32](#), [3.3.2](#), [3.3.3](#), [3.5.22](#), [3.5.26](#), [3.5.54](#)
- [247] ROBERTSON, A. P. und WENDY ROBERTSON: *Topologische Vektorräume*. B.I.-Hochschultaschenbücher 164/164a. Bibliographisches Institut, Mannheim; Hochschultaschenbücher-Verlag, 1967. [2.4.3](#), [2.4.8](#), [2.4.16](#), [2.5.1](#)
- [248] RODRÍGUEZ PALACIOS, ÁNGEL: *Banach space characterizations of unitaries: A survey*. Journal of Mathematical Analysis and Applications, 369:168–178, 2010. [4.7.3](#)
- [249] ROSENTHAL, HASKELL: *The Lie Algebra of a Banach space*. In: *Banach spaces (Columbia, Missouri, 1984)*, Band 1166 der Reihe *Lecture Notes in Mathematics*. Springer Berlin, 1984. [2.9.19](#), [2.11.22](#), [4.1.9](#)
- [250] ROWEN, LOUIS H.: *Ring theory, volume I*, Band 127 der Reihe *Pure and Applied Mathematics*. Academic Press, Inc., 1988. [1.1.13](#), [2.1.9](#), [2.1.34](#), [2.1.34](#), [2.1.43](#)
- [251] RUDIN, WALTER: *Functional Analysis*. McGraw-Hill, 1973. [2.2.44](#), [2.4.4](#), [2.4.8](#), [2.11.14](#), [3.2.30](#)
- [252] SAKAI, SHŌICHIRO: *Weakly compact operators on operator algebras*. Pac. J. Math., 14(2):659–664, 1964. [3.5.55](#)
- [253] SAKAI, SHŌICHIRO: *C^* -algebras and W^* -algebras*. Springer-Verlag, Berlin, Heidelberg, New York, 1971. [3.7.1](#), [3.7.24](#), [3.7.42](#)
- [254] SATAKE, ICHIRO: *Algebraic structures of symmetric domains*, Band 4 der Reihe *Kanô Memorial Lectures*. Iwanami Shoten, Publishers and Princeton University Press, 1980. Publications of the Mathematical Society of Japan. [4.7.59](#)
- [255] SCHAEFER, HELMUT HANS: *Topological vector spaces*. The MacMillan Company, New York; Collier-MacMillan Limited, London, 1967. 2nd print. [2.2.6](#), [2.2.11](#), [2.4.10](#), [2.4.14](#), [2.4.15](#), [2.4.30](#), [2.9.9](#)
- [256] SCHAFER, RICHARD D.: *An introduction to nonassociative algebras*. Academic Press, 1966. [2.1.30](#), [2.1.32](#), [2.1.34](#), [2.3.1](#), [2.3.13](#), [2.3.14](#), [2.3.19](#), [2.3.20](#), [2.3.31](#), [3.3.7](#)
- [257] SCHRÖDER, HERBERT: *K -theory for real C^* -algebras and applications*. Longman Scientific & Technical, UK, 1993. [3.5.22](#)
- [258] SCHULZE, VOLKER: *Einführung in die Algebra und Zahlentheorie*. Skript zur Vorlesung im Wintersemester 2005/2006 an der Freien Universität Berlin, 2006. [2.1.1](#), [2.1.22](#), [2.1.43](#), [3.3.1](#), [4.3.4](#), [4.3.9](#)
- [259] SHERMAN, S.: *Order in operator algebras*. Amer. Journal of Mathematics, 73(1):227–232, Jan. 1951. [3.5.48](#)
- [260] SHVYDKOY, ROMAN V.: *Geometric aspects of the Daugavet property*. J. Func. Anal., 176:198–212, 2000. [5.1.2](#), [5.1.3](#)
- [261] SIMONS, STEPHEN: *From Hahn-Banach to monotonicity*, Band 1693 der Reihe *Lecture Notes in Mathematics*. Springer Science+Business media B.V., 2., expanded Auflage, 2008. (1. edition: Minimax and monotonicity, LNM; 1693, Springer-Verlag Berlin Heidelberg, 1998). [2.11.4](#)
- [262] SINCLAIR, ALLAN M.: *The norm of a hermitian element in a Banach algebra*. Proceedings of the American Mathematical Society, 28(2):446–450, May 1971. [2.12.15](#)

- [263] SMITH, MARK A. und FRANCIS SULLIVAN: *Extremely smooth Banach spaces*. In: *Banach spaces of analytic functions. Proceedings of the Pelczynski Conference. Held at Kent State University, July 12-16, 1976. Edited by J. Baker and C. Cleaver and J. Diestel*, Band 604 der Reihe *Lecture Notes in Mathematics*, Seiten 125–137. Springer-Verlag, Berlin, Heidelberg, New York, 1977. [2.11.13](#)
- [264] SMITH, PETER: *The Galois Connection between Syntax and Semantics*. Download vom 2. April 2012. www.logicmatters.net/resources/pdfs/Galois.pdf, 2010. [1.1.18](#)
- [265] STACHÓ, LÁSZLÓ L.: *On the manifold of tripotents in JB^* -triples*. Webseite von L. L. Stachó, 2007. [4.7.32](#), [4.7.59](#)
- [266] STACHÓ, LÁSZLÓ L. und JOSÉ M. ISIDRO: *Algebraically compact elements of JBW^* -triples*. *Acta. Sci. Math.*, 54:171–190, 1990. [4.6.12](#), [4.7.21](#)
- [267] STACHÓ, LÁSZLÓ L.: *A counterexample concerning contractive projections of real JB^* -triples*. *Publicationes Mathematicae Debrecen*, 58(1-2):223–230, 2001. [4.7.63](#)
- [268] STORCH, UWE und HARTMUT WIEBE: *Lehrbuch der Mathematik in vier Bänden*. Spektrum Akademischer Verlag; Heidelberg, Berlin, 1993-2003. [4.3.9](#)
- [269] STORCH, UWE und HARTMUT WIEBE: *Analysis einer Veränderlichen*, Band 1 der Reihe *Lehrbuch der Mathematik in vier Bänden*. Spektrum Akademischer Verlag; Heidelberg, Berlin, 2003. 3. Auflage. [1.1.2](#)
- [270] STRĂTILĂ, SERBAN und LÁSZLÓ ZSIDÓ: *Lectures on von Neumann algebras*. Editura Academici, Bucuresti, Romania, Abacus Press, Turnbridge Wells, Kent, England, 1979. Revised and updated version of: *Lectii de algebre von Neumann*. [2.2.20](#), [2.4.35](#), [3.4.21](#), [3.7.31](#), [3.7.56](#), [3.7.59](#)
- [271] SUNDER, VIAKALATHUR SHANKAR: *An invitation to von Neumann algebras*. Springer-Verlag, New York, Berlin, Heidelberg, London, Paris, Tokyo, 1987. [3.2.39](#), [3.4.24](#), [3.6.2](#), [3.6.10](#), [3.7.5](#), [3.7.31](#), [3.7.46](#), [5.2.4](#)
- [272] SUNDER, VIAKALATHUR SHANKAR: *Functional Analysis: Spectral Theory*. Webseite von V. S. Sunder, July 2000. www.imsc.res.in/~sunder/fa.pdf. [3.4.21](#)
- [273] SUPPES, PATRICK: *Axiomatic Set Theory*. Dover Publications, Inc., New York, 1972. Unabridged and corrected republication of the work originally published by D. Van Nostrand Company in 1960. [1.1.1](#)
- [274] TAKESAKI, MASAMICHI: *Theory of operator algebras*. Springer-Verlag, Berlin, Heidelberg, New York, 1979, 2001. [2.2.17](#), [3.4.26](#), [3.5.48](#), [3.6.5](#), [3.6.7](#), [3.6.11](#), [3.7.22](#), [3.7.27](#), [3.7.59](#), [5.2.4](#)
- [275] TAYLOR, ANGUS E. und DAVID C. LAY: *Introduction to Functional Analysis*. John Wiley & Sons, New York, Chichester, Brisbane, Toronto, 2. Auflage, 1980. [2.2.35](#), [2.2.36](#), [2.2.44](#), [2.4.27](#), [2.5.5](#), [2.5.6](#), [2.5.11](#), [2.6.2](#), [2.6.6](#), [2.6.22](#), [2.12.8](#), [3.4.19](#)
- [276] THRON, WOLFGANG J.: *Topological Structures*. Holt, Rinehart and Winston, Inc., New York, Chicago, San Francisco, Toronto, London, 1966. [1.1.20](#)
- [277] TIEL, JAN VAN: *Convex analysis*. John Wiley & Sons Ltd. Chichester, New York, Brisbane, Toronto, Singapore, 1984. An introductory text. [2.1.27](#), [2.9.31](#)
- [278] TOPPING, DAVID MORRIS: *Lectures on von Neumann algebras*. Nostrand Reinhold Company, London, 1971. [3.2.38](#), [3.7.4](#), [3.7.15](#)
- [279] UPMEIER, HARALD: *Symmetric Banach manifolds and Jordan C^* -algebras*, Band 104 der Reihe *North-Holland Mathematics studies*. Elsevier Science Publishers, Netherlands, 1985. [2.3.2](#), [2.3.12](#), [2.3.31](#), [2.4.15](#), [2.7.3](#), [2.11.29](#), [2.12.15](#), [3.2.19](#), [3.5.26](#), [4.1.1](#), [4.1.2](#), [4.1.3](#), [4.1.13](#), [4.2.1](#), [4.2.2](#), [4.2.3](#), [4.2.4](#), [4.2.5](#), [4.2.6](#), [4.2.7](#), [4.2.8](#), [4.2.9](#), [4.2.10](#), [4.2.11](#), [4.2.11](#), [4.2.12](#), [4.2.13](#), [4.3.7](#), [4.3.8](#), [4.3.9](#), [4.3.10](#), [4.4.2](#), [4.4.3](#), [4.5.1](#), [4.5.11](#), [4.5.14](#), [4.5.16](#), [4.5.17](#), [4.5.18](#), [4.5.28](#), [4.6.6](#), [4.6.7](#), [4.6.10](#), [4.6.14](#), [4.6.15](#), [4.6.16](#), [4.6.22](#), [4.6.24](#), [4.7.3](#), [4.7.7](#), [4.7.10](#), [4.7.11](#), [4.7.14](#), [4.7.16](#), [4.7.21](#), [4.7.36](#), [4.7.39](#), [4.7.58](#), [4.7.59](#)

- [280] UPMEIER, HARALD: *Jordan algebras in analysis, operator theory, and quantum mechanics*. Nummer 67 in *Regional conference series in mathematics*. American Mathematical Society, 1987. Expository lectures from the CBMS regional conference held at the University of California, Irvine, July 15-19, 1985. [4.1.2](#), [4.5.12](#), [4.5.17](#), [4.7.58](#)
- [281] USMANOV, SHUKHRAT M.: *Takesaki's duality theorem and continuous decomposition for real factors of type III*. *Izvestiya: Mathematics / Izvestiya RAN: Ser. Mat.*, 64(4):827–845 (RAN:183–200), 2000. [3.7.3](#)
- [282] VIGUÉ, JEAN-PIERRE: *Le groupe des automorphismes analytiques d'un domaine borné d'un espace de Banach complexe. Application aux domaines bornés symétriques*. *Annales Scientifiques de l'École Normale Supérieure*, 4e series, 9(2):203–281, 1976. [4.5.8](#), [4.5.11](#), [4.5.16](#), [4.5.28](#)
- [283] VIGUÉ, JEAN-PIERRE: *Les automorphismes analytiques isométriques d'une variété complexe normée*. *Bulletin de la S.M.F. (La Société Mathématique de France)*, 110:49–73, 1982. [4.5.9](#), [4.5.16](#), [4.5.18](#)
- [284] WAERDEN, BARTEL LEENDERT VAN DER: *Algebra II*. Springer-Verlag Berlin, Heidelberg, New York, 1967. 5. Auflage der Modernen Algebra von 1930/1931. [2.1.34](#), [2.1.57](#)
- [285] WAERDEN, BARTEL LEENDERT VAN DER: *Algebra I*. Springer-Verlag Berlin, Heidelberg, New York, 1971. 8. Auflage der Modernen Algebra von 1930/1931. [2.1.7](#), [2.1.22](#)
- [286] WAHRIG, GERHARD, HILDEGARD KRÄMER und HARALD ZIMMERMANN (Herausgeber): *Brockhaus Wahrig : Deutsches Wörterbuch in sechs Bänden*, Band 15-20. F. A. Brockhaus, Wiesbaden, und Deutsche Verlags-Anstalt, Stuttgart, 1980-1984. 18., völlig neu bearbeitete Auflage. [2.1.34](#)
- [287] WEIDMANN, JOACHIM: *Lineare Operatoren in Hilberträumen*. Mathematische Leitfäden. B. G. Teubner Stuttgart, 1976. [2.2.42](#)
- [288] WEIDMANN, JOACHIM: *Lineare Operatoren in Hilberträumen Teil 1 Grundlagen*. Mathematische Leitfäden. B. G. Teubner Stuttgart, Leipzig, Wiesbaden, 2000. [2.2.42](#), [3.4.5](#), [3.4.9](#), [3.4.14](#), [3.4.16](#), [3.4.21](#), [3.5.47](#)
- [289] WERNER, DIRK: *An elementary approach to the Daugavet equation*. *Lect. Notes Pure Appl. Math.*, 175:449–454, 1996. [5.1.2](#)
- [290] WERNER, DIRK: *Recent progress on the Daugavet property*. *Irish Math. Soc. Bulletin*, Seiten 77–97, 2001. [5.1.7](#)
- [291] WERNER, DIRK: *Funktionalanalysis*. Springer-Verlag Berlin, Heidelberg, 2002. [2.4.23](#), [2.4.34](#), [2.5.1](#), [2.5.12](#), [2.5.15](#), [2.5.17](#), [2.6.11](#), [2.9.9](#), [3.4.5](#), [3.4.6](#), [3.4.17](#), [3.5.14](#), [3.6.2](#), [3.6.6](#)
- [292] WERNER, DIRK: *Daugavet's proof of Daugavet's theorem*. Webseite von Dirk Werner an der Freien Universität Berlin, 2011. <http://page.mi.fu-berlin.de/werner99/preprints/preprints.html>, Download vom 29.04.2011. [5.1.2](#)
- [293] WOJTASZCZYK, P.: *Some remarks on the Daugavet equation*. *Proc. Amer. Math. Soc.*, 115:1047–1052, 1992. [5.1.2](#)
- [294] ZEIDLER, EBERHARD: *Vorlesungen über nichtlineare Funktionalanalysis I — Fixpunktsätze* —, Band 1. C BSB B. G. Teubner Verlagsgesellschaft, Leipzig, 1. Auflage, 1976. [295](#)
- [295] ZEIDLER, EBERHARD: *Nonlinear functional analysis and its applications. Volume I: Fixed-point theorems*, Band I. Springer-Verlag, New York, Berlin, Heidelberg, Tokyo, 1986. Translation by Peter R. Wadsack from a completely newly manuscript of [294] which represents a significant enlargement and revision of the original version [294]. [4.1.6](#)
- [296] ZEIDLER, EBERHARD: *Applied Functional Analysis. Applications to Mathematical Physics*, Band 108 der Reihe *Applied Mathematical Sciences*. Springer, 1995. [3.4.21](#)
- [297] ŻELAZKO, WIESŁAW TADEUSZ: *Banach algebras*. Elsevier Publishing Company, Amsterdam, The Netherlands. (In co-edition with PWN - Polish Scientific Publishers, Warszawa, Poland), 1973. [2.3.2](#)

Symbol-, Namen- und Sachverzeichnis

- (A) (das von A erzeugte Ideal), 34
 $(A)_\ell$ (das von A erzeugte Linksideal), 34
 $(A)_r$ (das von A erzeugte Rechtsideal), 34
 1 (multipl. mnemotechn. Eins), 33
 1 -komplementiert, 85
 A^I (Menge aller Abb. von I nach A), 3
 A^Δ , 229, 232, 238, 239, 242
 $A^\times$ (depunktierte Menge), 3
 A_{red} (reduzierende Ideal), 160
 A_p (Kompression von A unter p), 158
 $B(x; r)$ (abg. Kugel mit Radius r um x), 37
 B_X (abgeschlossene Einheitskugel), 58
 $\text{Gal}(E : K)$ (Galoisgruppe), 210
 $J^\perp$ (zu J komplementärer $\mathbb{P}$ -Smd.), 102
 S_X (Einheitssphäre), 58
 $T^\#$ (algebr. duale Abb. von T), 61
 T_p (Kompression von T), 158
 $U(x; r)$ (offene Kugel mit Radius r um x), 37
 $U_a: x \mapsto axa$, 158
 $U_{F,\varepsilon}$ (Nullumgebung), 91
 X -beschränkt, 104
 X^* (topol. Dualraum von X), 82
 $X^\#$ (algebr. Dualraum von X), 61
 X_* (vermeintliches Prädual von X), 123
 X_{sym} (symmetr. Teil e. Banachr.), 221
 X_x (von x erzeugtes Untertripel), 229
 $X_{(*)}$ (selbstadj. El. von X), 139
 $X_{(-*)}$ (schiefs.adj. El. von X), 139
 X_{herm} (hermit. El. von X), 131
 $[E : K]$ (Körpererweiterungsgrad), 210
 $[\cdot, \cdot]$ (LUMER-Produkt), 130
 $[a, b]$, 58, 68
 $[p]_\ell, [p]_r$ (Äquivalenzklassen), 47, 53
 Γ_A (Gelfand-Raum), 136, 161, 164
 $\Phi_A(e)$, 124, 127, 131, 164, 166, 169
 $\Phi_X(\cdot)$, 124
 $*$ verschiedene Produkte
 AB, A^n , 34, 41, 70
 $A \cdot B$ (Modultheorie), 41
 $A \circ B$ ($\circ$ eine Verknüpfung), 28
 M^n (kartes. Produkt e. Moduls), 42
 $X \cong Y$ (isometrische Isomorphie), 58
 $\dashv$ (Kugelannihilator), 124
 $\equiv_\ell, \equiv_r$ (Äquivalenzrelationen), 47
 $\hat{F}$ (das zu F assoz. homog. Polynom), 67, 68
 $\leq$, 79
 $\leq^*$, 146, 188
 $\leq_\ell$, 47, 50, 53, 79, 155
 $\leq_{\text{UR}}$, 155
 $\leq_r$, 47, 49, 50, 79, 155
 $\leq_{\ell r}$, 47, 79, 155
 $\partial f(x_0)$ (Subdifferential von f bei x_0), 125
 $\perp$, 30, 89
 π_{X^*} (kanon. Proj. von X^{***} auf X^*), 101
 $\prod_{\nu \in I} A_\nu$ (kartesisches Produkt), 3
 $\rho(a)$ (Spektralradius von a), 134
 $M \simeq N$ (Isomorphie), 42
 $\tilde{p}$ (symmetrisiertes Polynom), 67, 203, 220, 222, 235
 $\tau(X, Y)$ (Mackey-Topologie), 105, 178
 $\text{Ran}(f)$ (Bild e. mengenwertigen Abb.), 6
 $\text{dist}(A, B)$ (Abstand zweier Mengen A, B), 37, 119
 $\text{dom}_{\text{eff}}(f)$ (effektiver Def.bereich von f), 119
 $\text{epi}(f)$ (Epigraph der Funktion f), 37, 119–121
 $\text{graph}(f)$, 5
 $\text{ran}(f)$ (Bild von f), 3
 $\text{spt}_X(\ell)$ (Träger von $\ell \in X_*$ in X), 254
 $\text{spt}_X(x)$ (Träger von $x \in X$ in X), 254
 $\text{spt}_{X^{**}}(x)$ (Träger von $x \in X$ in X^{**}), 255
 $\tilde{\nabla}_R f(a)$ (R -Quasigradient von f bei a), 202
 $\nabla_R f(a)$ (R -Gradient von f bei a), 119, 198
 $a|b$, 33, 36
 $a|_\ell b$, 33
 $a|_r b$, 33
 $f: A \rightsquigarrow B$ (mengenwertige Abbildung), 6
 $f: A \rightrightarrows B$ (mengenwertige Abbildung), 6
 $f\langle M \rangle$, 6
 f^{-1} (Inverse e. mengenwertigen Abb.), 6
 g_* , 206
 i_X (kan. Inklusionsabb.), 61, 82, 101, 237
 $i_{X,U}$, 61, 82, 83, 101, 102, 121
 $p \sim q$ (äquivalente Proj.), 189
 $p^\perp$, 78
 $s_\ell(T)$ (Linksträger von T), 155
 $s_r(T)$ (Rechtsträger von T), 155
 sp (Spurfunktional), 152
 x^* (ein Element von X^*), 82
 $x^\#$ (ein Element von $X^\#$), 61
 $y \otimes f$ (lineare Abbildung), 64, 83
 $\ominus$, 72, 78, 109, 209
 $\oplus$, 74, 75
 S' (Kommutante von S), 32

- $A^1 (A \times \mathbb{K})$, 76
- $R^1 (R \times \mathbb{Z})$, 45
- $X^1 (X \times \mathbb{K})$, 61
- Abbildung, 3
 - 1-1, 3
 - adjungierte, 78
 - analytische, 197, 204
 - antiisotone, 11
 - antiordnungsisomorphe, 11
 - antitone, 11
 - assoziierte bilineare, 67
 - assoziierte quadratische, 67
 - bianalytische, 197, 204
 - bilineare, 62
 - biunivoque, 3
 - duale, 82, 99, 101, 102, 130, 152, 164
 - algebraische, 61
 - einen Raum fixierende, 83
 - identische, 4
 - injektive, 3
 - isotone, 11
 - kompakte, 83
 - komplex-analytische, 197
 - konvexe, 119
 - n -lineare, 62
 - n -lineare stetige, 83
 - R -lineare, 41
 - mengenwertige, 6
 - monotone, 10
 - nach oben halbstetige, 20, 21
 - nach unten halbstetige, 20, 21
 - oberhalbstetige, 20, 21, 121, 203
 - $(\tau_X - \overline{\tau_Y})$ -oberhalbstetige, 95, 129, 203, 258
 - $(\tau - \overline{d})$ -oberhalbstetige, 95
 - offene, 19
 - one-to-one, 3
 - ordnungserhaltende, 11
 - ordnungsisomorphe, 11
 - ordnungsumkehrende, 11
 - positiv definite, 63, 142
 - positive, 63, 142
 - propere, 119
 - quadratische, 67
 - Rang-Eins, 62, 64, 83, 255
 - reell-analytische, 197
 - schmale, 263
 - schwach kompakte, 91, 261
 - selbstadjungierte, 139
 - semilineare, 42, 63, 138, 140, 152, 222
 - Ψ -semilineare, 41, 112
 - sesquilineare, 63
 - sestertilineare, 222
 - stark Radon-Nikodým, 108, 261
 - stetige, 19, 23, 91, 106, 212
 - streng ordnungserhaltende, 11
 - surjektive, 3
 - trilineare, 62
 - unitale, 34, 35, 137
 - unterhalbstetige, 20, 21, 120, 121, 203, 264
 - $(\tau_X - \overline{\tau_Y})$ -unterhalbstetige, 95, 203
 - $(\tau - \overline{d})$ -unterhalbstetige, 95
- Abbildungen, zueinander duale, 92
- abelsch, 28, 193
- abgeschlossene Hülle, 18
- Ableitung, 18, 199
- Abschluss, 18
- Abschlussoperator, 14
- Abschnittsfilter, 11
- absolut konvergent, 81
- absolut summierbar, 81
- Absolutbetrag, 170
- absolutkonvex, 59
- absolutkonvexe Hülle, 59
- Absolutpolare, 96
- absorbierend, 59
- absorbiert, 59
- Abstand (zweier Mengen), 37
- abzählbar kompakt, 24
- Abzählbarkeitsaxiom, erstes, 21
- Abzählbarkeitsaxiom, zweites, 21
- ACC, 8
- accessible point, 127
- accumulation point, 18
- Ad (adjungierte Darstellung), 213
- ad (adjungierte Abbildung), 78
- adherence point, 18
- Adjungierte, 78, 139
 - Hilbertraum-, 152
- adjungierte Abbildung, 78
- adjungierte Darstellung, 213
- Adjungierung, 143
- affiner Untervektorraum, 66
- Aktion einer Gruppe, 212
- Algebra, 70
 - alternative, 73, 75, 151
 - Banach-, 70
 - halbnormierbare, 70
 - hermitesche, 147
 - Involution einer, 143
 - involutive, 143
 - Jordan-, 74, 226
 - komplex strukturierbare, 114
 - Komplexifizierung e. norm., 118
 - nichtassoziative, 70
 - normierbare, 70
 - normierte, 70
 - *-normierte, 145
 - symmetrische, 143, 147, 148, 162–164, 166
 - topologische, 109, 179
 - unitale, 70
 - Vidav-, 132, 167
 - vom komplexen Typus, 114
 - vom reellen Typus, 114
 - vom stark reellen Typus, 114, 137
- Algebra-Ideal, 71

- Algebrahalbnorm, 70, 149
 algebraisch gesättigt, 66
 algebraisch irreduzibel, 111, 165
 algebraische Struktur, 28
 Algebranorm, 70, 119
 analytische Automorphismengruppe, 204
 analytischer Fluss, 208
 analytisches Vektorfeld, 206
 Anfangsmenge, 156
 Anfangsprojektion, 157, 164, 168, 252
 Annihilator, 30, 89, 99
 annihilatorabgeschlossen, 31
 annihilatorabgeschlossene Hülle, 31
 Annihilatorsystem, 29, 32, 42, 63–65, 77, 89, 124, 126, 142, 165, 209–211, 243, 254
 Annulator, 30
 Anti-Kommutativität, 72
 Anti-Kommutator-Produkt, 74
 Antiautomorphismus, 35, 142
 Antidarstellung, 110
 Antihomomorphismus, 34, 110
 antimultiplikativ, 34
 Antiordnungsisomorphismus, 11, 127, 253
 antiselbstadjungiert, 139
 antisymmetrisch, 5
 Antiverband, 7, 171, 185
 approximative Eins, 75, 172, 173, 182, 230
 äquilibriert, 59
 äquivalente Abbildungen, 204
 äquivalente Darstellungen, 111
 äquivalente Projektionen, 189
 Äquivalenz zweier äquiv. Darst., 111
 Äquivalenzrelation, 5, 47, 53, 157, 189, 209
 artinsch, 9, 51, 71
 Asplundraum, 201
 assoziativ, 28
 Assoziator, 38
 Astriktion, 4, 84, 158
 asymmetrisch, 5
 Atlas, 204
 Atom (einer C^* -Algebra), 173, 194, 196, 252, 265, 268
 Atom (eines JBW^* -Tripels), 253, 268
 Atom (für ein Maß), 107, 261
 atomar, 257
 atomare Zerlegung, 255, 257
 atomloses Maß, 107
 aufsteigende Ketten-Bedingung, 8
 aufwärts filtrierend, 5
 ausgeartet, 30
 ausgeglichen, 60
 Auswahlaxiom, 5, 7
 Auswertungsabbildung, 61, 208, 209, 217
 lineare, 205
 mit einer Darstellung einer Gruppe assoziierte, 212
 Auswertungsfunktional, 61, 82, 103, 137
 $Aut(M)$, 31, 109, 204
 $aut(M)$, 109, 208
 Automorphismus, 31
 Algebra-, 109
 induzierter innerer, 212
 axiom of restriction, 7
 Bahn, 209
 balanced, 59
 Banach-Jordan-Algebra, 74, 234, 239, 244, 246
 Banach-Jordan-Tripel, 223, 246
 allgemeines, 223
 Banach-Lie-Algebra, 72, 109
 Banach-Lie-Gruppe, 208
 Banach-Mannigfaltigkeit, 204
 normierte, 215
 Banach-* -Algebra, 145
 Banach-Tripel, 223
 allgemeines, 223
 Banachalgebra, 70
 involutive, 145
 *-normierte involutive, 145
 Banachraum, 81
 base norm spaces, 89
 Basis, 44
 Basis einer Topologie, 21
 Basispunkt, 217
 Berührungspunkt, 18
 Bergmann-Metrik, 219
 Bergmann-Operator, 223, 232
 berühren, 198
 beschränkt (Abbildung), 87, 99, 104, 115, 186, 259
 beschränkt (Gebiet), 214–217, 219–222, 227, 234–236
 beschränkt (Netz), 10, 11, 60, 61, 75, 76, 175, 177, 185, 187, 188, 230
 beschränkt (Topol. Vektorraum), 104–109, 125, 128, 129, 171, 176, 186, 197
 beschränkt (Vektorfeld), 206
 $\beta(X, Y)$ (starke Topologie), 105
 Betrag, 149
 Betragsstützpunkt, 124
 Betragsstützfunktional, 124
 beulbar, 106, 269
 Beulpunkt, 106, 107, 128
 bianalytisch äquivalent, 197
 Biball, 218
 bikontraktiv, 101, 141
 Bild, 3
 Bildprojektion, 155
 bilineare Abbildung
 zu einer Darstellung assoziierte, 110
 Bilinearform, 62, 67, 119
 kanonische, 63
 Bilinearsystem, 63, 110
 Bipolarensatz, 97, 183
 Bogenlänge, 215
 Box-Operator, 224

- $\mathbb{C}$ (komplexe Zahlen), 36
- c (konvergente Folgen), 88, 260
- $C(X, Y)$, 87
- $C^b(X, Y)$, 87
- c_0 (Nullfolgen), 70, 82, 88, 103, 174, 260, 269
- $C_0(X, \mathbb{K})$, 87
- Cantor-Bernstein-Eigenschaft, 190
- Carathéodory-Norm, 219
- Cartan-Faktor, 244, 257
- Cauchy-Filter, 80
- Cauchy-Folge, 80
- Cauchy-Netz, 80
- Cauchy-Riemann-Konzept, 204
- Cayley-Dickson, 151
- circled, 59
- $\text{cl}(\tau; A)$, 18
- $\text{cl}(s^*o; A)$, 178
- $\text{cl}(so; A)$, 177
- $\text{cl}(uso; A)$, 177
- $\text{cl}(uwo; A)$, 176
- $\text{cl}(wo; A)$, 175
- cluster point, 18
- Co-Isometrie, 164
- conditionally compact, 24
- $C_p(H_1, H_2)$, 173
- C^* -Algebra, 160, 172, 173, 193, 229, 231–233, 238, 239, 251, 252, 265
- C^* -Algebra von Operatoren, 160
- C^* -Algebranorm, 161
 - schwache, 161
- C^* -Bedingung, 161
 - schwache, 161
- C^* -Bedingung für Tripel, 229
- C^* -Norm, 161
 - schwache, 161
- C^* -Tripel, 229
- $D(\cdot, \cdot)$, 223
- Darstellung, 110
 - adjungierte, 213
 - analytische, 212
 - nicht in X ausgearbeitete, 111
 - stetige, 212
 - treue, 212
- Darstellung einer Gruppe, 212
- Darstellungssatz
 - Fréchet-Riesz'scher, 152
- Daugavet-Eigenschaft, 260
- Daugavet-Gleichung, 260
- Daugavet-Operatoren, 260
- Daugavetraum, 260
- DCC, 8
- dentable, 106
- $\text{Der}(A)$ (Derivationen von A), 78
- Derivation, 78
 - Tripel-, 224, 225, 233, 242
- Derivationen-Algebra von A , 78
- Derivierte, 18
- Diagonale, 4
- dicht, 92, 98, 99, 111, 159, 160, 167, 180
- dicht (bezüglich), 18
- Differential, 205, 208, 212
- R -differenzierbar, 198
- Dimension, 44, 152, *siehe* Hilbertraumdimension
- dimensioniert, 44
- dimensionierter Ring, 44
- Distanzfunktion, 59
- division tripotent, 251
- Divisions-Jordan-Algebra, 74, 251
- Divisionsalgebra, 70, 149
- divisionsminimal tripotent, 251, 252, 255, 258, 259, 265, 268
- Divisionsring, 36
- Divisorenring, 36
- Dixmier-Projektion, 101
- domain of integrity, 36
- 3-Kugel-Eigenschaft, 103
- Dual bezüglich einer Bilinearform, 93
- duales Paar, 65
- Dualität, 65, 99
- Dualraum
 - algebraischer, 61
 - eines Produktes, 95
 - stetiger, 82
- Dualsystem, 63
 - Produkt mehrerer, 94
 - stetiges, 93
- Durchmesser
 - beulbar, 106
 - Beulpunkt, 107
 - C^* -Algebra, 173
 - Daugaveträume, 262
 - ℓ_∞ , 269
 - Prädual, 268
 - δ -rau, 202
 - stark exponierter Punkt, 128
- duxial, 83
- e (Eins), 32
- e (neutrale Element; Eins), iii, 28
- echte Teilmenge, 3
- echter Kegel, 60
- echtes Idempotent, 33
- effektiver Definitionsbereich, 119
- $\text{Eig}(T, \lambda)$ (Eigenraum), 62
- Eigenraum, 62
- Eigenvektor, 62
- Eigenwert, 62
- Einbettung, 61
- eindeutige Hahn-Banach-Fortsetzung, 129, 232
- eindeutiges Prädual, 121
- einfach, 41, 70
- einfach zusammenhängend, 27
- Einheit, 35
- Einheitengruppe, 35

- Einheitskugel, abgeschlossene, 58
- Einheitskugelannihilatorsystem, 124
- Einheitssphäre, 58
- Einparametergruppe, 208
- Eins, 28, 32, 45, 76, 142–147, 161, 165, 168, 184
siehe auch approximative Eins, 75
- Elementarfilter, 11
- endlich, 4
- endlich erzeugt (Ideal), 34, 54
- Endmenge, 156
- Endomorphismus, 31
- Endprojektion, 157, 164, 168
- Epigraph, 37, 119–121
- Epimorphismus, 31
- ε -Beulpunkt, 106
- ε -schwach*-Beulpunkt, 106
- equilibrated, 59
- erstes Element, 6
- erweiterte reelle Zahlen, 36
- erweiterte reelle Zahlengerade, 36
- Erweiterung, 76
- erzeugend, 60
- Erzeugendensystem, 42
- $\text{ex}_{\mathbb{C}}(K)$, 69
- $\text{ex}_{\mathbb{R}}(K)$, 69
- $\text{ex}(K)$, 69
- $\text{Expfaces}(\tau; C)$, 126
- $\text{expo}(A)$, 127
- Exponential, 208
- Exponentialabbildung, 208, 209
- τ -exponierte Seite, 126
- exponierter Punkt, 127, 128
- extremale Teilmenge, 69
- Extremalpunkt, 69, 89, 106, 107, 128, 160, 167, 168, 248, 253, 264, 268, 269
 komplexer, 69, 248
 reeller, 69, 248
- exzeptionelle Jordan-Algebra, 75
- $F(X, Y)$ (lin. Abb. endl. Ranges), 62, 82, 83
- $\mathcal{F}(X, Y)$ (stet. lin. Abb. endl. Ranges), 82, 83
- facear operation, 127
- $\text{Faces}(\tau; C)$, 126
- Faktor, 184, 185, 192, 193, 195, 196, 265–267
- Faktor (JBW^* -Tripel), 257
- Faktor vom Typ I (JBW^* -Tripel), 257
- fallende Ketten-Bedingung, 8
- Faserung, 5
- fast alle, 4
- Filter
 - feinerer, 11
 - konvergenter, 24
 - Limespunkt, 24
- Finaltopologie, 24
- Fixgruppe, 210
- Fixpunkt, 139
 isolierter, 216
- Fixpunktkörper, 210
- Fixpunktmenge (einer Abbildung), i, 4, 39, 76, 139
- Fixpunktmenge (Fixsystem), 210
- Fixsystem, 210
- Fixsystem einer auf einer Gruppe operierenden Menge, 209
- Fluss
 - analytischer, 208
 - lokaler analytischer, 207
- Folgenraum, 88, 260, 269
- folgenvollständig, 80
- Form
 - assoz. quadratische, 67
 - quadratische, 67, 150, 151
 - reelle, 116, 141, 240, 242
- Fortsetzung, kanonische komplex-lin., 117
- Fréchet-Ableitung, 198, 199
- Fréchet-differenzierbar, 198
- Fréchet-Raum, 82, 84
- Fréchet-Riesz, 186
- Fréchet-Riesz'scher Darstellungssatz, 152
- frei, 44
- freie Gruppe, 212
- freier Modul, 43
- Frobenius, 149
- F_{σ} -Menge, 18
- Fundamentalformel, 225
- fundierte Menge, 8
- Fundierungsaxiom, 1, 2, 7
- 5-lineare Gleichung, 223
- Funktion, im Unendl. verschwindende, 87
- Funktional, Norm-annahmendes, 124
- Funktionenräume, 86
- $\mathfrak{g}(G)$ (Lie-Alg. e. Banach-Lie-Grp.), 208
- $G(\cdot)$ (Menge d. invertierb. El.), 33
- $G(\cdot)$ (Menge der Inversen), 35, 144, 197
- Gâteaux-Ableitung, 200
 komplexe, 201
- Gâteaux-differenzierbar, 199, 200
- Gâteaux-holomorph, 201
- Gâteaux-Variation, 199
- Galois(unter)körper, 211
- Galois-Verbindung, 12
 antitone, 12
 monotone, 12
- Galoisgruppe, 210
- galoissch, 211
- Galoissystem, 210
- G_{δ} -Menge, 18
- Gebiet eines Banachraumes, 197
- Gelfand-Homomorphismus, 137
- Gelfand-Naimark-Segal-Konstrukt., 165
- Gelfand-Raum, 137
 erweiterter, 137
- Gelfand-Topologie, 137
- Gelfand-Transformierte, 137
- generalized sequence, 10

- geordnet (*-Algebra), 172
 geordneter Vektorraum, 58
 gerichtete Familie, 10
 gerichtete Menge, 5
 gesättigt, 66
 getrennt stetig, 109
 $GL(X)$ (in $L(X)$ inv.-bare El.), 109
 $\mathfrak{gl}(X)$, 109
 glatte Norm, 124, 202
 glatter Punkt, 124, 127
 gleichgradig stetig in e. Punkt, 108
 gleichstetig, 108, 109
 $\mathcal{GL}(X)$ (in $\mathcal{L}(X)$ inv.-bare El.), 109
 $G^q(\cdot)$ (Menge der Quasi-Inversen), 33, 35, 56
 R -Gradient, 198
 Gradientenabbildung, 63, 152
 rechte, 63
 Graduierung
 additive, 73
 multiplikative, 73
 Peirce-, 245
 Graph einer Abbildung, 5
 Grenzwert, 21
 größtes Element, 6
 Gruppe, 29
 Darstellung einer, 212
 Gruppenisomorphismus, 35
 Gruppoid, 28
 gutes Produkt, 167

 $\mathbb{H}$ (Quaternionen), 149
 Häufungspunkt, 18
 Hüllenoperator, 14
 halbeinfach, 50, 51
 Halbgruppe, 28
 Halbnorm, 57, 58, 91, 119
 Algebra-, 70
 spektrale, 229
 Halbordnung, 5
 halbprim, 45, 172
 halbprimitiv, 49
 halbproper (*-Algebra), 171, 172
 Halbraum, 66
 Hamilton-Algebra der Quaternionen, 149
 Harish-Chandra-Realisierung, 219, 221
 Hauptgleichung, 223
 Hauptideal, 34
 Hauptidealbereich, 36
 Hauptidealring, 36
 Hausdorffraum, 25–27, 64, 80, 82, 90, 91, 109, 161, 163, 175
 Hermitefizierung, 242
 $\bar{\cdot}$ hermitesch, 227
 $^+$ hermitesch, 228
 hermitesch (Banach-Jordan-Tripel), 228
 hermitesch (Element e.unit. Alg. ü. $\mathbb{C}$), 131, 132, 135, 136, 140, 145, 166, 167, 226–229
 hermitesch (Involution; Algebra), 147, 148, 162
 hermitesch (Metrik), 219
 hermitesch (nach Vidav), 132, 167
 hermitesch (sesquilin. Abb.), 63, 142, 227
 1-hermitesch, 131
 Hermitezität (Ban.-Jord.-Trip. ü. $\mathbb{C}$), 228
 Hermitezität (norm. Vektorräume), 131
 Hilbert'sche direkte Summe, 154, 165
 Hilbert-Schmidt-Norm, 174
 Hilbert-Schmidt-Operator, 174
 Hilbertraum, 81, 152, 232, 242, 248
 Hilbertraumadjungierte, 152, 179–181, 232
 Hilbertraumdimension, 152
 holomorph, 199
 homöomorph, 19
 Homöomorphismus, 19, 58, 102, 122
 homogen, 211
 Homomorphiesatz für Ringe, 47
 Homomorphismus, 31, 34, 41, 58, 79, 110, 112
 Anti-, 34
 Gruppen-, 41
 Tripel-, 223
 homotop, 26
 Homotopie, 26
 Homotopieformel, 225
 Hyperebene, 66

 ICC-Gruppe, 212, 266
 Id_A , 4
 Ideal, 34, 38, 71, 223
 beidseitiges, 34, 71
 einer Menge, 211
 endlich erzeugtes, 34, 54
 inneres, 223, 224, 232, 246
 Idem($\cdot$), 32
 Idempotent, 32
 idempotent, 32, 39, 47–49, 52, 110, 168, 223
 Identitätssatz, 215
 Im (Imaginärteilbldg. Kplxfg.), 116, 139
 Im (Imaginärteilbldg.), 116, 139
 imitieren eines Raumes, 82
 Immersion, 214
 Index, 10
 Induktionsbedingung, 9
 induktiv geordnet, 7
 Infimum, 7, 16, 17, 36, 155, 187
 infinitesimale Automorphismen, 208
 infinitesimale Drehung, 220
 infinitesimale Transformation, 206
 infinitesimale Transvektion, 220
 infinitesimaler Generator, 208
 Initialtopologie, 24
 Injektion, 3
 Inklusionsabbildung
 kanon., 61, 82, 101
 Innere, 18
 innerer Punkt, 18
 Int (Konjugation), 212
 integral domain, 36

- integrierbares Vektorfeld, 208
- Integritätsbereich, 36
- interior, 18
- Intervall, 68
- intrinsische Norm, 162
- invariant, 110
- Inverse, 29, 33
- invertierbar, 28, 33, 74
- Involution, 143
 - gespiegelte, 139
 - kanonische, 143, 145, 169
 - kartesische, 140, 141, 144, 163, 167, 227
 - Vidav-, 145
- Involution einer Algebra, 143
- Involution eines Ringes, 142
- Involution eines Vektorraumes, 138
- involutive Algebra, 143
- involutive Banachalgebra, 145
- irreduzibel, 78, 111, 165
- irreduzibel idempotent, 49
- irreduzibel tripotent, 250
- irreflexiv, 4
- isolierter Punkt, 19, 27
- Isometrie, 58, 83, 94, 133, 162, 164, 214, 220, 221, 231–234
 - partielle, 156, 157, 164, 165, 168, 189, 195
- isometrisch, 58
- isometrisch komplex strukturierbar, 114
- Isomorphie, isometrische, 76
- Isomorphismus, 31, 58, 94
 - Tripel-, 223
- Isotropie-Unteralgebra, 221
- Isotropie-Untergruppe, 210, 219, 236
- Jacobi-Gleichung, 72, 236
- Jacobson-halbeinfach, 56, 78, 160, 171, 172
- Jacobson-Radikal, 56, 57, 112, 136, 144, 160
- James-Raum, 115
- JB -Raum, 234
- JB^* -Algebra, 232–234, 246
- JB^* -Tripel, 141, 184, 229, 257
 - reelle Analoga, 239
- JBW^* -Algebra, 232
- JBW^* -Tripel, 238, 242
- JC^* -Algebra, 233
- JC^* -Tripel, 232
- Jordan C^* -Algebra, 232
- Jordan W^* -Algebra, 232
- Jordan-Algebra, 74, 226, 245, 246
- Jordan-Banachraum, 234
- Jordan-Funktor, 235
- Jordan-Gleichung, 223
- Jordan-Multiplikationsoperator, 224
- Jordan-Paar, 224
- Jordan-Tripel, 223
 - allgemeines, 223
 - lineares, 223
- Jordan-Tripel-Gleichung, 223
- Jordan-Tripelprodukt, 223, 225
 - partiell, 235
 - symmetrisiertes, 243
- Jordanprodukt
 - komplexes, 144
 - reelles, 144, 233
- J^* -Algebren, 232
- J^* -Ideal, 229
- J^* -Tripel, 228, 229
- J^* -Tripel
 - lineares, 228, 229
- J^*B -Tripel, 238
- $\mathbb{K}$, 36
- $K(H)$, 103, 123, 164, 173, 174, 176
- $K(X, Y)$ (kompakte lin. Abb.), 83
- $K_w(X, Y)$ (schwach komp. lin. Abb.), 91
- Körper, 210
- Kalotte, 66
- kanonische Bilinearform, 63
- kanonische Inklusionsabbildung, 82, 101
- kanonische Involution, 143, 145, 169
- kanonische komplex-lin. Fortsetzung, 117
- kanonische Projektion, 101
- kanonische Quotientenabbildung, 82
- Kaplansky, Formel von, 189
- Karte, 204
- kartesische Involution, 140–142, 144, 163, 167, 227, 256
- Kegel, 60, 146, 169, 208
- Keil, 60, 79
- Keim, 204
- Kern, 18, 31, 35, 62, 112
- Kernprojektion, 155
- Kette, 5
- kleiner als relativ zu, 190
- kleinstes Element, 6
- KMP (Krein-Milman-Eigenschaft), 128
- kofinal, 5
- Kommutante, 32, 77, 110, 148, 212
- kommutativ, 28, 32, 77, 136, 171, 185, 193, 195, 223
- kommutativ konvergent, 81
- kommutatives Diagramm, 109
- Kommutator, 38, 77, 206
- kompakt, 24
- Kompaktifizierung, 137
- kompatibel, 204, 214
- Kompatibilitätseigenschaft, 69
- Komplement, 84
 - topologisches, 84
- komplex strukturierbar, 114, 242
- komplexe Struktur, 114, 115
- komplexe Strukturierung, 114
- komplexer Typus, 114
- Komplexifizierung, 116, 141, 240–242
- Komplexifizierung e. norm. Algebra, 118
- Komplexifizierung nach LI, 141, 142, 144, 164, 241

- Komplexprodukt, 28
 Komposition, erlaubte, 150
 Kompositions-Algebra, 151
 Kompression, 158
 Kondensationspunkt, 18
 Konjugation, 141, 153, 212, 240
 Konjugation e. norm. Vektorraumes, 140
 Konjugationsklasse, 212
 konjugiert-lin. Tripel-Homomorph., 223
 konjugiert-linear, 42, 224, 226, 237, 239–241
 konjugierte Quaternion, 149
 Kontraktion, 101, 214
 kontraktiv, 101, 123, 141, 162, 235
 konvergente Reihe, 81
 Konvergenz bzgl. e. Filters, 24
 Konvergenzradius, 197
 konvex, 59, 60
 konvexe Funktion, 119
 konvexe Hülle, 59
 Körper, 36
 K_p (Isotropie-Untergruppe), 219, 236
 Krein-Milman-Eigenschaft, 128, 257
 kreisförmig, 59
 kreisförmige Hülle, 59
 Kreuzprodukt, 72
 Kronecker-Delta, 33
 kugelabgeschlossen, 124
 Kugelannihilator, 124
 Kugelannihilatorsystem, 124
 kugelorthogonal, 124
 Kugelsystem, 124

 $L(X, Y)$ (lineare Abbildungen), 61
 $L^n(X, Y)$ (n -lineare Abbildungen), 62
 $\mathcal{L}(X, Y)$ (stetige lin. Abbildungen), 82
 $\mathcal{L}^n(X, Y)$ (stetige n -lin. Abbildungen), 83
 $\mathcal{L}_p(X, \Sigma, \mu)$, 87
 $L_p(X, \Sigma, \mu)$, 87, 103
 $L_{p(\mathbb{C})}(X, \Sigma, \mu)$, 87
 $\mathcal{L}_{p(\mathbb{C})}(X, \Sigma, \mu)$, 87
 $L(X)$, 62, 78
 L -eingebetteter Raum, 103, 244, 264, 268
 L -Projektion, 102
 L -Summand, 102, 244, 255
 L_a (Links-Multiplikation), 35
 lattice, 7
 LBP, 235
 leere Indexmenge, 3, 29
 Lemma von Schwarz, 215
 Lie-Algebra, 72
 assoziierte, 72, 77
 Lie-Algebra e. Banach-Lie-Gruppe, 208
 Lie-Tripel, 235
 Lie-Tripel-Produkt, 236
 Liften eines Moduls, 117
 Limes inferior, 37
 Limes superior, 37
 limit points, 18

 lin. bihl. Eigenschaft, 235
 Lindelöf-Raum, 24
 linear unabhängig, 44
 R -linear unabhängig, 42
 lineare Hülle, 42
 lineares Jordan-Tripel, 223
 lineares Tripel, 222
 linke Gradientenabbildung, 63
 links, 63
 links ausgeartet, 29
 links-invertierbar, 28, 33
 Links-Multiplikation, 35, 76, 114, 130, 165, 226
 links-quasi-invertierbar, 33
 links-quasi-regulär, 55
 links-reguläre Darstellung, 111
 Links-Translation, 208
 Linksadjungierte (Galois), 12
 Linksannihilator, 30, 224
 linksartinsch, 41, 51, 71
 Linkseins, 28, 32
 Linkseins modulo e. Rechtsideals, 55
 Linksideal, 34, 71
 Linksinverse, 28, 33
 Linksmodul, 40
 linksneutrales Element, 28
 linksnoethersch, 41, 51, 54, 71
 Linksnulleiler, 33
 Linkssockel, 56
 Linksteiler, 33
 Linksträger, 155, 158
 lokal dicht, 19
 lokal kompakt, 24
 lokal konvex, 79, 82, 109
 lokal linearisierbar, 213
 lokal wegzusammenhängend, 26
 lokal zusammenhängend, 26
 lokaler analytischer Fluss, 207
 ℓ_p , 88, 118, 174, 260, 269
 $L_p(\mu)$, 70, 84, 141, 174
 LUMER-Produkt, 130, 131
 Norm bestimmendes, 131

 $M_{\mathbb{K}}(n)$, 149
 M^m , M_m , 10
 M^m , M_m , 187
 M -eingebetteter Raum, 103, 193
 M -Ideal, 102, 261
 M -Projektion, 102
 M -Summand, 102, 238, 243, 247, 256
 Mackey-Topologie, 105, 178
 Magma, 28
 Majorante, 7
 masa, 266
 Maßtheorie, 87
 Matrizenring, 34, 46, 53–55, 72, 148, 149
 maximal kommutativ, 32, 77, 110
 Maximal-Prinzip, 7
 Maximalbedingung, 8

- maximales Element, 6
- maximales Ideal, 52, 71, 144
- Metrik, 37
 - hermitesche, 219
 - Riemann'sche, 219
- metrischer Raum, 37
- minimal idempotent, 49, 52, 87, 193, 250
- minimal tripotent, 250–252, 254, 255, 258, 268
- Minimalbedingung, 8
- minimales Element, 6
- minimales Ideal, 52, 71, 83, 150
- Minkowski-Funktional, 59, 118
- Minkowski-Norm, 118
- Minorante, 6
- mnemotechnische Eins, 33
- Modul, 40
- modular, 55, 144
- modulus support functional, *siehe* Betragsstütz-funktional
- Monoid, 29
- Monomorphismus, 31, 94
- monoton fallend vollständig, 10
- monoton steigend vollständig, 10
- monoton vollständig, 61, 184, 186, 187
- monoton vollständig nach KRT, 61
- Moore-Smith-Folge, 10
- Morphismus, 217
- Multiplikation
 - a -, 223, 226, 234, 239
 - getrennt stetige, 109
- Multiplikationsoperator
 - Jordan-, 224
- Murray-von Neumann-Äquivalenz, 189
- $\mathbb{N}$ (natürliche Zahlen), 3
- narrow operator, 263
- Netz, 10
 - feineres, 10
 - Konvergenz, 24
 - monoton fallend, 10
 - monoton steigend, 10
 - nach oben beschränkt, 10
 - nach oben ordnungsbeschränkt, 10
 - nach unten beschränkt, 10
 - nach unten ordnungsbeschränkt, 10
 - streng monoton fallend, 10
 - streng monoton steigend, 10
- von Neumann-Algebra, 179, 202
 - atomare, 195
 - diskrete, 193
 - echt unendliche, 192
 - endliche, 192
 - halbdiskrete, 194
 - halbendliche, 193
 - halbkontinuierliche, 194
 - kontinuierliche, 194
 - rein unendliche, 192
 - Typ H, 196
 - Typ I, 193, 195, 267
 - Typ I_∞ , 195, 267
 - Typ I_{fin} , 195
 - Typ I_n , 267
 - Typ II, 194, 195, 265, 267
 - Typ II_1 , 195, 266, 267
 - Typ II_∞ , 195, 267
 - Typ III, 192, 195, 265, 267
 - unendliche, 192
- von Neum.-Alg. $\ddot{u}.\mathbb{R}$ n.STØRMER-AYUPOV, 180
- neutrales Element, 28
- nicht ausgeartet, 63, 111
- nicht in X ausgeartet, 111
- nicht links ausgeartet, 63
- nicht rechts ausgeartet, 63
- nicht trennend, 30
- nil, 45
- nil-Radikal, 51, 52, 54, 56
- niles Ideal, 45
- nilpotent, 45, 51, 52, 56
- nirgends dicht, 19
- noethersch, 9, 51, 54, 71
- Norm, 58, 214
 - δ -raue, 202
 - äquivalente, 118
 - Algebra-, 70, 119
 - extrem raue, 202, 262
 - Fréchet-Differenzierbarkeit, 202
 - Fréchet-glatte, 108, 202
 - Gâteaux-Differenzierbarkeit, 202
 - glatte, 124, 202
 - intrinsische, 162
 - Minkowski-, 118
 - raue, 202
 - reguläre, 76
 - stark subdifferenzierbare, 258
 - Unterhalbstetigkeit, 203
- Norm-1-komplementiert, 85
- Norm-annehmend, 126
- Norm-bestimmend, 83
- Norm-Topologie, 58
- normal, 144, 166
- normale Linearform, 188, 233
- normaler Raum, 26
- Normalisator, 212
- Normalteiler, 29
- normierend, 83, 98
- normiert, 111
- normierte *-Algebra, 145
- p -Norm, 88
- nuklear, 174
- nuklearer Operator, 174
- nuklearer Raum, 82
- nullhomotop, 27
- Nullteiler, 33
- nullteilerfrei, 33, 36
- numerischer Radius, 130
- numerischer Wertebereich, 130, 242

- bzgl. e. Semiskalarpr., 130
 obere Schranke, 7
 Operation e. Gruppe auf e. Menge, 209
 Operator, 37
 K -Operator-Ring, 37
 Operatorenbereich, 28, 37
 $\text{OProj}(H)$ (Orthog.proj. von $\mathcal{L}(H)$), 155
 Orbit, 209
 ordnungsbeschränkt, 10, 171
 Ordnungsintervall, 58, 171
 Ordnungsisomorphismus, 11, 253
 orthogonal, 30–32, 77, 89, 165, 224, 229, 247
 orthogonal (Funktionale), 255
 orthogonal (Operator), 157
 orthogonalabgeschlossen, 31, 65–67, 93
 orthogonalabgeschlossene Hülle, 31
 Orthogonalmenge, 31
 Orthogonalprojektion, 154, 155
 Orthogonalraum, 64
 Orthonormalbasis, 152
 Orthonormalsystem, 152
 orthosymmetrisch, 30, 77, 224, 228, 247

 $\mathbb{P}$ (L oder M), 102
 $P^n(X, Y)$ (n -homogene Polynome), 67
 $\mathcal{P}^n(X, Y)$ (stet. n -homog. Polyn.), 83
 $\mathcal{P}(M)$ (Potenzmenge der Menge M), 3
 $\mathbb{P}$ -eingebetteter Raum, 103
 $\mathbb{P}$ -Summand, 102
 Paarung, 65
 Parallelotop, 24
 Parseval'sche Gleichung, 153
 partielle Fréchet-Ableitung, 199
 partielle Ordnung, 5
 $\text{pcran}(T)$ (Bildproj. von T), 155
 Peirce
 -Raum, nicht-diagonaler, 75
 -Unteralgebra, diagonale, 75
 Peirce-Raum, 245, 248, 249
 Peirce-Zerlegung, 40, 48, 73, 75, 245
 Permutation, 4, 39, 81
 Permutationsoperator, 62
 $\text{pker}(T)$ (Kernproj. von T), 155
 Polare, 96
 Polarisationsgleichung, 68
 Polarisierung, 68
 Polynom
 assoziertes n -homogenes, 67
 n -homogenes, 67
 stetiges n -homogenes, 83
 polynomiales Vektorfeld, 207
 Polynomring, 54, 211
 Polyradius, 218
 Polyscheibe, 217
 Polyzyylinder, 218
 $\text{Pos}(\text{idem}; \cdot)$, 79
 $\text{Pos}(\cdot; \cdot)$, 79
 $\text{Pos}(\sigma; \cdot)$, 135
 $\text{Pos}(*\sigma; \cdot)$, 144
 $\text{Pos}(*; \cdot)$, 146
 $\text{Pos}(*W; \mathcal{L}(H))$, 154
 $\text{Pos}(V\sigma; \cdot)$, 135, 136, 229
 $\text{Pos}(V; \cdot)$, 132, 136
 $\text{Pos}(W; \mathcal{L}(H))$, 154
 positiv, 36, 227, 229
 echt, 36
 positive Linearform, 146
 Positivität, Konzept der, 79, 144, 146
 Potenzmenge, 3
 Potenzreihe, 197
 Prädual, 121, 128, 142, 201, 202, 232, 237, 242, 244, 251, 253, 257, 258, 261, 268
 eindeutiges, 121
 stark eindeutiges, 121
 prägeordneter Vektorraum, 58
 Prähilbertraum, 64
 Präordnung, 5, 47
 pre-facear operation, 127
 Primelement, 36
 Primideal, 45
 primitiv idempotent, 49
 Primring, 46, 171
 principal ideal domain, 36
 principal ideal ring, 36
 Prinzip d. gleichmäß. Beschränktht., *siehe* Satz von Banach-Steinhaus
 Prinzip der lokalen Reflexivität, 237
 Produkt
 direktes, 42, 61, 116
 inneres direktes, 32
 kartesisches, 3–5, 24, 42
 verschränktes, 266
 Produkttopologie, 24, 89, 118
 Produkttopologie auf $\bigoplus_{\nu \in I} X_\nu$, 85
 $\text{Proj}(A)$ (Proj. e. $*$ -Algebra A), 144, 155
 Projektion
 A_p , 183
 C^* -Norm, 161
 L - und M -, 102
 $L(X)$, 78
 $*$ -Algebra, 144
 $*$ -Darstellung, 160
 äquivalente, 189
 abelsche, 193
 atomare, 173, 265, 268
 auf $\text{cl}(AH)$, 183
 bikontraktive, 101, 141
 diskrete, 193
 echt unendliche, 192
 Eins, 144
 endliche, 192
 erzeugende, 189
 Extremalpunkt, 168
 halbabelsche, 194
 halbkontinuierliche, 194
 Hilbertraum, 155

- innere topol. Summe, 84
- komplementäre, 78
- Kompression, 158
- kontinuierliche, 194
- kontraktive, 101, 123, 237
- rein unendliche, 192
- Tangentialbündel, 205
- treue, 193
- unendliche, 192
- Verband, 186
- zentrale, 78, 192
- proper
 - *-Algebra, 171, 172
 - Idempotent, 33
 - Teilmenge, 3
- Ψ (Ringhom.), 112
- Pták, 148
- Pullback, 112
- punktierte Menge mit Basispunkt, 3, 29, 31, 32, 40, 58
- punktierter Kegel, 60
- $Q(\cdot, \cdot)$, 223
- r^q (Quasi-Inverse von r), 33
- quadratische Form, 67, 150, 151
 - assoziierte, 67
- Quadratwurzel, 154, 169, 170
- Quasi-Inverse, 33
- quasi-Inverse
 - Links-, 33
 - Rechts-, 33
- quasi-invertierbar, 33, 55, 114, 143
- Quasi-Produkt, 33, 48
- quasi-proper (*-Algebra), 172
- quasi-regulär, 55
- Quasi-Spektrum, 134
- R -Quasigradient, 202
- Quasiordnung, 5
- Quaternion
 - Betrag, 149
 - konjugierte, 149
- Quaternionen, 149
- Quotient, 67, 82, 89, 94, 100, 235
- Quotientenabbildung, kanonische, 82
- Quotiententopologie, 82, 100
- $\mathbb{R}$ (reelle Zahlen), 36
- $\overline{\mathbb{R}}$ (erweiterte reelle Zahlen), 36
- $R(\cdot)$, 63, 142
- R -minimales Element, 7
- R_a (Rechts-Multiplikation), 35
- radial, 59
- Radikal, 50
 - $\mathcal{S}$ -Radikal, 50, 51, 71
- radikal, 50
- Radikal eines Ideals, 46
- Radikal-Eigenschaft, 50, 51, 56, 71
- Radikalring, 50
- Radon-Nikodým-Eigenschaft, 107, 108, 123, 128, 173, 201, 257, 259, 261, 262
- Rand, 18
- Randpunkt, 18
- Rang, 253
- Rauheit, 202
- Re (Funktionale), 113
- Re (Realteilbldg. Kplxfzg.), 116, 139
- Re (Realteilbldg.), 116, 139
- realification, 113
- rechts, 63
- rechts ausgeartet, 29
- rechts-invertierbar, 28, 33
- Rechts-Multiplikation, 35, 76
- rechts-quasi-invertierbar, 33
- rechts-quasi-regulär, 55
- rechts-reguläre Darstellung, 111
- Rechts-Translation, 208
- Rechtsadjungierte (Galois), 12
- Rechtsannihilator, 30, 224
- rechtsartinsch, 51, 71
- Rechtseins, 28, 32
- Rechtsideal, 34, 71
- Rechtsinverse, 28, 33
- Rechtsmodul, 40
- rechtsneutrales Element, 28
- rechtsnoethersch, 51, 71
- Rechtsnullteiler, 33
- Rechtsshiftoperator, 157
- Rechtssockel, 56
- Rechtsteiler, 33
- Rechtsträger, 156, 158
- reduzierende Ideal, 172
- reduzierendes Ideal, 160
- reduziert, 160, 172
- reell, 132, 180
- reelle Form, 116, 141, 240, 242
- reelle Strukturierung, 113
- reeller Ring, 114
- reeller Typus, 114, 119, 135, 137
- reelles JB^* -Tripel, 239
- reflexiv, 4, 100, 126, 173
- regulär, 171, 172, 223, 248
- reguläre Norm, 76
- regulärer Raum, 25
- Regularitätsaxiom, 7
- Reihe, 81
- rein atomares Maß, 107, 261
- Relation, 4
 - duale, 4
- relativ kompakt, 24
- residuation theory, 13
- Restriktion, 4, 158
- Restriktion des Homomorphismus, 112
- Richtung, 5
- Richtungsableitung, 200
- Riemann'sche Metrik, 219
- Riesz, 152

- Riesz'scher Raum, 58
 Ring, 32, 180
 $\mathcal{S}$ - K -Ring, 51
 $\mathcal{S}$ -Ring, 50
 alternativer, 38
 Involution eines, 142
 nichtassoziativer, 32
 nichtassoziativer K -, 37
 nichtassoziativer $\mathcal{S}$ - K -, 51
 selbstadjungierter Teil, 142
 vom komplexen Typus, 114
 Ring mit einem Ring von Operatoren, 38
 Ring mit Operatoren, 37
 Ring mit Operatorenbereich, 37
 Ring-Ideal, 71
 rough norm, 202

 $\sigma(\mathcal{S})$ (Klasse aller $\mathcal{S}$ -Ringe), 50, 51
 $\mathcal{S}$ (Radikal-Eigenschaft), 50
 $\mathfrak{S}_n$ (Permutationsgruppe), 29
 $\mathfrak{S}_\infty$, 29, 212
 $S(\cdot)$, 63, 142
 $\mathfrak{S}$ -Topologie, 104, 105
 Satz von
 Alaoglu-Bourbaki, 100, 109, 122, 123, 126, 203
 Artin, 39
 Banach-Steinhaus, 104, 177
 Banach-Stone (Kauf), 234
 Bishop-Phelps, 125, 129
 Cartan (Eindeutigkeit), 215
 Cartan (Klassifizierung), 219
 Dixmier-Ng, 121
 Eberlein-Šmulian, 173
 Goldstine, 98, 252, 264
 Hahn-Banach, 62, 90, 94, 98–101, 105, 106, 126, 129, 135, 136, 232, 254
 Harish-Chandra (Einbettung), 219
 Hellinger-Toeplitz, 153
 Jordan-Hölder, 53, 54
 Kaplansky (Dichte), 181
 Krein-Milman, 123, 128, 168, 264, 268, 269
 Krein-Smuljan, 101
 Lindenstrauss, 128
 Mackey (beschränkte Mengen), 105
 Mazur, 149, 150
 Riemann (Abbildungen), 217
 Straszewicz-Klee, 128
 Vidav-Palmer, 162, 167, 169
 Vigier, 185, 186
 von Neumann (Dichte), 180
 von Neumann (Doppel-Kommut.), 180
 Wedderburn-Artin, 54
 Satz von der offenen Abbildung, 99
 Schattenklasse, 173
 Scheibe, 106, 173
 schiefhermitesch (Algebra), 147, 148
 schiefhermitesch (El.e.unit.Alg.ü. $\mathbb{C}$), 115, 132
 Schiefkörper, 36, 70, 150
 schiefminimal idempotent, 47, 49, 50, 52, 53, 71, 87, 150, 164, 173, 193–195, 251
 schiefselbstadjungiert, 139
 schwach Hahn-Banach-glatt, 129
 schwache C^* -Algebrannorm, 161
 schwache C^* -Bedingung, 161
 schwache C^* -Norm, 161
 schwache Operator-Topologie, 175
 schwache Topologie, 90
 schwacher Asplundraum, 201
 schwach*-Asplund-Raum, 201
 schwach*-Beulpunkt, 106
 schwach*-Operatortopologie, 175
 schwach*-Scheibe, 106
 schwach*-Topologie, 90, 92
 Segment, 68
 sehr proper ($*$ -Algebra), 171, 172
 Seite, 69, 126, 252, 253
 Selbstabbildung, 4
 selbstadj. Teil e. Vektorraumes, 139
 selbstadjungiert, 139, 142
 τ -semi-exponierte Seite, 127
 Semiexpfacs($\tau; C$), 127
 Semiskalarprodukt, 63, *siehe* LUMER-Prod.
 separabel, 18
 Sesquilinearform, 63
 Sesquilinearsystem, 63
 σ -kompakt, 24
 σ -schwache Operator-Topologie, 176
 σ -starke Operator-Topologie, 177
 σ -stark* Operator-Topologie, 178
 Skalar (Moduln), 40
 Skalarbereich, Einschränkung des, 112
 Skalarenring, 40
 Skalarprodukt, 64, 146
 VR-wertiges, 142
 skew minimal idempotent, 49
 slice, 106
 smooth, 124, 202
 so (starke Op.-Topologie), 176
 $\text{soc}(R)$ (Sockel), 56
 Sockel, 56
 span, 66, 92, 145, 159, 160, 167, 169, 175, 188, 189
 span(S) (erzeugte Untermodul von S), 42
 spektrale Halbnorm, 229
 Spektralradius, 134
 Spektrum, 133
 spezielle Jordan-Algebra, 74
 Spin-Faktor, 244
 spitzer Kegel, 60
 Spur, 146, 152, 174
 Spurabbildung, 152
 Spurform, 227
 Spurfunktional, 174
 Spurklasse-Operatoren, 174
 s^*o (stark* Op.-Topologie), 177

- stärkste Operator-Topologie, 177
- Stützpunkt, 124
- Stabilisator, 210, 212
- Stabilitätsuntergruppe, 210
- Standuntergruppe, 210
- stark eindeutiges Prädual, 237, 238
- stark exponierter Punkt, 128, 202, 258, 261, 268, 269
- stark extremaler Punkt, 107, 128
- stark subdifferenzierbar, 198, 201, 258
- starke Operator-Topologie, 176
- starke Topologie, 105
- stark* Operator-Topologie, 177
- *-Abbildung, 139
- *-Algebra, 143
 - normierte, 145
- *-Antidarstellung, 159
- *-Antihomomorphismus, 159
- *-antisymmetrisch, 139
- *-Banachalgebra, 145
- *-Darstellung, 159
- *-Homomorphismus, 159
- *-Ideal, 144
- *-normierte involutive Banachalgebra, 145
- *-radikal, 160
- *-symmetrisch, 139
- *-Teilmenge, 139, 142
- sternförmig, 69
- *JB-Tripel, 239, 243
- *JBW-Tripel, 242
- stetiges n -homogenes Polynom, 83
- Stetigkeitspunkt, 106, 107
- Strecke, 68
- streng komplexe Struktur, 115, 138
- streng monoton fallend vollständig, 10
- streng monoton steigend vollständig, 10
- strikte partielle Ordnung, 5
- strikte Halbordnung, 5
- strongly subdifferentiable, 198
- strukturierbar, komplex, 114, 242
- Strukturierung
 - komplexe, 114
 - reelle, 113
- Strukturtheorem, 53
- Stützfunktion, 124
 - Halbstetigkeit, 203
- Stützfunktional, 124, 254
- Stützpunkt, 128
- Subbasis einer Topologie, 22, 23
- Subdarstellung, 111
- Subdifferential, 125
- subdifferenzierbar, 198
- R -subdifferenzierbar, 197
- R -Subgradient, 197
- Submersion, 213, 217
- Summe, 81
 - algebraische, 42
 - äußere direkte, 42
 - direkte, 35, 42
 - direkte algebraische, 42
 - erzeugte (Ideale), 34
 - innere direkte, 32, 35, 84
 - innere direkte algebraische, 42, 67
 - innere topologische direkte, 84
 - lokal konvexe direkte, 86
 - topologische direkte, 84, 95
- summierbar, 79
- Support, *siehe* Stütz..., *siehe* Träger (eines Operators)
- Supremum, 7, 16, 17, 36, 155, 187
- Surjektion, 3
- Symmetrie (Mannigfaltigkeit), 216, 217, 234, 236
- symmetrisch (n -lin. Abb.), 62, 63, 67, 68
- symmetrisch (*-Algebra), 143, 147, 148, 162–164, 166
- symmetrisch (Mannigfaltigkeit), 216, 217, 219, 221, 222, 227, 228, 234–236
- symmetrisch (Relation), 5
- symmetrische Gruppe, 29
- symmetrischer Teil, 221, 235
- System, 29
- $\mathcal{T}(M)$ (VR aller analyt. VF auf M), 206
- T_0 -Raum, 25, 109
- T_1 -Raum, 25, 37, 109
- T_2 -Raum, 25, 37, 109
- T_3 -Raum, 25, 37
- T_{3a} -Raum, 25, 37
- T_4 -Raum, 26, 37
- Tangentialbündel, 205
- Tangentialnorm, 214
- Tangentialraum, 205
- Tangentialvektor, 205
- Teiler, 33
- Teilnetz, 10
- Tensorprodukt, 117
- Topologie, 17
 - analytische, 217
 - diskrete, 17, 87, 88
 - feinere, 23
 - gröbere, 23
 - indiskrete, 17
 - leere, 17
 - Produkt-, 24
 - Tychonoff-, 24
 - von einer Subbasis erzeugte, 22, 23
- Topologie der
 - glm. Konv. auf Kompakta, 220
 - lokal konvexen direkten Summe, 86
 - lokalen glm. Konv., 220
 - punktw. schwachen Konv., 175
 - punktweisen Konvergenz, 90, 137, 176
- topologisch äquivalent, 111
- topologisch irreduzibel, 111, 165
- topologisch komplementär, 85

- topologisch komplementierbar, 85
- topologisch komplementiert, 85
- topologisch zyklisch, 111
- topologisch zyklischer Vektor, 111
- topologische Algebra, 109, 110, 179, 182, 184
- topologische Äquivalenz, 111
- topologische Gruppe, 79
- topologischer Raum, 17
- topologischer Vektorraum, 79
- Torsionselement, 42
- torsionsfrei, 42
- total, 66, 91
- Totalordnung, 5
- Träger einer Relation, 4
- trägerabgeschlossen, 4
- trägerabgeschlossene Hülle, 4
- traceform, 227
- Träger (eines Operators), 156
- Träger eines $\ell \in X_*$ in X , 254
- Träger eines $x \in X$ in X^{**} , 255
- Träger eines $x \in X$ in X , 254
- transfinite Induktion, 10
- transfinite Rekursion, 10
- transitiv, 5, 209
- Transitivitätsgebiet, 209
- trennen, die Punkte, 66
- trennend, 30, 110, 111
- Trennungssaxiome, 25, 109
- treu, 42, 110, 172, 212
- treue Linearform, 146
- Tri($\cdot$), 223
- Tripel, 222
 - allgemeines, 222
 - lineares, 222
 - trilineares, 222
- Tripel-Derivation, 224, 225, 233, 242
- Tripel-Homomorphismus, 223
 - konjugiert-linearer, 223
- Tripel-Isomorphismus, 223, 231
- Tripelideal, 223, 238, 243
- Tripelprodukt, 223
 - projiziertes, 237
- Tripelsystem, 222
- Tripotent, 223
- tripotent, 223, 230, 234, 238
- trivial, 111
- Tychonoff-Raum, 26
- Tychonoff-Topologie, 24
- Tychonoff-Würfel, 24
- Ultrafilter, 11, 237
 - fixiert, 11
 - freier, 11
- Ultrapotenz, 90, 237
- Ultraprodukt, 90
- ultraschwache Operator-Topologie, 176
- ultraschwache Topologie, 188
- ultrastarke Operator-Topologie, 177
- ultrastark* Operator-Topologie, 178
- Umgebung, 18
- Umgebungsbasis, 21, 23
- Umgebungsfilter, 24
- Umgebungssubbasis, 22, 23, 91
- Umkehrabbildung, 3
- unbedingt konvergent, 81
- unendlich, 4
- Ungleichung von Cauchy-Schwarz-Bunyakovsky, 147
- uniform boundedness theorem, *siehe* Satz von Banach-Steinhaus
- unitär, 144, 145, 157, 166, 168, 189, 223, 226, 230, 232, 233, 242, 245, 246
 - exponentiell, 145
- unitär (Operator), 157
- unitär äquivalent, 157, 189
- unital, 40, 70, 165
- universell, 189
- universelle Darstellung, 165
- unpunktierter Kegel, 60
- Unter- JB^* -Tripel, 232
- Unteralgebra, 70
- untere Schranke, 6
- Untergruppe, 29
- Unterm modul, 41
- Unterraum, 57
- Unterring, 32
- Untertripel, 223, 232
 - von x erzeugtes, 229
- Untervektorraum, 57
- upward-filtering, 5
- uso* (ultrastarke Op.-Topologie), 177
- us*o* (ultrastark* Op.-Topologie), 178
- uwo* (ultraschwache Op.-Topologie), 176
- $V(1; x)$ (num. Ber. von x bzgl. 1), 130
- $V(a)$ (algebr. num. Wertebereich), 130, 202
- $V_{\text{spatial}}(T)$ (räuml. num. Wertebereich), 130, 153, 202, 261
- Varietät, affine, 211
- Vektor, 57
- Vektorfeld
 - integrierbares, 208
 - polynomiales, 207
 - vollständiges, 208
- Vektorprodukt, 72
- Vektorräume in Dualität setzen, 65
- Vektorraum, 57
 - antiselbstadjungierter Teil, 139
 - geordneter, 58, 146
 - Involution eines, 138
 - komplex strukturierbarer, 114, 242
 - komplexe Strukturierung, 114
 - Komplexifizierung, 116
 - Konjugation eines, 140
 - konjugiert-komplexer, 113, 152, 239
 - prägeordneter, 58

- reelle Strukturierung, 113
 - selbstadjungierter Teil, 139
 - vom komplexen Typus, 114
- Vektorverband, 58
- verallgemeinerte Folge, 10
- Verband, 7, 58, 155, 171, 186
 - vollständiger, 7, 17, 18
- Verknüpfung, 28
- Verknüpfung e. leeren Familie, 28
- Vertex, 124, 128, 133, 168
- Vertex nach Ingelstam, 118
- Vertex-Eigenschaft nach Ingelstam, 119
- Vidav, hermitesch nach, 132
- Vidav-Algebra, 132, 167
- Vidav-Involution, 145
- vollständig, 80, 81, 105, 245, 248
- vollständig regulärer Raum, 25
- vollständiges Vektorfeld, 208

- $W(T)$ (num. Ber. von T bzgl. $[\cdot, \cdot]$), 130
- Würfel, 24
 - euklidischer, 24
 - Tychonoff-, 24
- Weg, 26
 - geschlossener, 26
- wegzusammenhängend, 26, 27, 60
- Weierstrass-Konzept, 204
- Wertebereich
 - numerischer, 130
 - algebraischer, 130, 202
 - räumlicher, 130, 202
- wesentlich, 111
 - Algebra, 159, 180–183
 - Darstellung, 111, 159, 160
- wo (schwache Op.-Topologie), 175
- wohlfundiert, 7
- Wohlordnung, 10
- W^* -Algebra, 180, 238, 243, 268
- W^* -Algebra über $\mathbb{K}$ auf einem Hilbertraum H , 179
- W^* -Algebra über $\mathbb{R}$ auf einem Hilbertraum H über $\mathbb{C}$, 180
- w^*o (schwach* Op.-Topologie), 175

- $\mathbb{Z}$ (ganze Zahlen), 36
- Zählmaß, 260
- Zariski-Abschluss, 211
- zentral, 53
- zentral idempotent, 32
- zentral in S , 32
- zentraler Träger, 190
- Zentralisator, 212
- Zentrum, 32, 110, 183–185, 212
- zerstörter Teil e. Ringes, 39
- zirkulär, 59
- Zorn, Lemma von, 7, 9, 11, 52, 55
- zulässig, 104, 105
- K -zulässig, 38
- zusammenhängend, 26, 27, 60, 87
- Zustand, 146, 164
 - x -unitaler, 124
 - unitaler, 124, 164
- Zustandsraum, 124
- zyklisch, 111
- zyklischer Vektor, 111