

Grundlagen der Analysis

Prof. Dr. Rainer Weissauer

Ruprecht-Karls-Universität Heidelberg
Fakultät für Mathematik und Informatik
Mathematisches Institut

Vorlesungsskriptum SS/WS 2011/12

Bearbeitungsstand: 26. November 2014

Adresse des Dozenten:

Ruprecht-Karls-Universität Heidelberg
Mathematisches Institut
Im Neuenheimer Feld 288
D-69120 Heidelberg, Deutschland
Raum 205

`weissaue@mathi.uni-heidelberg.de`
`http://www.mathi.uni-heidelberg.de/~weissaue/`

Vorwort

Dieses Skript richtet sich als Begleitmaterial der Vorlesung *Höhere Mathematik für Physiker II+III* vorrangig an Studenten der Fachrichtung Physik. In dem zweisemestrigen Zyklus werden die für Physikstudenten relevanten Methoden der Analysis dargestellt.

Die Vorlesung deckt dabei in zwei Semestern mathematische Inhalte ab, die normalerweise in den drei Vorlesungen *Analysis I–III* dargestellt werden. Dabei werden notgedrungen Dinge ausgelassen, da auf die mathematische Strenge der Darstellung nicht verzichtet werden soll. Kenntnisse aus der Vorlesung Lineare Algebra wurden vorausgesetzt. Die Hörer der Vorlesung sollten nämlich die Grundvorlesung *Lineare Algebra I* im Semester davor gehört haben.

Begleitend zu der Vorlesung und dem Übungsbetrieb wurden einmal wöchentlich in einer zusätzlichen Großübung Beispiele behandelt, die in der Vorlesung selbst aus zeitlichen Gründen nicht diskutiert werden konnten.

Das vorliegende Skript folgt in seinem Aufbau keineswegs konsequent der Vorlesung. Auch innerhalb der einzelnen Kapitel wurden in der Vorlesung Teile des Stoffes manchmal geringfügig umgestellt, um vom Timing die Übungsaufgaben so effizient wie möglich mit der Vorlesung abzustimmen.

Kapitel V hat einen sehr speziellen Charakter. Hier werden an einer Stelle des Skriptes spezifische Anwendungen der Analysis gebündelt, die in der Vorlesung und zum Teil in der Großübung gestreut vorgestellt wurden. Ähnliches gilt für Kapitel XI. Das (noch ziemlich unvollständige) Kapitel XII wurde in der Vorlesung nicht behandelt. Es ist gedacht für interessierte Leser und gibt einen kleinen Ausblick.

Inhaltsverzeichnis

Vorwort	i
1 Der Konvergenzbegriff	1
1.1 Angeordnete Körper	1
1.2 Die Euklidsche Norm	4
1.3 Metrische Räume	6
1.4 Folgen in metrischen Räumen	7
1.5 Die geometrische Reihe	9
1.6 Vollständige metrische Räume	10
1.7 Der Banachsche Fixpunktsatz	11
1.8 Quaderschachtelung	13
1.9 Reelle Zahlen	16
1.10 Infimum und Supremum	17
2 Stetige Abbildungen	21
2.1 Stetigkeit	21
2.2 Eigenschaften stetiger Funktionen	23
2.3 Der Zwischenwertsatz	24
2.4 Das ε - δ -Kriterium	25
2.5 Gleichmässige Stetigkeit	26
2.6 Reellwertige stetige Funktionen	27
2.7 Gleichmässige Konvergenz	29
2.8 Vollständigkeit von $C(X)$	30
2.9 Monotone Folgen stetiger Funktionen	30
3 Integration	33
3.1 Vorbemerkungen	33
3.2 Treppenfunktionen	34
3.3 Das reelle Standardintegral	36
3.4 Eigenschaften des Standardintegrals	38
3.5 Der Logarithmus	38
3.6 Das mehrdimensionale Integral $\int_{\mathbb{R}^n} f(x)dx$	40
3.7 Monotone Hüllen	41
3.8 Abstrakte Integrale	42

4	Differentiation	45
4.1	Das Landausymbol	45
4.2	Differenzierbarkeit	46
4.3	Die Jacobi-Matrix	49
4.4	Extremwerte	50
4.5	Symmetrie der Hessematrix	52
4.6	Lokale Maxima	53
4.7	Der Hauptsatz	54
4.8	Differentialgleichungen	55
4.9	Stetig partiell differenzierbare Funktionen	60
4.10	Der Umkehrsatz	61
4.11	Substitutionsregel	64
4.12	Differentialformen	67
4.13	Beweis des Poincare Lemmas	72
4.14	Satz von Stokes für Quader	74
4.15	Analytische Funktionen	76
5	Ausgewählte Themen I	81
5.1	Wegintegrale	81
5.2	Holomorphe Funktionen	83
5.3	Vektorfelder I	84
5.4	Vektorfelder II	87
5.5	Poissonklammer	89
5.6	Variationsrechnung	91
5.7	Satz von Darboux	93
5.8	Kanonische Transformationen	94
5.9	Harmonische Funktionen	95
5.10	Harmonische Polynome	97
5.11	Drehimpuls Operatoren	99
5.12	Taylor Koeffizienten	100
5.13	Orthogonale Gruppen	101
5.14	Fourier-Graßmann Transformation	103
5.15	Laplace Operatoren	104
5.16	Maxwell Gleichungen	105
5.17	Partitionen der Eins	107
6	Lebesgue Integration	111
6.1	Übersicht	111
6.2	Das Lebesgue Integral	112
6.3	Der Verband $L(X)$	114
6.4	Vertauschungssätze	115
6.5	Anwendungen	117
6.6	Nullmengen	118
6.7	Messbare Funktionen	119

7	Verallgemeinerte Funktionen	121
7.1	Basics	121
7.2	Distributionen	122
7.3	Faltung	124
7.4	Coulomb Distribution	126
7.5	Wellengleichung	126
8	Hilberträume	129
8.1	Vorbemerkung	129
8.2	L^2 -Räume	131
8.3	Satz von Fischer-Riesz	132
8.4	Der Folgenraum $L^2(\mathbb{Z})$	133
8.5	Orthonormalbasen	134
8.6	Fourier Reihen	135
8.7	Stone-Weierstraß	137
8.8	Fourier Transformation	138
8.9	Der harmonische Oszillator	142
9	Integration auf Mannigfaltigkeiten	143
9.1	Untermannigfaltigkeiten mit Rand	143
9.2	Randintegrale	145
9.3	Der Satz von Stokes	146
9.4	Standardintegral auf der Kugeloberfläche	147
9.5	Greensche Formel	149
10	Harmonische Analysis	151
10.1	Der Hilbertraum $L^2(S)$	151
10.2	Poisson Kern	152
10.3	Orthogonalität	154
10.4	Harmonische Funktionen sind analytisch	156
10.5	Entwicklung auf Kugelschalen	157
10.6	Die Potential Gleichung	159
11	Ausgewählte Themen II	161
11.1	Kugelvolumina	161
11.2	Überdeckungskompaktheit	163
11.3	Residuensatz	164
11.4	Wärmeleitungskern	165
11.5	Spinordarstellung	166
11.6	Oszillatordarstellung	168
11.7	Spinor-Matrizen	170
11.8	Heisenberggruppe	171
11.9	Vertauschungslemma	172

Die Anwendungen in Kapitel V benötigen gewisse Voraussetzungen über die Differentiation. Die Abhängigkeiten sind wie folgt:

Leitfaden für Kapitel V.

- $1.3 \implies 5.12$
- $3.4 \implies 5.19$ (Anfang) + 5.13
- $4.5 \implies 5.3 \implies 5.12$
- $4.9 \implies 5.4 \implies 5.5 \implies 5.6 \implies 5.7$
- $4.12 \implies 5.13 \implies 5.14$
- $4.14 \implies 5.1 \implies 5.2$
- $4.13 + 4.14 + 5.14 \implies 5.15$

Im ersten Semester habe ich Kapitel I-IV behandelt (ausschließlich der Sektionen 4.14 und 4.15, die ich zu Beginn des zweiten Teils nach Kapitel VI bewiesen habe, da in diesen Abschnitten der Satz von der dominierten Konvergenz benutzt wird; man könnte natürlich hier die benutzten Vertauschungssätze auch erst einmal annehmen) und teilweise die Anwendungen 5.3-5.12. Die Behandlung der Abschnitte aus Kapitel V in der Vorlesung wurde meistens durch Übungsaufgaben vorbereitet und in der großen Übung vertieft.

Das Kapitel XI bestand zum Teil aus Übungsmaterial, Themen der Großübung und der Vorlesung. Kapitel XII und Teile von Kapitel XI wurden nicht in der Vorlesung behandelt, und sind gedacht als Lesestoff zur Anregung und weiteren Vertiefung.

Leitfaden für Kapitel XI.

- $9.3 + 9.4 + 4.11 \implies 11.1$
- $11.2 \implies 8.7$
- $10.5 \implies 11.3$
- $5.13 \implies 11.5 \implies 11.7$
- $8.8 \implies 11.4$
- $8.8 \implies 11.6 \implies 11.8$

In der ersten Vorlesung habe ich Kapitel I-IV behandelt (ausschließlich Abschnitt 4.13 und 4.15) sowie Kapitel V (ausschließlich Abschnitt 4.13 und 4.15). In den Abschnitten 4.13 und 4.15 wird die Vertauschung von Limesprozessen benötigt. Deshalb habe ich sie erst im Wintersemester nach dem Kapitel VI diskutiert, da die benötigten Vertauschungssätze sich dann unmittelbar aus dem Satz von der dominierten Konvergenz ergeben. Die erste Hälfte von Abschnitt 4.12 hatte ich in der Vorlesung bereits im unmittelbaren Anschluß an Abschnitt 4.9 dargestellt, die zweite Hälfte dann nach Abschnitt 4.11 um in der Zwischenzeit das Kalkül in der Großübung und in den Übungsaufgaben etwas vertrauter zu machen.

Ein kurzer Überblick

In Kapitel I studieren wir den *Körper $\mathbb{R}$ der reellen Zahlen* zusammen mit den *Euklidischen Räumen $\mathbb{R}^n$* . Eine naive Definition der reellen Zahlen mit Hilfe von Dezimalbruchentwicklungen wird vermieden. Daß man reelle Zahlen durch eine Dezimalbruchentwicklung beschreiben kann, ergibt sich erst am Ende und eher beiläufig. Einer der Gründe für diese Vorgehensweise ist, daß eine reelle Zahl mehrere Dezimalbruchentwicklungen besitzen kann wegen

$$0,999\dots = 1.$$

Deshalb wird der Körper der reellen Zahlen wie in Mathematikvorlesungen üblich axiomatisch eingeführt als ein *archimedisch angeordneter Cauchy vollständiger* (pythagoräischer) Körper. Dieser Zugang ist sehr natürlich, denn es wird dabei automatisch der Begriff der *konvergenten* beziehungsweise *Cauchy konvergenten* Folgen in metrischen Räumen eingeführt. Diese zentralen Begriffsbildungen sind unverzichtbar und erweisen sich als grundlegend für alle weiteren Aussagen. Die wichtigsten Resultate des Kapitels, neben der Einführung der reellen Zahlen, sind die Sätze über *geometrische Folgen und Reihen* mit dem *Banachschen Fixpunktsatz*, dem Satz von *Bolzano-Weierstraß* und das *Prinzip der monotonen Konvergenz*. Alle hier genannten Resultate sind ihrer Natur nach Konvergenzaussagen. Der Banachsche Fixpunktsatz ist nützlich zur Bestimmung von Umkehrfunktionen und zur Lösung von Differentialgleichungen, der Satz von Bolzano-Weierstraß hängt eng zusammen mit Extremwertproblemen. Diese Sätze werden soweit möglich in der Sprache der *metrischen Räume* behandelt, so daß sie dann sowohl für $\mathbb{R}$ als auch für Euklidische Vektorräume $\mathbb{R}^n$ gelten. Das Prinzip der monotonen Konvergenz erweist sich später als das eigentliche Fundament der Integrationstheorie.

Im Kapitel II untersuchen wir *stetige Funktionen*, also Funktionen die anschaulich gesprochen keine Sprünge machen. Dies macht man am einfachsten präzise mit Hilfe der Folgendefinition der Stetigkeit, die am Anfang des Kapitels II eingeführt wird

$$\left(\lim_{n \rightarrow \infty} x_n = x \right) \implies \left(\lim_{n \rightarrow \infty} f(x_n) = f(x) \right).$$

Wir formulieren diesen Stetigkeitsbegriff später mit Hilfe des sogenannten ε - δ -*Kriteriums* um. Ein Hauptgrund dafür, daß sich dadurch erst der Begriff der *gleichmässigen Stetigkeit* motiviert. Dieser ist viel subtiler als der Begriff der Stetigkeit und wird erst in der ε - δ Formulierung richtig begreifbar. Gleichmässige Stetigkeit ist von grundlegender Bedeutung in vielen Bereichen der Mathematik. Ein zentrales Resultat ist der Satz von Heine, daß eine stetige reellwertige Funktion auf einem folgenkompakten metrischen Raum automatisch gleichmässig stetig ist. Daraus leitet sich später z.B. die Integrierbarkeit stetiger Funktionen ab. Am Ende von Kapitel II beschäftigen wir uns mit (den für diesen Zeitpunkt) recht abstrakt wirkenden Aussagen über gleichmässige oder monotone Konvergenz von Funktionenfolgen auf einem (kompakten) metrischen Raum. Diese abstrakten Sätze werden die spätere Grundlage zur Lösung von Differentialgleichungen (gleichmässige Konvergenz) sein bzw. für die spätere Begründung der *Lebesgue Integrations-theorie* (monotone Konvergenz). Der Leser sollte daher diese Dinge zu diesem frühen Zeitpunkt einfach geduldig zur Kenntnis nehmen.

Im Kapitel III wird das *Euklidische Integral*

$$\int_{\mathbb{R}^n} f(x) dx$$

definiert für eine stetige Funktion $f(x)$ auf $\mathbb{R}^n$ mit kompaktem (d.h. beschränktem) Träger. Die Bedeutung des Integralbegriffes muß sicherlich nicht erläutert werden. Die grundlegende Idee ist, daß man stetige (oder allgemeinere) Funktionen durch *Treppenfunktionen* approximiert um deren Integral zu definieren. Die explizite Berechnung von Integralen ist ein generell schwieriges Problem, und die wichtigste Methode dafür ist der später zu beweisende *Hauptsatz*. Im Kapitel III beschränken wir uns daher auf die Diskussion des *logarithmischen Integrals* $\log(x) = \int_1^x \frac{dt}{t}$. Am Ende des Kapitels diskutieren wir eine Erweiterung des Integrationsbegriffs, welche auf dem *Prinzip der monotonen Konvergenz* beruht. Dies bereitet einerseits die spätere Einführung der Lebesgue Integration vor, erlaubt es andererseits bereits Integrale

$$\int_A f(x) dx$$

für nichtnegative stetige Funktionen f zu definieren, deren Definitionsbereich eine beliebige (!) kompakte Teilmenge A im $\mathbb{R}^n$ ist. Diese Überlegungen zeigen insbesondere, daß jede kompakte Teilmenge $A \subset \mathbb{R}^n$ ein wohldefiniertes Volumen $\text{vol}(A) = \int_A dx$ besitzt. Dies wird im Kapitel IV beim Beweis der allgemeinen n -dimensionalen *Substitutionsformel für Integrale* benutzt.

Kapitel IV ist dann der *Differentialrechnung* gewidmet. Der Begriff einer differenzierbaren Funktion wird von Anfang an für beliebige Funktionen

$$f : \mathbb{R}^n \rightarrow \mathbb{R}^m$$

oder allgemeinere f , definiert auf gewissen zulässigen Teilmengen des $\mathbb{R}^n$, eingeführt. Daß das Differential von f in einem Punkt ξ des Definitionsbereiches

$$Df(\xi) : \mathbb{R}^n \rightarrow \mathbb{R}^m$$

für die Funktion $f(x)$ die beste affin lineare Approximation $f(\xi) + Df(\xi)(x - \xi)$ bei ξ darstellt, liefert rasch die wichtigsten Sätze. Wir diskutieren dann die *Kettenregel*, *Extremwertprobleme*, beweisen den *Hauptsatz* und wenden diesen an auf die Theorie der Differentialgleichungen einer Variable. Im Anschluß diskutieren wir die schwierigeren Sätze der Analysis mehrerer Variablen, den *Satz von der Umkehrfunktion* und die *Substitutionsregel* für mehrdimensionale Integrale. Unser Ziel ist es dabei so schnell wie möglich den Begriff der *Differentialformen* einzuführen und zu motivieren. Erst mit Hilfe dieser Differentialformen kann man Verallgemeinerungen des Hauptsatzes der Analysis für höhere Dimensionen überhaupt formulieren. Ein Teil dieses allgemeinen Hauptsatzes ist das Poincare Lemma (eine weitreichende Verallgemeinerung der *Hauptsätze der klassischen Vektoranalysis*), ein anderer Teil ist der *Satz von Stokes*

$$\int_{\partial Q} \omega = \int_Q d\omega ,$$

der am Ende von Kapitel IV zuerst einmal nur für Quader Q bewiesen wird. Für die meisten lokalen Anwendungen reicht dies bereits aus.

Im Kapitel VI wird die *Lebesgue Integrationstheorie* entwickelt. Diese Theorie, allem voran der *Satz von der dominierten Konvergenz* und der *Satz von Beppo Levi*, stellt die Grundlage für das spätere Kapitel über Hilberträume dar. Beide Sätze sind Vertauschungssätze, die garantieren daß unter gewissen Voraussetzungen Limiten mit der Integration vertauschen

$$\lim_{n \rightarrow \infty} \int f_n(x) dx = \int \lim_{n \rightarrow \infty} f_n(x) dx .$$

Nur mit Hilfe dieser Sätze sind die später wichtigen $L^2(X)$ -Hilberträume definierbar. Kapitel VI ist zum Teil sehr technisch, aber letztlich auch sehr einfach. Zum Verständnis ist hier vor allem Geduld erforderlich. Einige Beweise von Kapitel IV (Beweis des Poincare Lemmas und die Diskussion analytischer Funktionen) hängen bei dem von uns gewählten Zugang vom Satz der dominierten Konvergenz ab. Der abschliessende Abschnitt über *messbare Funktionen* enthält in kondensierter Form eigentlich alle wesentlichen Resultate von Kapitel VI.

Im Kapitel VII studieren wir *Distributionen* und bestimmen *Fundamentallösungen* von der *Laplace Gleichung* und der *D'Alembert Gleichung*.

Das Kapitel VIII ist von Bedeutung für die *Quantentheorie*. Aus diesem Grund geben wir einen kurzen Überblick über die Querverbindungen am Anfang des Kapitels. Im Kapitel VIII definieren wir den Begriff des (separablen) Hilbertraumes und beweisen in diesem Kontext die Sätze über *Fourierzerlegung*. Ein besonderer Augenmerk wurde dabei auf eine vollständige und ausführliche Diskussion der *reellen Fourier Transformation* gelegt. Es gibt einen engen Zusammenhang zur *Quantentheorie* (Stichwort *harmonischer Oszillator*).

Im Kapitel IX diskutieren wir die Analysis auf eingebetteten *Mannigfaltigkeiten* und beweisen in diesem Kontext den *Satz von Stokes* (und damit den Satz von *Gauß*) in der Sprache der Differentialformen. Das für uns wichtigste Beispiel einer eingebetteten Mannigfaltigkeit ist die $(n - 1)$ -dimensionale *Sphäre* S^{n-1} im Euklidischen Raum $\mathbb{R}^n$. Dieses Beispiel behandeln wir ausführlich. Insbesondere definieren wir das rotationsinvariante Standardintegral¹ auf der Sphäre und diskutieren die *Greenschen Formeln*.

Im Kapitel X wenden wir die Ergebnisse aus Kapitel VIII und betreiben Analysis auf der Sphäre unter Berücksichtigung der Operation der Drehgruppe. Wir finden in den *sphärischen Kugelfunktionen* eine Hilbertraum Basis von $L^2(S^{n-1})$ in beliebiger Dimension n und beweisen die *Poissonformel*. Als Anwendung erhalten wir die *Laurent-Entwicklung* von harmonischen Funktionen auf Kugelschalen (*Multipolentwicklungen*) mit genauer Diskussion von Konvergenzfragen, und beweisen damit als Spezialfall die *Analytizität harmonischer Funktionen*. Dies umfaßt im Spezialfall $n = 2$ die Grundlagen der *komplexen Funktionentheorie*.

¹Die Sphäre S^{n-1} ist eine kompakte Teilmenge im $\mathbb{R}^n$. Das Standardintegral auf der Sphäre $A = S^{n-1}$ ist nicht das in Kapitel III diskutierte Euklidische Integral $\int_A f(x) dx$, denn dieses ist identisch Null (die Sphäre ist eine Nullmenge im $\mathbb{R}^n$). Das Standardintegral auf der Sphäre schreiben wird deshalb $\int_{S^{n-1}} f(x) \sigma_{n-1}(x)$.

1 Der Konvergenzbegriff

1.1 Angeordnete Körper

Wir wiederholen an dieser Stelle den aus der Linearen Algebra bekannten Begriff des *Körpers*. Es handelt sich dabei um einen Rechenbereich mit Multiplikation und Addition.

Genauer gilt: Ein Körper ist ein Tupel $(K, +, \cdot, 0, 1)$ bestehend aus einer Menge K , zwei Verknüpfungen $+: K \times K \rightarrow K$, genannt *Addition*, und $\cdot: K \times K \rightarrow K$, genannt *Multiplikation*, sowie zwei verschiedenen Elementen 0 (Nullelement) und 1 (Einselement) mit gewissen Eigenschaften. So soll zum einen das Tupel $(K, +, 0)$ eine abelsche Gruppe mit neutralem Element 0 sein. D. h. man kann beliebige Elemente $a, b \in K$ addieren, das heißt durch $+$ verknüpfen, so daß gilt $a + b = b + a \in K$ sowie $a + 0 = a$, und jede Gleichung

$$x + a = b$$

hat für gegebenes $a, b \in K$ eine eindeutige Lösung x . Wir schreiben diese in der Form $x = b - a$.

Zum anderen soll die Menge der von Null verschiedenen Elemente $K^* \subseteq K$ eine abelsche Gruppe $(K^*, \cdot, 1)$ definieren. Insbesondere ist daher für alle $a, b \in K$ mit $a \neq 0, b \neq 0$ die Gleichung

$$x \cdot a = b$$

eindeutig lösbar. Deren Lösung schreiben wir als $x = a/b = b \cdot a^{-1}$. Für gewöhnlich lassen wir den Punkt für die Multiplikation meist weg und schreiben kurz ab statt $a \cdot b$. Ist $b = 0$ und $a \neq 0$, dann ist übrigens $x = 0$ die einzige Lösung der Gleichung $x \cdot a = b$. Dies folgt aus dem *Distributivgesetz*, das in einem Körper erfüllt sein soll. Im Distributivgesetz wird $a(b + c) = ab + ac$ gefordert für alle $a, b, c \in K$ und es impliziert $0 \cdot a = a \cdot 0 = 0$ für alle $a \in K$.

Der Begriff des Körpers ist bereits aus der Linearen Algebra bekannt. Typische Beispiele sind: der Körper $\mathbb{Q}$ der *rationalen Zahlen*, der Körper $\mathbb{R}$ der *reellen Zahlen* sowie der Körper $\mathbb{C}$ der *komplexen Zahlen*. Für die Analysis spielt der Körper $\mathbb{R}$ der reellen Zahlen eine fundamentale Rolle. Seine Elemente stellen wir uns intuitiv vor als die Punkte auf einer lückenlosen Geraden. Wir sind von der Schule gewohnt in diesem Körper zu rechnen.

Eine sehr wichtige Eigenschaft des Körpers der reellen Zahlen besteht darin, daß dieser Körper eine *Anordnung* besitzt. Eine Anordnung ist eine Relation $x < y$: Alle x und y aus einem Körper K lassen sich also in Bezug auf diese Anordnung vergleichen. Man setzt formal $y > x \Leftrightarrow x < y$ und schreibt

$$x \leq y \quad :\Longleftrightarrow \quad x < y \text{ oder } x = y.$$

1 Der Konvergenzbegriff

Der Begriff der Anordnung ist auch in der Physik sehr wesentlich, wenn es um die Parametrisierung der Zeit geht. Das Vorher und Nachher von Ereignissen spielt eine fundamentale Rolle bei der Kausalität und dem physikalischen Begriff der Entropie.

Der Begriff eines angeordneten Körpers lässt sich mathematisch in axiomatischer Weise definieren. Ein angeordneter Körper $(K, <)$ ist ein Körper K zusammen mit einer ausgezeichneten Teilmenge $P \subseteq K^*$. Man nennt dann P den „Kegel der positiven Zahlen“ des angeordneten Körpers. Dies ist ein eindimensionales Analogon des in der Physik auftretenden vorderen Lichtkegels im Minkowskiraum. Eine Zahl $x \in K$ nennt man negativ oder man schreibt $x \in -P$, wenn ihr negatives $-x$ in P liegt.

Definition 1.1. Ein Tupel (K, P) bestehend aus einem Körper K und einer Teilmenge P von K^* heißt **angeordneter Körper**, wenn

- (1) $K = P \cup \{0\} \cup -P$, d. h. K zerlegt sich disjunkt in P , $-P$ und Null.
- (2) $P + P \subseteq P$, d. h. die Summe zweier Zahlen aus P ist wiederum in P
- (3) $P \cdot P \subseteq P$, d. h. das Produkt zweier Zahlen aus P ist wieder in P ; also $1 \in P$ gilt und wir schreiben $x < y$ genau dann wenn $y - x \in P$ gilt.

Die Menge P definiert damit eine Relation $<$ auf K und wir schreiben auch $(K, <)$ anstelle von (K, P) , denn per Definition gilt

$$P = \{x \in K \mid 0 < x\}.$$

Aus dem ersten Axiom angeordneter Körper folgt für zwei Zahlen $x, y \in K$ entweder $x < y$ oder $x = y$ oder $y < x$ im ausschliesslichen Sinn. Für $y = 0$ folgt unmittelbar $-P = \{x \in K \mid x < 0\}$. Wir bemerken folgende Eigenschaft:

- Jedes Quadrat x^2 einer Zahl x aus K^* ist positiv, kurz: $x^2 \in P$.

Dies ist klar für $x \in P$ nach dem dritten Axiom. Ist x nicht in P , dann ist $-x \in P$ und damit $(-x)^2 \in P$ nach dem ersten Axiom. Also $x^2 \in P$ wegen $x^2 = (-1)^2 \cdot x^2 = (-x)^2 \in P$. Hier haben wir benutzt $-x = (-1) \cdot x$ und $(-1) \cdot (-1) = 1$. Diese Eigenschaften gelten in jedem Körper [benutze dazu das Distributivgesetz].

Man sieht daher, daß der Körper der komplexen Zahlen keine Anordnung besitzen kann, denn $-1 = i^2$ ist ein Quadrat in $\mathbb{C}^*$, aber liegt nicht in P wegen $1 \in P$.

Bemerkung 1.2. Sei $(K, <)$ ein angeordneter Körper. Dann gelten wegen Definition 1.1 folgende Eigenschaften:

- (1) Es gilt entweder $x < y$ oder $x = y$ oder $x > y$ (im ausschließlichen Sinn)
- (2) Ist $x < y$ und $y < z$, dann ist $x < z$.
- (3) Ist $x < y$, dann gilt $x + z < y + z$ für alle $z \in K$.

Beweis. (1) ist klar. Zu (2) beachte: Aus $y - x \in P$ und $z - y \in P$ folgt $z - x = (z - y) + (y - x) \in P + P \subseteq P$. Zu (3) beachte: Aus $y - x > 0$ folgt $(y + z) - (x + z) = y - x > 0$ und damit auch $x + z < y + z$. $\square$

In einem angeordneten Körper definiert man den Betrag $|x|$ eines Element $x \in K$ wie folgt: $|x| = 0$ gilt genau dann wenn $x = 0$; und für $x \neq 0$ sei per Definition $|x| = x$ resp. $-x$ je nachdem ob $x \in P$ oder $x \notin P$. Dann gilt nach Definition $|x| \in P$ für $x \neq 0$, und man sieht sofort

$$|x \cdot y| = |x| \cdot |y|.$$

Natürliche Zahlen. Jeder angeordnete Körper enthält die natürlichen Zahlen in der Form

$$\mathbb{N} = \{0, 1, 2, 3, \dots\} := \{0, 1, 1 + 1, 1 + 1 + 1, \dots\}.$$

Beachte nämlich $0 < 1$, und wegen Eigenschaft (3) folgt dann durch Addition $1 = 0 + 1 < 1 + 1 = 2$ und dann analog $2 = 1 + 1 < 3 = 1 + 1 + 1$ und so weiter. Die Zahlen $0, 1, 2, 3, \dots$ sind wegen Eigenschaft (2) paarweise verschieden. Die so definierte Teilmenge $\mathbb{N} \subseteq K$ ist unter Multiplikation und Addition abgeschlossen, wie man sofort mit Hilfe des Distributivgesetzes in K zeigt, und kann mit den natürlichen Zahlen identifiziert werden. Wir benutzen folgende Notation: Das Produkt der ersten n positiven natürlichen Zahlen sei $n! = n \cdot (n-1) \cdots 2 \cdot 1$.

Wegen der Körperaxiome liegen daher die ganzen Zahlen $\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\}$ als paarweise verschiedenen Zahlen in einem angeordneten Körper, und damit auch die Quotienten a/b ganzer Zahlen a und $b \neq 0$. Also ist der Körper der rationalen Zahlen ein Teilkörper jedes angeordneten Körpers: $\mathbb{Q} \subseteq K$. Insbesondere enthält K unendlich viele Elemente. Endliche Körper besitzen daher keine Anordnung.

Es stellt sich nun die Frage, ob die Anordnung die charakteristischen Eigenschaften der reellen Zahlen $\mathbb{R}$ bereits vollständig beschreibt. Das ist nicht der Fall, denn der Körper der rationalen Zahlen $\mathbb{Q}$ ist auch ein angeordneter Körper, aber verschieden vom Körper der reellen Zahlen. Wir wollen daher weitere Eigenschaften suchen, die $\mathbb{R}$ charakterisieren.

Definition 1.3. Ein angeordneter Körper $(K, <)$ heißt **archimedisch**, wenn gilt: Für jedes $x \in K$ existiert eine natürliche Zahl $n \in \mathbb{N}$ mit der Eigenschaft $x < n$.

Definition 1.4. Ein archimedischer Körper $(K, <)$ heißt **pythagoräisch**, wenn gilt: Jede Zahl aus P ist ein Quadrat in K .

Der Körper $\mathbb{Q}$ der rationalen Zahlen ist archimedisch, aber nicht pythagoräisch. 2 ist positiv aber kein Quadrat in $\mathbb{Q}$, weil die Gleichung $n^2 = 2m^2$ keine ganzzahligen Lösungen m, n besitzt [Benutze die Eindeutigkeit der Primfaktorzerlegung].

Für jede Zahl $y \in P$ eines angeordneten Körpers mit $y \geq \frac{1}{4}$ gilt $0 < \eta \leq \frac{1}{4}$ für $\eta := 1/16y$. Ist η ein Quadrat $\eta = \xi^2$, dann auch $y = (1/4\xi)^2$. Dies zeigt uns später in Lemma 1.24, daß ein vollständiger archimedischer Körper auch pythagoräisch ist.

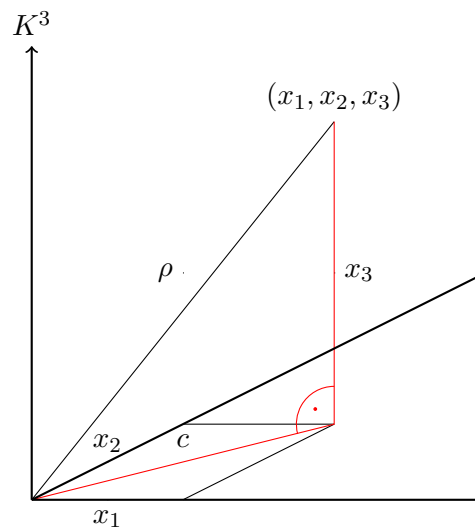
In einem pythagoräischen Körper besitzt jede nicht negative Zahl $y \in K$ eine eindeutig bestimmte nicht negative Quadratwurzel $x_1 = +\sqrt{y}$, d.h. eine eindeutig bestimmte nicht negative Lösung $x = x_1$ der Gleichung $x^2 - y = 0$. [Für $y = 0$ ist das klar. Ist $y \in P$, gibt es eine Lösung x_1 , da K pythagoräisch ist. Dann ist auch $x_2 = -x_1$ eine Lösung mit $x_2 \neq x_1$, und damit oBdA $x_1 \in P$ und $x_2 \in -P$. In einem Körper hat aber die Gleichung $x^2 - y = (x - x_1)(x - x_2) = 0$ dann höchstens die Lösungen $x = x_1, x_2$.] Der **Absolutbetrag** $|x|$ einer Zahl $x \in K^*$ kann daher in der Form $|x| = +\sqrt{x^2}$ geschrieben werden.

1.2 Die Euklidische Norm

Wir wollen für einen pythagoräischen Körper K die **Norm** (oder auch Länge) eines Vektors im r -dimensionalen Vektorraum K^r definieren. Betrachte den r -dimensionalen Standardvektorraum

$$K^r = \{(x_1, \dots, x_r) \mid x_1, \dots, x_r \in K\}$$

für einen pythagoräischen Körper K . Motivation: Für einen beliebigen Punkt $x = (x_1, x_2, x_3)$ im Anschauungsraum K^3 würde der Satz von Pythagoras den Abstand ρ von x zum Nullpunkt liefern durch die Formel $\rho^2 = c^2 + x_3^2 = (x_1^2 + x_2^2) + x_3^2$



Dadurch motiviert, definiert man die **Standardnorm** oder **Euklidische Norm** für $x = (x_1, \dots, x_r)$ aus K^r entsprechend als

$$\|x\| = +\sqrt{x_1^2 + \dots + x_r^2}.$$

Wegen Satz 1.5 ist $x_1^2 + \dots + x_r^2 \geq 0$, und daher existiert die (!) positive Wurzel aus dieser Zahl (siehe letzter Abschnitt von §1.1). Dies macht $\|x\|$ wohldefiniert als Zahl in $P \cup \{0\} \subseteq K$, und $\|x\| = 0$ gilt genau dann wenn $x = 0$ gilt (Satz 1.5). Beachte

$$\|\lambda \cdot x\| = |\lambda| \cdot \|x\|$$

für alle Skalare λ aus K , denn beide Seiten sind ≥ 0 und haben dasselbe Quadrat. [Benutze Übungsaufgabe]. Im eindimensionalen Fall $r = 1$ ist $\|x\| = |x|$.

Satz 1.5. Sei K pythagoräisch und ein Vektor $x = (x_1, \dots, x_r) \in K^r$ gegeben. Dann gilt: $x_1^2 + \dots + x_r^2 \in P$ oder $x_1^2 + \dots + x_r^2 = 0$. Letzteres gilt genau dann, wenn

$$x_1 = \dots = x_r = 0.$$

Beweis. Den Beweis reduziert man durch Induktion nach r auf den Fall $r = 2$. Dieser Fall sei als Übungsaufgabe gestellt. $\square$

Definition 1.6. Sei K pythagoräisch. Für $x, y \in K^r$ nennt man

$$(x, y) = x \cdot y = x_1 y_1 + \cdots + x_r y_r$$

das **Standard-Skalarprodukt** von x und y .

Insbesondere gilt $x \cdot x = \|x\|^2$ für $x = (x_1, \dots, x_r) \in K^r$.

Satz 1.7 (Ungleichung von Schwarz). Seien $x, y \in K^r$. Dann ist

$$|x \cdot y| \leq \|x\| \cdot \|y\|.$$

Gleichheit gilt genau dann, wenn x und y proportional sind.

Beweis. ObdA sei $x - t \cdot y \neq 0$ für alle $t \in K$ (d. h. x und y seien nicht proportional). Dann gilt

$$0 < \|x - ty\|^2,$$

d. h.

$$0 < (x - ty, x - ty) = \sum_{i=1}^r (x_i - ty_i)^2 = \|x\|^2 - 2t(x, y) + t^2\|y\|^2,$$

wegen $(x_i - ty_i)^2 = x_i^2 - 2tx_i y_i + t^2 y_i^2$. Sei nun obdA $y \neq 0$. Dann folgt

$$t^2 - \frac{2t(x, y)}{\|y\|^2} + \frac{\|x\|^2}{\|y\|^2} > 0$$

$$t^2 - \frac{2t(x, y)}{\|y\|^2} + \left(\frac{(x, y)}{\|y\|^2} \right)^2 > \frac{(x, y)^2}{\|y\|^4} - \frac{\|x\|^2}{\|y\|^2}$$

$$\left(t - \frac{(x, y)}{\|y\|^2} \right)^2 > \frac{(x, y)^2 - \|x\|^2 \cdot \|y\|^2}{\|y\|^4}.$$

Setzt man $t := \|y\|^{-2}(x, y)$, dann folgt $(x, y)^2 - \|x\|^2 \cdot \|y\|^2 < 0$. $\square$

Satz 1.8 (Dreiecksungleichung im K^r). Sei K pythagoräisch und seien $x, y \in K^r$ und $\lambda \in K$. Dann gilt $\|\lambda \cdot x\| = |\lambda| \cdot \|x\|$, sowie $\|x\| = 0 \iff x = 0$ und

$$\|x + y\| \leq \|x\| + \|y\|.$$

1 Der Konvergenzbegriff

Beweis. Nach dem Übungsblatt 1 genügt es $\|x + y\|^2 \leq (\|x\| + \|y\|)^2$ zu zeigen. Die linke Seite ist

$$(x + y, x + y) = \|x\|^2 + 2(x, y) + \|y\|^2,$$

und die rechte Seite ist $(\|x\| + \|y\|)^2 = \|x\|^2 + 2\|x\|\|y\| + \|y\|^2$. Die Behauptung folgt daher aus $2(x, y) \leq 2\|x\|\|y\|$ und der Schwarzschen Ungleichung

$$2|(x, y)| \leq 2\|x\|\|y\|.$$

□

Folgerung. Die Funktion $\|\cdot\| : V \rightarrow K$ definiert eine **Norm** auf dem K -Vektorraum V , d.h. es gilt: a) $\|x\| \geq 0$ und $\|x\| = 0 \iff x = 0$, b) $\|\lambda x\| = |\lambda| \cdot \|x\|$ für alle $\lambda \in K$ und alle $x \in V$, c) $\|x + y\| \leq \|x\| + \|y\|$ für alle $x, y \in V$.

1.3 Metrische Räume

Im Folgenden sei $(K, <)$ ein fest gewählter archimedischer Körper (später dann immer der Körper der reellen Zahlen).

Definition 1.9. Sei X eine beliebige Menge. Das Tupel (X, d) mit einer Abbildung

$$d: X \times X \rightarrow K, (x, y) \mapsto d(x, y)$$

heißt **metrischer Raum** (bezüglich K), falls für die Abbildung d gilt:

- (1) (Positivität) $d(x, y) \geq 0$ für alle $x, y \in X$ und $d(x, y) = 0$ gilt genau dann, wenn $x = y$.
- (2) (Symmetrie) Für alle $x, y \in X$ ist $d(x, y) = d(y, x)$.
- (3) (Dreiecksungleichung) $d(x, z) \leq d(x, y) + d(y, z)$ gilt für alle $x, y, z \in X$.

Man nennt $d(x, y)$ die **Abstandsfunktion** oder **Metrik** des metrischen Raumes (X, d) . In einem metrischen Raum gilt automatisch die folgende **untere Dreiecksungleichung**

$$|d(x, z) - d(x', z)| \leq d(x, x'),$$

denn die Dreiecksungleichung $d(x, z) \leq d(x, x') + d(x', z)$ impliziert $d(x, z) - d(x', z) \leq d(x, x')$, und durch Vertauschung von x und x' folgt daraus die Behauptung.

Ist $\|\cdot\| : V \rightarrow K$ eine **Norm** auf einem K -Vektorraum V , dann definiert $d(x, y) = \|x - y\|$ eine **Metrik** auf V . [Beachte $d(y, x) = \|y - x\| = \|-(x - y)\| = |-1| \cdot \|x - y\| = \|x - y\| = d(x, y)$ sowie $d(x, z) = \|x - z\| = \|(x - y) + (y - z)\| \leq \|x - y\| + \|y - z\| = d(x, y) + d(y, z)$.]

Das für uns wichtigste Beispiel eines metrischen Raumes ist der archimedische Körper K selbst mit seiner Metrik $d(x, y) = |x - y|$. Ist K ein pythagoräischer Körper und ist $\|\cdot\|$ die **Euklidische Norm** auf dem r -dimensionalen K -Vektorraum K^r , dann definiert

$$d(x, y) = \|x - y\|$$

die sogenannte **Standardmetrik** auf $V = K^r$. Den so definierten metrischen Raum nennt man den r -dimensionalen **Euklidischen Raum**.

1.4 Folgen in metrischen Räumen

Fast alle Aussagen der Analysis bauen auf den in diesem Abschnitt erläuterten Konzepten auf. Wir beginnen mit dem Begriff einer Folge:

Definition 1.10. Sei X eine beliebige Menge. Eine **Folge in X** ist eine Abbildung

$$x: \mathbb{N} = \{0, 1, 2, 3, \dots\} \rightarrow X.$$

Anschaulich läßt sich eine Folge als eine unendliche „Durchnumerierung“ von Elementen interpretieren. Dies wirkt sich auch auf die Notation aus: Statt einer Abbildungsbeziehung, also einer Auflistung der Art

$$\begin{aligned} 0 &\mapsto x(0), \\ 1 &\mapsto x(1), \\ 2 &\mapsto x(2), \\ &\vdots \end{aligned}$$

verwenden wir Indizierungen zur Numerierung der betroffenen Elemente von X , um die Folge zu beschreiben:

$$x = (x_0, x_1, x_2, x_3, \dots).$$

Die Elemente x_0, x_1, x_2 , etc. heißen die *Folgenglieder*, bzw. kurz die *Glieder* der Folge x .

Bisher haben wir keine näheren Anforderungen an die Menge X gestellt. Wir nehmen jetzt an, daß X ein metrischer Raum ist. Wir wollen uns daher mit den Abständen zwischen Folgegliedern befassen und durch folgende Definition insbesondere ganz bestimmte Folgen behandeln:

Definition 1.11. Sei (X, d) ein metrischer Raum. Eine Folge $x_0, x_1, x_2, \dots$ in (X, d) heißt **Cauchyfolge**, wenn zu jedem $\varepsilon > 0$ aus K eine natürliche Zahl $N = N(\varepsilon)$ existiert, so daß für alle $n, m \in \mathbb{N}_0$ gilt

$$n, m \geq N \quad \Rightarrow \quad d(x_n, x_m) < \varepsilon.$$

Zur anschaulichen Bedeutung. Zunächst taucht hierbei die Zahl ε auf. Diese steht intuitiv gesprochen für etwas „beliebig Kleines“. Man stellt sich dabei vor, dass egal wie klein ε gewählt wird, man trotzdem noch davon abhängende Zahlen $N(\varepsilon)$ wie behauptet finden kann. Man kann, wenn man nur weit genug mit dem Index geht, den Abstand zwischen Folgengliedern unter jede noch so kleine Schranke drücken. Anschaulich besteht das Wesen einer Cauchyfolge also darin, daß die Abstände zwischen den Gliedern immer enger werden. Dies hängt substanziell von der gewählten Abstandsfunktion d ab.

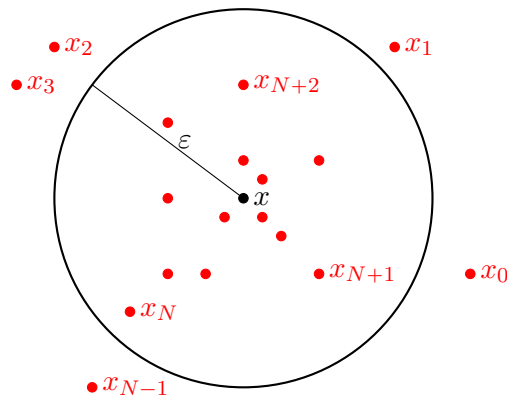
Definition 1.12. Eine Folge $x_0, x_1, \dots$ in (X, d) heißt **konvergent** gegen einen Grenzwert $x \in X$, wenn zu jedem $\varepsilon > 0$ ein $N = N(\varepsilon) \in \mathbb{N}_0$ existiert, so daß für alle $n \geq N(\varepsilon)$ gilt $d(x_n, x) < \varepsilon$.

1 Der Konvergenzbegriff

Zur Veranschaulichung. Wir fixieren ein $\varepsilon > 0$ und betrachten die offene Kugel

$$B_\varepsilon(x) = \{y \in X \mid d(x, y) < \varepsilon\}$$

um x mit dem Radius ε . Für eine gegen x konvergierende Folge x_k liegen alle x_k mit $k \geq N(\varepsilon)$ innerhalb von $B_\varepsilon(x)$. Dies sind fast alle (d.h. alle bis auf endlich viele) Folgenglieder der Folge, insbesondere immer unendlich viele. Dass immer nur endlich viele außerhalb einer beliebigen offenen ε -Kugel, also in $X \setminus B_\varepsilon(x)$ liegen können, soll die folgende Graphik veranschaulichen:



Die Folgenglieder sammeln sich immer mehr in der Nähe von x . Egal wie klein ε wird, alle bis auf höchstens endlich viele Folgenglieder haben einen Abstand zu x , der kleiner als ε ist.

Nun nennen wir eine Folge $(x_n)_{n \in \mathbb{N}_0}$ **beschränkt**, wenn es $y \in X$ und ein $C \in K$ gibt, so daß für alle n gilt $d(x_n, y) \leq C$. Diesen Begriff wollen wir im Folgenden mit den bekannten Begriffen der Cauchyfolge und der konvergenten Folge verknüpfen:

Lemma 1.13. *Jede konvergente Folge ist eine Cauchyfolge. Jede Cauchyfolge ist beschränkt.*

Beweis. Zunächst beweisen wir die erste Aussage. Gegeben sei eine Folge $(x_n)_{n \in \mathbb{N}_0}$ mit ihrem Grenzwert $x \in X$. Dann ist $d(x_n, x) < \frac{1}{2}\varepsilon$ für alle $n \geq N := N(\frac{1}{2}\varepsilon)$. Aus der Dreiecksungleichung

$$d(x_n, x_m) \leq d(x_n, x) + d(x, x_m)$$

folgt dann $d(x_n, x_m) < \frac{1}{2}\varepsilon + \frac{1}{2}\varepsilon$, also $d(x_n, x_m) < \varepsilon$, für alle $n, m \geq N$. Somit ist (x_n) eine Cauchyfolge.

Kommen wir nun zum zweiten Teil. Im Falle $\varepsilon = 1$ gilt $d(x_n, x_m) < 1$ für $n, m \geq N$ nach der Cauchy-eigenschaft. Setze nun $y := x_N$. Dann ist

$$d(x_n, y) = d(x_n, x_N) < 1$$

für alle $n \geq N$. Also ist $d(x_n, y) \leq C$ für $C = \max(d(x_0, y), \dots, d(x_{N-1}, y), 1)$. □

Lemma 1.14. *(Eindeutigkeit des Grenzwertes) Sei $x_0, x_1, \dots$ eine Folge in (X, d) , welche gegen $x \in X$ und $y \in X$ konvergiert. Dann ist $x = y$.*

Beweis. Wir führen einen Widerspruchsbeweis. Wäre $d(x, y) > 0$, dann existiert wegen der Konvergenz der Folge x_n für $\varepsilon = d(x, y)$ ein $N = N(\frac{1}{2}\varepsilon)$ aus $\mathbb{N}_0$ so daß $d(x_n, x) < \frac{1}{2}\varepsilon$ gilt für $n \geq N(\frac{1}{2}\varepsilon)$. Analog existiert ein $M = M(\frac{1}{2}\varepsilon) \in \mathbb{N}$ mit $d(x_n, y) < \frac{1}{2}\varepsilon$ für $n \geq M$. Aufgrund der Dreiecksungleichung $d(x, y) \leq d(x, x_n) + d(x_n, y)$ und der Symmetrie $d(x, x_n) = d(x_n, x)$ folgt für alle $n \geq \max(N, M)$ dann

$$d(x, y) < \varepsilon.$$

Wir erhalten einen Widerspruch zu der Annahme $\varepsilon = d(x, y)$. Es folgt $x = y$. $\square$

Dieses Lemma rechtfertigt es von dem Grenzwert einer Folge zu sprechen. Daß eine Folge $(x_n)_{n \in \mathbb{N}}$ gegen einen Grenzwert $x \in X$ konvergiert, wird häufig durch folgende Schreibweisen

$$\lim_{n \rightarrow \infty} x_n = x$$

oder

$$x_n \xrightarrow{n \rightarrow \infty} x$$

angedeutet. Bei letzterer Schreibweise wird der Ausdruck $n \rightarrow \infty$ teilweise auch über den Pfeil geschrieben oder ganz weggelassen.

Eine *Teilfolge* einer gegebenen Folge $(x_n)_{n \in \mathbb{N}}$ ist eine Auswahl $(x_{n_k})_{k \in \mathbb{N}}$, die ihrerseits auch wiederum eine Folge ist und deren Glieder allesamt auch in dieser Reihenfolge (jedoch mit beliebig großen Lücken dazwischen) Glieder der Folge $(x_n)_{n \in \mathbb{N}}$ sind. Wir fordern dabei, daß $k \mapsto n_k$ eine Injektion $\mathbb{N} \hookrightarrow \mathbb{N}$ ist. Zum Beispiel ist die Folge

$$x_0, x_2, x_4, x_6, \dots$$

eine Teilfolge einer Folge $(x_n)_{n \in \mathbb{N}}$, bei der jedes zweite Glied (immer genau die mit ungeradem Index) herausgenommen wurde. Diesen Begriff werden wir in Kürze (Satz 1.26) mit dem Begriff einer beschränkten Folge verknüpfen. Nicht jede beschränkte Folge ist konvergent. So hat beispielsweise die Folge $x_n = (-1)^n$ keinen Grenzwert, ist aber beschränkt.

1.5 Die geometrische Reihe

Lemma 1.15. In einem archimedischen Körper K konvergiert im Fall $|q| < 1$ jede geometrische Folge $x_n = C \cdot q^n$ gegen Null.

Beweis. Sei $\varepsilon > 0$ gegeben. Nach Annahme gilt $|q| < 1$ und damit $|q|^{-1} > 1$. Somit gilt $|q|^{-1} = 1 + x$ für ein $x > 0$. Die Ungleichung $d(C \cdot q^n, 0) < \varepsilon$ ist dann äquivalent zu

$$C/\varepsilon < (1 + x)^n.$$

Man zeigt nun leicht die Bernoulli Ungleichung $(1 + x)^n \geq 1 + n \cdot x$ mittels Induktion nach n . Im Fall $n = 0$ und $n = 1$ ist dies trivialerweise richtig. Ist $n \geq 1$, dann ist $(1 + x)^n$ größer als $1 + n \cdot x$, wie die Induktionsannahme $(1 + x)^{n-1} \geq 1 + (n-1) \cdot x$ zeigt:

$$(1 + x) \cdot (1 + x)^{n-1} \geq (1 + x) \cdot (1 + (n-1) \cdot x) = 1 + n \cdot x + (n-1) \cdot x^2 \geq 1 + n \cdot x.$$

1 Der Konvergenzbegriff

Das Archimedische Axiom garantiert die Existenz einer natürlichen Zahl $N > (C/\varepsilon - 1)/x$. Für alle $n \geq N$ gilt dann

$$C/\varepsilon < 1 + n \cdot x.$$

Daraus folgt wegen $1 + n \cdot x \leq (1 + x)^n$ das Lemma. $\square$

Dies hat die folgende Konsequenz: Die *geometrische Reihe* $s_n = 1 + q + q^2 + \cdots + q^n$ konvergiert für $|q| < 1$ und hat in diesem Falle den Grenzwert

$$\lim_{n \rightarrow \infty} \sum_{i=0}^n q^i = \frac{1}{1-q}.$$

Dies folgt aus der verallgemeinerten Binomialformel

$$(1 - q) \cdot (1 + q + \cdots + q^n) = 1 - q^{n+1},$$

die man leicht durch Induktion nach n beweist. Diese Formel zeigt $s_n - \frac{1}{1-q} = \frac{-q^{n+1}}{1-q}$. Also

$$d\left(s_n, \frac{1}{1-q}\right) = C \cdot |q|^n$$

für $C = |q/(1 - q)|$. Aus dem letzten Lemma folgt daher $d(s_n, \frac{1}{1-q}) < \varepsilon$ für $n \geq N(\varepsilon)$. Das zeigt die Behauptung. Analog zeigt man

Lemma 1.16. *In einem archimedischen Körper K konvergiert im Fall $|q| < 1$ die geometrische Reihe $s_n = \sum_{i=0}^n c \cdot q^i$ gegen den Grenzwert $\frac{c}{1-q}$.*

1.6 Vollständige metrische Räume

Definition 1.17. *Ein metrischer Raum (X, d) heißt **vollständig**, wenn jede Cauchyfolge in (X, d) konvergiert.*

Definition 1.18. *Ein metrischer Raum (X, d) heißt **folgenkompakt**, wenn jede Folge aus (X, d) eine konvergente Teilfolge besitzt.*

Ein folgenkompakter metrischer Raum ist automatisch vollständig, denn eine Cauchyfolge x_n konvergiert gegen x genau dann wenn eine Teilfolge der Cauchyfolge gegen x konvergiert. [Benutze $d(x, x_m) \leq d(x, x_n) + d(x_n, x_m)$ für $m \geq n$ und geeignete x_n aus der Teilfolge.]

Definition 1.19. *Eine Teilmenge A eines metrischen Raumes (X, d) heißt **abgeschlossen**, wenn gilt: Für jede in (X, d) konvergente Folge $x_n \in A$ mit Grenzwert $x = \lim_{n \rightarrow \infty} x_n \in X$ gilt $x \in A$. Der Abschluß $\overline{A}$ einer Menge $A \subseteq X$ ist die kleinste abgeschlossene Menge in X , welche A enthält (d.h. der Durchschnitt aller abgeschlossenen Mengen Y mit $A \subseteq Y \subseteq X$).*

Beispiel 1.20. Die Intervalle $[a, b]$, oder auch $[a, \infty) = \{x \in K | x \geq a\}$ oder $(-\infty, a] = \{x \in K | x \leq a\}$, sind abgeschlossene Teilmengen in K .

Beweis. Wir zeigen pars pro toto, daß für eine Folge $x_0, x_1, \dots$ von Zahlen in K mit dem Grenzwert x gilt: Aus $x_n \geq a$ für $n = 0, 1, \dots$ folgt auch $x \geq a$. Dies sieht man wie folgt: Wäre $x < a$, dann gilt $d(x_n, x) < \varepsilon$ für fast alle n bei Wahl von $\varepsilon := a - x > 0$. Andererseits gilt dann aber auch

$$d(x_n, x) = x_n - x = \underbrace{x_n - a}_{\geq 0} + \underbrace{a - x}_{=\varepsilon} \geq \varepsilon$$

für alle n , und wir erhalten einen Widerspruch. □

Analog sind Quader der Gestalt $A = [a_1, b_1] \times \dots \times [a_r, b_r]$ abgeschlossene Teilmengen des Euklidischen Raumes $\mathbb{R}^r$.

Satz 1.21. Jede abgeschlossene Teilmenge A eines vollständigen metrischen Raumes (X, d_X) versehen mit der eingeschränkten Metrik ist ein vollständiger metrischer Raum (A, d_X) .

Beweis. Sei x_n eine Cauchyfolge in (A, d) . Dann ist per Definition x_n eine Cauchyfolge in (X, d_X) . Nach Annahme existiert also $x = \lim_{n \rightarrow \infty} x_n$ in (X, d_X) . Weil A abgeschlossen ist, gilt $x \in A$. Also ist per definitionem $x \in A$ der Grenzwert von x_n in (A, d_X) . □

Satz 1.22. Jede folgenkompakte Teilmenge A eines metrischen Raumes (X, d_X) ist beschränkt und abgeschlossen in (X, d_X) .

Beweis. Wäre A nicht beschränkt, gäbe es eine Folge $x_1, x_2, \dots$, aus A mit $d_X(x_0, x_n) \geq n$. Für jede Teilfolge $\tilde{x}_n$ einer solchen Folge gilt erst recht $d_X(x_0, \tilde{x}_n) \geq n$. Somit besäße x_n keine konvergente (und damit beschränkte) Teilfolge. Ein Widerspruch zur Folgenkompaktheit von A !

Um zu zeigen, daß A abgeschlossen ist, betrachten wir eine beliebige Folge x_n aus A mit Grenzwert x in (X, d_X) . Jede Teilfolge $\tilde{x}_n$ der Folge x_n konvergiert gegen den Grenzwert x in (X, d_X) . Nach Annahme ist (A, d_X) folgenkompakt. Somit existiert eine konvergente Teilfolge $\tilde{x}_n$ der Folge x_n mit einem Grenzwert a in A . Aus der Eindeutigkeit des Grenzwertes (Lemma 1.14) folgt $x = a$. Somit ist $x \in A$, d.h. A ist eine abgeschlossene Teilmenge von (X, d_X) . □

1.7 Der Banachsche Fixpunktsatz

Satz 1.23. Sei (X, d) ein vollständiger metrischer Raum und $F: X \rightarrow X$ eine **kontraktive** Abbildung eines metrischen Raumes (X, d) in sich, d. h. es gebe eine reelle Konstante $0 < \kappa < 1$ mit

$$d(F(x), F(y)) \leq \kappa \cdot d(x, y)$$

1 Der Konvergenzbegriff

für alle $x, y \in X$. Dann besitzt die Abbildung F einen eindeutig bestimmten Fixpunkt $x \in X$, d. h. einen eindeutig bestimmten Punkt x mit der Eigenschaft

$$\boxed{F(x) = x}.$$

Beweis. Wir müssen die Existenz und die Eindeutigkeit des Fixpunktes $x \in X$ zeigen. Wir wollen mit der Eindeutigkeit beginnen. Seien x und ξ Fixpunkte von F . Aus der Kontraktivität $d(F(x), F(\xi)) \leq \kappa \cdot d(x, \xi)$ und der Fixpunkteigenschaft $F(x) = x, F(\xi) = \xi$ folgt

$$d(x, \xi) \leq \kappa \cdot d(x, \xi).$$

Wäre $x \neq \xi$, könnte man durch $d(x, \xi) > 0$ teilen und erhielte den Widerspruch $1 \leq \kappa$.

Nun zeigen wir die Existenz. Wähle hierzu ein beliebiges $x_0 \in X$ und setze $x_1 = F(x_0)$, $x_2 = F(x_1), \dots, x_n = F^n(x_0)$ als Folge in (X, d) . Diese Folge ist beschränkt, denn

$$\begin{aligned} d(x_0, x_n) &\leq d(x_0, x_1) + d(x_1, x_2) + \dots + d(x_{n-1}, x_n) \\ &= d(x_0, x_1) + d(F(x_0), F(x_1)) + \dots + d(F^{n-1}(x_0), F^{n-1}(x_1)) \\ &\leq d(x_0, x_1) + \kappa d(x_0, x_1) + \dots + \kappa^{n-1} d(x_0, x_1) \\ &\leq \frac{d(x_0, x_1)}{1 - \kappa} = C. \end{aligned}$$

Als nächstes zeigen wir, daß x_n eine Cauchyfolge ist. Sei hierzu oBdA $m \geq n$. Dann ist

$$d(x_n, x_m) = d(\underbrace{F^n(x_0)}_{=x_n}, \underbrace{F^n(x_{m-n})}_{=x_m}) \leq \kappa^n \underbrace{d(x_0, x_{m-n})}_{\leq C} \leq \kappa^n C.$$

Da $\kappa^n \xrightarrow{n \rightarrow \infty} 0$, folgt daraus wie behauptet $d(x_n, x_m) < \varepsilon$ falls $m \geq n \geq N(\varepsilon)$. Die Cauchyfolge x_n konvergiert gegen einen Grenzwert $x \in X$, denn nach Annahme ist (X, d) vollständig. Es bleibt die Fixpunkteigenschaft von x zu zeigen. Hierzu stellen wir fest

$$\begin{aligned} d(x, F(x)) &\leq d(x, x_n) + d(x_n, F(x)) \leq d(x, x_n) + \kappa d(x_{n-1}, x) \\ &\leq d(x, x_n) + d(x, x_{n-1}) < \frac{1}{2}\varepsilon + \frac{1}{2}\varepsilon < \varepsilon \end{aligned}$$

für $n \geq N(\frac{1}{2}\varepsilon)$, resp. $n - 1 \geq N(\frac{1}{2}\varepsilon)$. Ein solches $n \in \mathbb{N}$ existiert natürlich. Also gilt $d(x, F(x)) < \varepsilon$ für alle $\varepsilon > 0$. Mit anderen Worten: Es gilt $d(x, F(x)) = 0$ bzw. $F(x) = x$. $\square$

Eine Anwendung. Sei $0 \leq \eta \leq 1/4$ gegeben in einem vollständigen archimedischen Körper K . Wähle ein ε in K mit $0 < \varepsilon \leq \eta$. Dann ist

$$X = \left[0, \frac{1}{2} - \varepsilon \right]$$

abgeschlossen im metrischen Raum (K, d) . Versehen mit der Metrik $d(x, y) = |x - y|$ von (K, d) ist (X, d) daher vollständig. Die Abbildung

$$F(x) = x^2 + \frac{1}{4} - \eta$$

ist kontraktiv auf X ; beachte $d(F(x), F(y)) = |x^2 - y^2| = |x + y| \cdot d(x, y) \leq \kappa \cdot d(x, y)$ mit $|x + y| \leq \kappa := 2(\frac{1}{2} - \varepsilon) < 1$. Weiterhin gilt $F(x) \geq 0$ und F ist monoton auf $[0, \frac{1}{2} - \varepsilon]$ mit $F(\frac{1}{2} - \varepsilon) = \frac{1}{2} - \varepsilon - (\eta - \varepsilon^2) \leq \frac{1}{2} - \varepsilon$, denn $0 < \varepsilon < 1$ impliziert $\varepsilon^2 \leq \varepsilon \leq \eta$. Folglich ist

$$F: X \rightarrow X$$

wohldefiniert wegen $F(X) = F([0, \frac{1}{2} - \varepsilon]) \subseteq X$. Dies zeigt

Lemma 1.24. *Ein vollständiger archimedischer Körper ist pythagoräisch.*

Beweis. Nach der Bemerkung auf Seite 3 genügt es zu zeigen: Jedes $\eta \in [0, \frac{1}{4}]$ ist ein Quadrat in K . Aus dem Banachschen Fixpunktsatz folgt die Existenz eines Fixpunktes $x \in X$ von F . Aus $F(x) = x \implies (x - \frac{1}{2})^2 = \eta$ folgt dann $\eta = \xi^2$ für $\xi = x - \frac{1}{2} \in K$. $\square$

1.8 Quaderschachtelung

Sei K ein pythagoräischer Körper. Im Euklidischen Standardraum K^r der Dimension r gelten die **Schachtelungs-Ungleichungen**

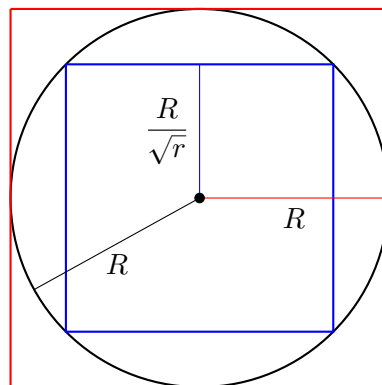
$$\max_{i=1,\dots,r} |x_i| \leq \|x\| \leq \sqrt{r} \cdot \max_{i=1,\dots,r} |x_i|$$

für einen Vektor $x = (x_1, \dots, x_r) \in K^r$. Beachte: $\max_{i=1,\dots,r} |x_i|$ definiert auch eine Norm auf K^r , die sogenannte **Quadernorm**. Die Formel läßt sich durch Quadrieren beweisen, denn

$$\max_{i=1,\dots,r} |x_i|^2 \leq |x_1|^2 + \dots + |x_r|^2 \leq r \cdot \max_{i=1,\dots,r} |x_i|^2$$

gilt aus offensichtlichen Gründen.

Was bedeutet dies anschaulich? Die Norm $\|\cdot\|$ gibt den Abstand eines Punktes von Null an. Wir betrachten die „Kugel“ aller Punkte mit Abstand kleinergleich R vom Ursprung. Diese Kugel B liegt in einem Quader mit der Seitenlänge $2R$. Umgekehrt liegt der Quader mit der Seitenlänge $\sqrt{r}^{-1} \cdot 2R$ in der Kugel B .



1 Der Konvergenzbegriff

Sei nun

$$Q = [c, d]^r = \underbrace{[c, d] \times \cdots \times [c, d]}_{r \text{ mal}}.$$

ein würfelförmiger Quader im K^r mit der Seitenlänge $l(Q) = |d - c|$.

Lemma 1.25. Für beliebige Punkte ξ, η aus einem Quader Q mit den Seitenlängen $l(Q)$ gilt

$$d(\xi, \eta) \leq l(Q)\sqrt{r}.$$

Beweis. Durch Verschieben des Quaders kann man o. B. d. A. annehmen $\eta = 0$. Dann gilt $d(\xi, \eta) = \|\xi\| \leq \sqrt{r} \cdot \max_{i=1, \dots, r} |\xi_i|$ und es genügt zu zeigen $|\xi_i| \leq |d - c|$. Wegen $c \leq \xi_i \leq d$ und $c \leq 0 \leq d$ folgt aber $|\xi_i| \leq \max(d, -c) \leq d - c = |d - c|$. $\square$

Satz 1.26 (Bolzano-Weierstraß). Jede beschränkte Folge im Euklidischen Raum $(K^r, \|\cdot\|)$ besitzt eine Teilfolge, die eine Cauchyfolge ist.

Beweis. Ist die Folge $x_0, x_1, \dots$ beschränkt in K^r , so liegt sie in einer Kugel und damit auf Grund der Schachtelungsungleichungen in einem geeigneten Quader

$$Q = [a, b]^r = \underbrace{[a, b] \times \cdots \times [a, b]}_{r \text{ mal}}.$$

Hierbei ist $[a, b] = \{x \in K \mid a \leq x \leq b\}$ ein geeignetes Intervall in K . Man teilt nun den Quader in 2^r Teilquader, indem man jedes der Intervalle $[a, b]$ in zwei gleich lange Teile unterteilt. In mindestens einem der Teilquader müssen unendlich viele Folgenglieder sein. Dies liefert eine Teilfolge $x_{1,0}, x_{1,1}, x_{1,2}, \dots$, die vollständig in einem der Teilquader liegt. Dieses Verfahren setzt man iterativ fort und erhält eine absteigende Folge von Quadern:

$$Q_0 \supseteq Q_1 \supseteq Q_2 \supseteq \cdots$$

Hierbei ist $Q = Q_0$.

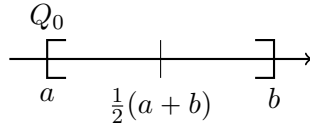
In k -ten Quader Q_k liegt dann vollständig enthalten die Teilfolge $x_{k,0}, x_{k,1}, x_{k,2}, \dots$ der Folge $x_{k-1,0}, x_{k-1,1}, x_{k-1,2}, \dots$ usw.

Diese Folgen ordnet man nun in einer Tabelle an:

$$\begin{array}{cccc} x_{0,0} & x_{0,1} & x_{0,2} & \cdots \\ x_{1,0} & x_{1,1} & x_{1,2} & \cdots \\ x_{2,0} & x_{2,1} & x_{2,2} & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{array}$$

Diagonalfolgentrick. Man bildet jetzt die Diagonalfolge $\xi_k := x_{k,k}$ für alle $k \in \mathbb{N}$ und betrachtet die dadurch neu entstandene Folge $\xi_0, \xi_1, \xi_2, \dots$. Dies ist eine Teilfolge der ursprünglichen Folge $x_0, x_1, \dots$, und es gilt $\xi_i \in Q_n$ für alle $i \geq n$. Aus $i, j \geq n$ folgt daher $\xi_i, \xi_j \in Q_n$.

Die Seitenlänge der Quader Q_n halbiert sich mit jeder Unterteilung. Wir zeigen dies oBdA im Fall $n = 0$. Es gilt $l(Q_0) = |b - a|$.



Wie in der Zeichnung angedeutet sei Q_1 eine der beiden Teilhälften von Q_0 . Wegen

$$l(Q_0) = \begin{cases} |b - \frac{a+b}{2}| = |b-a|/2, & \text{falls } Q_1 = [\frac{a+b}{2}, b] \\ |\frac{a+b}{2} - a| = |b-a|/2, & \text{falls } Q_1 = [a, \frac{a+b}{2}] \end{cases}$$

ist die Länge von Q_1 in beiden möglichen Fällen $l(Q_1) = l(Q_0)/2$, halbiert sich also bei der Teilung. Durch Induktion folgt daher $l(Q_n) = 2^{-n} \cdot l(Q_0) = 2^{-n} \cdot |b-a|$.

Aus Lemma 1.25 folgt für beliebige Punkte $\xi_i, \xi_j \in Q_n$ dann mit der Konstante $C = |b-a|\sqrt{r}$ die Ungleichung $d(\xi_i, \xi_j) \leq \frac{C}{2^n}$. Wir erstellen ein vorläufiges Resümee: Wir haben eine Teilfolge $\xi_0, \xi_1, \dots$ der gegebenen Folge $x_0, x_1, \dots$ konstruiert, so daß für eine feste positive Konstante C in K gilt

$$d(\xi_i, \xi_j) \leq \frac{C}{2^n},$$

für alle $i, j \geq n$. Wir wollen daraus ableiten, daß ξ_i eine Cauchyfolge ist. Für beliebiges $\varepsilon > 0$ müssen wir zeigen, daß ein $N = N(\varepsilon)$ existiert mit

$$d(\xi_i, \xi_j) < \varepsilon$$

für $i, j \geq N(\varepsilon)$. Dazu genügt es $N \in \mathbb{N}$ wählen zu können mit

$$\frac{C}{2^n} < \varepsilon$$

für alle $n \geq N$. Die Existenz einer solchen Zahl $N \in \mathbb{N}$ folgt aus dem Lemma 1.15 für die Wahl $q = 1/2$. $\square$

Satz 1.27. *In einem pythagoräischen Körper ist jede monoton wachsende, nach oben beschränkte Folge eine Cauchyfolge. Dies gilt ebenso für jede monoton fallende, nach unten beschränkte Folge.*

Beweis. Wir betrachten nur den Fall der monoton wachsenden, nach oben durch eine Konstante C beschränkten Folgen. Der umgekehrte Fall ist völlig analog. Aus der Monotonie

$$x_n \leq x_{n+1}$$

der Folge folgt $x_n \in [x_0, C]$ für alle n . Also ist die Folge x_n beschränkt, und nach Satz 1.26 existiert damit eine Teilfolge $\tilde{x}_n$ der Folge, welche eine Cauchyfolge ist. D.h. für beliebiges $\varepsilon > 0$ existiert ein $\tilde{N}(\varepsilon)$ mit

$$d(\tilde{x}_i, \tilde{x}_j) = \tilde{x}_i - \tilde{x}_j < \varepsilon, \quad \forall i, j \geq \tilde{N}(\varepsilon)$$

1 Der Konvergenzbegriff

wobei hier stillschweigend $i \geq j$ angenommen werden kann. Wähle nun $N = N(\varepsilon) \geq \tilde{N}(\varepsilon)$ so groß, daß für alle $j \geq N = N(\varepsilon)$ gilt

$$x_j \leq \tilde{x}_{\tilde{N}(\varepsilon)} \quad \text{und damit} \quad -\tilde{x}_{\tilde{N}(\varepsilon)} \leq -x_j .$$

Die Existenz einer solchen Zahl N folgt aus der Monotonie der Folge x_j und der Tatsache, daß $\tilde{x}_i$ eine Teilfolge der Folge x_i ist. Dies impliziert natürlich auch für beliebiges i

$$x_i \leq \tilde{x}_i .$$

Für $i \geq j \geq N(\varepsilon)$ gilt daher nach Wahl von $N(\varepsilon)$ und wegen der Cauchyfolgeeigenschaft der Teilfolge $\tilde{x}_i$

$$d(x_i, x_j) = x_i - x_j \leq \tilde{x}_i - \tilde{x}_{\tilde{N}(\varepsilon)} < \varepsilon .$$

□

1.9 Reelle Zahlen

Konvergente Folgen sind Cauchyfolgen, wie wir gesehen haben, aber nicht jede Cauchyfolge ist konvergent. Beispielsweise ist $\mathbb{Q}$ ein archimedischer Körper, der aber nicht vollständig ist.

Jetzt kommt der Zeitpunkt, an dem wir permanent zu den reellen Zahlen übergehen wollen.

Definition 1.28. Wir fixieren ein für alle mal einen archimedischen vollständigen Körper¹ und nennen ihn den **Körper der reellen Zahlen**. Wir bezeichnen diesen Körper mit $\mathbb{R}$. Folgen in $\mathbb{R}$ nennen wir **reelle Folgen**.

Als vollständiger archimedischer Körper ist $\mathbb{R}$ pythagoräisch. Da $\mathbb{R}$ per Definition vollständig ist, sind in $\mathbb{R}$ die Begriffe der Cauchyfolge und der konvergenten Folge äquivalent. Solche Folgen sind immer beschränkt und haben einen eindeutig bestimmten Grenzwert. Liegt die Folge in einem abgeschlossenen Intervall I , so ist auch ihr Grenzwert in I enthalten. Weiterhin enthält jede reelle beschränkte Folge eine konvergente Teilfolge und jede monotone beschränkte Folge konvergiert. Zuletzt besitzt noch jede nach oben beschränkte Teilmenge $X \subseteq \mathbb{R}$ eine kleinste obere Schranke $\sup(X)$. Für den Beweis der letzten Aussage sei auf den nächsten Abschnitt verwiesen. Aus den Sätzen Satz 1.22, Satz 1.21 und Satz 1.26 folgt schliesslich die fundamentale Aussage

Satz 1.29. Eine Teilmenge A des Euklidischen Raumes $\mathbb{R}^r$ ist folgenkompakt genau dann, wenn sie beschränkt und abgeschlossen ist.

Wir wollen nun die Darstellung reeller Zahlen durch Dezimalbrüche betrachten. Nach dem Archimedischen Axiom ist jede nicht negative reelle Zahl y kleiner als eine geeignete natürliche

¹Die Existenz eines solchen Körpers kann man aus den Peano Axiomen für die natürlichen Zahlen ableiten.

Zahl n . Da nur endlich viele natürliche Zahlen vor n liegen, kann man obdA $n - 1 \leq y < n$ annehmen. Dann liegt $x = y - (n - 1)$ im Intervall $I_0 = [0, 1)$. Teilt man I_0 in 10 Teilintervalle, folgt analog

$$x \in I_1 = \left[\frac{a_0}{10}, \frac{a_0 + 1}{10} \right)$$

für ein $a_0 \in \{0, 1, \dots, 9\}$. Unterteilt man I_1 wieder in 10 Teilintervalle und fährt so fort, erhält man eine Approximation von x durch Zahlen

$$x_n := \frac{a_0}{10} + \frac{a_1}{100} + \dots + \frac{a_{n-1}}{10^n}$$

mit $a_0, a_1, \dots, a_{n-1} \in \{0, 1, \dots, 9\}$ und

$$x_n \leq x \leq x_n + \frac{1}{10^n}.$$

Daraus folgt $d(x, x_n) \leq \frac{1}{10^n}$, die Folge der x_n konvergiert also gegen die gegebene Zahl x (Lemma 1.15). Man nennt dies die **Dezimalbruchentwicklung** von x und schreibt (bekanntlich)

$$x \text{ „}=\text{“ } 0, a_0 a_1 a_2 \dots$$

Umgekehrt definiert jede solche Dezimalbruchentwicklung eine reelle Zahl im Intervall $[0, 1]$, denn die dadurch definierte Folge rationaler Zahlen

$$x_n := \frac{a_0}{10} + \frac{a_1}{100} + \dots + \frac{a_{n-1}}{10^n}$$

ist monoton wachsend

$$x_n \leq x_{n+1}.$$

Es gilt $x_n \in [0, 1]$, denn $x_n \leq \frac{9}{10} + \frac{9}{100} + \dots + \frac{9}{10^n} = \frac{9}{10} \cdot \frac{1-10^{-n}}{1-10^{-1}} = 1 - 10^{-n} \leq 1$. Da die Folge x_n monoton wachsend und nach oben beschränkt ist (hier durch 1), konvergiert sie nach Satz 1.27 und ihr Grenzwert x liegt im Intervall $I = [0, 1]$. Beachte $x \in [0, 1)$ außer im Fall $0,9999\dots$

1.10 Infimum und Supremum

Definition 1.30. Sei $X \subseteq \mathbb{R}$ eine nach oben beschränkte, nichtleere Teilmenge. Man nennt

$$Y := \{y \in \mathbb{R} \mid x \leq y \text{ für alle } x \in X\}$$

die Menge der **oberen Schranken** von X . Analog ist im Falle einer nach unten beschränkten Menge die Menge ihrer **unteren Schranken** definiert.

Es bezeichne $Y^c := \mathbb{R} \setminus Y$ das Komplement einer Teilmenge Y von $\mathbb{R}$.

1 Der Konvergenzbegriff

Bemerkung 1.31. Für die Menge Y der oberen Schranken einer nach oben beschränkten nicht leeren Menge $X \subseteq \mathbb{R}$ gelten folgende Eigenschaften:

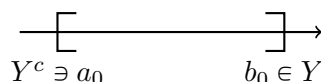
$$(1) Y \neq \emptyset$$

(2) Y ist abgeschlossen in $\mathbb{R}$.

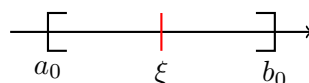
(3) Für $\xi \in Y^c$ gilt $\xi \leq y$ für alle $y \in Y$. Insbesondere gilt für konvergente Folgen $a_n \in Y^c$ mit Limes a daher $a \leq y$ für alle $y \in Y$ (denn $(-\infty, y]$ ist abgeschlossen).

Beweis. (1) gilt nach Annahme. Zum Beweis von (2) sei eine Folge $y_n \in Y$ gegeben mit $\lim_{n \rightarrow \infty} y_n = y$. Für $x \in X$ gilt $x \leq y_n$, d.h. $y_n \in [x, \infty)$ für alle $n \in \mathbb{N}$. Folglich ist $y \in [x, \infty)$ und damit $y \in Y$, denn $x \leq y$ gilt für alle $x \in X$. (3) $\xi \in Y^c$ bedeutet $\xi \notin Y$. Daraus folgt nach Definition von Y , daß es ein $x \in X$ gibt mit $\xi < x$. Ebenfalls nach Definition von Y gilt aber $x \leq y$ für alle $x \in X$ und $y \in Y$. Es folgt $\xi < x \leq y$ und damit $\xi \leq y$ für $y \in Y$. $\square$

Wähle ein $a_0 \notin Y$ und ein $b_0 \in Y$. Ein solches b_0 existiert nach (1) und a_0 existiert wegen $X \neq \emptyset$ (z.B. $a_0 = x - 1$ für ein $x \in X$):



Setze nun durch Halbieren des Intervalls $\xi := \frac{a_0 + b_0}{2}$



und führe eine Fallunterscheidung durch: Im Falle $\xi \in Y$ setze $b_1 := \xi$ und $a_1 := a_0$, im Fall $\xi \in Y^c$ setze $a_1 := \xi$ und $b_1 := b_0$. Iteriert man dies für alle $n \in \mathbb{N}$, so erhält man

$$Y^c \ni a_n \leq b_n \in Y$$

und nach Konstruktion gilt

$$a_0 \leq a_1 \leq a_2 \leq \dots \leq a_n \leq b_n \leq \dots \leq b_2 \leq b_1 \leq b_0.$$

Somit ist $a_0, a_1, \dots$ eine monoton steigende, nach oben beschränkte Folge und analog $b_0, b_1, \dots$ eine monoton fallende nach unten beschränkte Folge. Beide Folgen konvergieren also nach Satz 1.27. Setzt man $a = \lim_{n \rightarrow \infty} a_n$ und $b = \lim_{n \rightarrow \infty} b_n$, dann gilt

$$a_0 \leq \dots \leq a_n \leq \dots \leq a \leq b \leq \dots \leq b_n \leq \dots \leq b_0.$$

[Es gilt $a_j < b_k$ für alle $j, k \in \mathbb{N}$. Daraus folgt $a \leq b_k$ für alle k , und im Limes $k \rightarrow \infty$ schliesslich $a \leq b$.] Wir behaupten nun $a = b$. Zum Beweis fixieren wir ein $\varepsilon > 0$. Es gilt dann

$$0 \leq b - a \leq \underbrace{|b - b_n|}_{< \frac{1}{3}\varepsilon} + \underbrace{|b_n - a_n|}_{< \frac{1}{3}\varepsilon} + \underbrace{|a_n - a|}_{< \frac{1}{3}\varepsilon} < \varepsilon$$

für $n \geq N(\varepsilon)$. Dies folgt aus den Konvergenzaussagen $a_n \xrightarrow{n \rightarrow \infty} a$, $b_n \xrightarrow{n \rightarrow \infty} b$ sowie $|b_n - a_n| = 2^{-n}|b_0 - a_0| \xrightarrow{n \rightarrow \infty} 0$, wie im Beweis von Satz 1.26. Da $\varepsilon > 0$ beliebig gewählt werden kann, folgt daraus $a = b$.

Auf Grund der obigen Intervall Schachtelungen gilt weiterhin

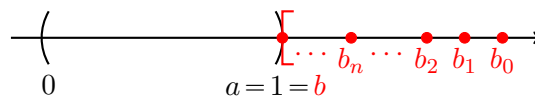
$$(4) \quad a \leq y \text{ für alle } y \in Y,$$

denn $a_n \in Y^c$ impliziert $a_n \leq y$ für alle $y \in Y$ und somit $a \leq y$ nach (3). Aus (2) folgt wegen $b_n \in Y$ dann

$$(5) \quad b = \lim_{n \rightarrow \infty} b_n \in Y.$$

Damit ist $a = b$ die kleinste obere Schranke von X nach (4) und (5).

Beispiel: $X = (0, 1)$ mit $\sup(X) := a = b = 1$



Dies zeigt

Satz 1.32. Jede nach oben beschränkte nichtleere Teilmenge $X \subseteq \mathbb{R}$ besitzt eine kleinste obere Schranke. Diese nennt man das **Supremum** $\sup(X)$ der Menge X .

Für die Punkte a_n gilt nach Konstruktion $a_n < a$, und daher wegen $a = \sup(X)$ also $a_n \notin Y$. Da a_n somit keine obere Schranke von X ist, gibt es ein $x_n \in X$ mit $a_n \leq x_n$. Aus der Ungleichung $a_n \leq x_n \leq a = \sup(X)$ folgt wegen $a_n \rightarrow a$ dann recht einfach $x_n \rightarrow a$. Also

Zusatz. Es gibt eine Folge in X , welche gegen $a = \sup(X)$ konvergiert.

Oft wird das Supremum formal auch für den Fall einer nach oben unbeschränkten Menge definiert. In diesem Falle schreibt man

$$\sup(X) = +\infty$$

für den Fall, daß die Menge X nach oben nicht beschränkt ist.

Völlig analog zum Supremum einer nach oben beschränkten Menge, existiert das **Infimum** einer nach unten beschränkten Menge $\inf(X) = -\sup(-X)$.

Satz 1.33. Jede nach unten beschränkte nichtleere Menge $X \subseteq \mathbb{R}$ besitzt ein Infimum $\inf(X)$, eine größte untere Schranke von X .

Sei nun $x_1 \leq x_2 \leq x_3 \leq \dots$ eine beliebige monoton steigende Folge reeller Zahlen. Dann ist entweder $X = \{x_n \mid n \in \mathbb{N}\}$ nach oben beschränkt und die Folge x_n konvergiert monoton gegen $\sup(X) \in \mathbb{R}$. Oder die Menge X ist nicht nach oben beschränkt und $\sup(X) = +\infty$. In diesem Fall konvergiert die Folge gegen $+\infty$ in folgendem Sinn: Für alle $C > 0$ existiert ein $N = N(C)$ so daß für alle $n \geq N$ gilt $x_n > C$.

1 Der Konvergenzbegriff

Prinzip der monotonen Konvergenz. Diese Betrachtungen kann man wie folgt zusammen fassen. Für $\mathbb{R}^+ = \mathbb{R} \cup \{+\infty\}$ gilt: *Jede (!) monoton steigende Folge x_n mit Werten in $\mathbb{R}^+$ besitzt einen Grenzwert x in $\mathbb{R}^+$ im obigen Sinne.* Wir schreiben auch

$$x_n \nearrow x$$

oder benutzen häufig die Schreibweise $x = \sup_n x_n = \sup x_n$ für den Grenzwert in diesem erweiterten Sinn. Analog existiert immer ein Grenzwert $\inf x_n$ in $\mathbb{R}^- = \mathbb{R} \cup \{-\infty\}$ für eine monoton fallende Folge $x_n \in \mathbb{R}^-$.

2 Stetige Abbildungen

2.1 Stetigkeit

Definition 2.1. Seien (X, d_X) und (Y, d_Y) metrische Räume. Eine Abbildung

$$f: (X, d_X) \rightarrow (Y, d_Y)$$

heißt **stetig**, wenn für jede Folge x_n in X gilt:

$$\boxed{x_n \rightarrow \xi \implies f(x_n) \rightarrow f(\xi)} .$$

Anders formuliert: f führt konvergente Folgen in konvergente Folgen über und vertauscht mit Limesbildung, d.h.

$$\boxed{f\left(\lim_{n \rightarrow \infty} x_n\right) = \lim_{n \rightarrow \infty} f(x_n)} .$$

Ein Spezialfall. Gibt es für $f: (X, d_X) \rightarrow (Y, d_Y)$ eine feste Konstante $C > 0$ mit

$$d_Y(f(x), f(\xi)) \leq C \cdot d_X(x, \xi) \quad , \quad \forall x, \xi \in X ,$$

nennt man f **Lipschitz-stetig** und C nennt man **Lipschitz-Konstante**. Jede Lipschitz-stetige Abbildung ist stetig, denn für $x_n \rightarrow \xi$ gilt $d_X(x_n, \xi) < \varepsilon/C$ für genügend großes n und damit $d_Y(f(x_n), f(\xi)) < \varepsilon$. Jede kontraktive Abbildung ist Lipschitz-stetig für $C = \kappa < 1$. Weitere Beispiele für Lipschitz-stetige Abbildungen sind:

Beispiel 2.2. (1) Die identische Abbildung $\text{id}_X: (X, d_X) \rightarrow (X, d_X)$, also $\text{id}_X(x) = x$ mit der Lipschitz-Konstante $C = 1$.

(2) Die konstante Funktion $f: (X, d_X) \rightarrow (Y, d_Y)$ definiert durch $f(x) = y_0$ für einen festen Punkt $y_0 \in Y$. Hier kann $C > 0$ beliebig gewählt werden.

(3) Sei (X, d_X) ein metrischer Raum und $\emptyset \neq A \subseteq X$ eine abgeschlossene Teilmenge. Sei $f(x) = d(x, A) := \inf_{a \in A} d(x, a)$ die **Abstandsfunktion** $f: (X, d_X) \rightarrow Y = \mathbb{R}$ zu A . Beachte $d(x, A) = 0 \iff x \in A$, da A abgeschlossen ist. Es gilt $d_Y(f(x), f(\xi)) \leq C \cdot d_X(x, \xi)$ für $C = 1$. [Für $\xi \in X$ und reelles $\varepsilon > 0$ gibt es ein $a \in A$ mit $d(\xi, a) < d(\xi, A) + \varepsilon$. Für alle $x \in X$ gilt $d(x, A) \leq d(x, a)$, also $f(x) - f(\xi) \leq d_X(x, a) - d_X(\xi, a) + \varepsilon \leq d_X(x, \xi) + \varepsilon$. Aus der Symmetrie in x, ξ folgt $d_Y(f(x), f(\xi)) = |f(x) - f(\xi)| \leq d_X(x, \xi) + \varepsilon$ für alle $\varepsilon > 0$ und damit die Behauptung.]

2 Stetige Abbildungen

(4) Jede $\mathbb{R}$ -lineare Abbildung $L: \mathbb{R}^r \rightarrow \mathbb{R}^s$ zwischen Euklidischen Räumen.

(5) Insbesondere als Spezialfall von (4) die Koordinatenprojektionen

$$p_i: \mathbb{R}^r \rightarrow \mathbb{R}, \quad p_i(x_1, \dots, x_r) = x_i$$

für $i = 1, \dots, r$.

Beweis der Lipschitz-Stetigkeit im Fall (4). Es gilt $d_{\mathbb{R}^s}(L(x), L(\xi)) = \|L(x) - L(\xi)\|_{\mathbb{R}^s} = \|L(x - \xi)\|_{\mathbb{R}^s}$ für lineare Abbildungen $L(y) = (\sum_{j=1}^r L_{ij}y_j)_{i=1,\dots,s}$. Wir suchen eine Konstante C mit der Eigenschaft

$$\|L(y)\|_{\mathbb{R}^s} \leq C \cdot \|y\|_{\mathbb{R}^r}.$$

Es gilt $\|L(y)\|_{\mathbb{R}^s} \leq \sqrt{s} \cdot \max_{i=1,\dots,s} |L(y)_i|$ wegen der Quaderungleichung. Es folgt $\|L(y)\|_{\mathbb{R}^s} \leq r\sqrt{s} \cdot \max_{i=1,\dots,s; j=1,\dots,r} |L_{ij}| \cdot \max_{j=1,\dots,r} |y_j| \leq C \|y\|_{\mathbb{R}^r}$ für $C = r\sqrt{s} \cdot \max_{i,j} |L_{ij}|$. Diese Wahl von C ist nicht optimal, aber offensichtlich gilt

Lemma 2.3. Für $\mathbb{R}$ -lineare Abbildungen $L: \mathbb{R}^r \rightarrow \mathbb{R}^s$ ist

$$\|L\| := \sup_{v \neq 0} \left(\frac{\|L(v)\|_{\mathbb{R}^s}}{\|v\|_{\mathbb{R}^r}} \right) \leq C < \infty$$

eine wohldefinierte reelle (!) Zahl ≥ 0 , und es gilt

$$\|L(v)\| \leq \|L\| \cdot \|v\|.$$

Bemerkung 2.4. Eine beliebige Funktion $f: (X, d) \rightarrow \mathbb{R}^r$ schreibt sich in der Form

$$y = f(x) = \begin{pmatrix} f_1(x) \\ \vdots \\ f_r(x) \end{pmatrix} \in \mathbb{R}^r.$$

Die reellwertigen Funktionen $f_i = p_i \circ f$ sind wegen (5) als Zusammensetzung stetiger Funktionen stetig (siehe Korollar 2.7). Umgekehrt definieren je r stetige reellwertige Funktionen $f_1(x), \dots, f_r(x)$ auf (X, d) eine stetige Funktion $(X, d) \rightarrow \mathbb{R}^r$, denn $x_n \rightarrow x$ impliziert $f_i(x_n) \rightarrow f_i(x)$ in $\mathbb{R}$, und wegen $d(f(x_n), f(x)) \leq \sqrt{r} \cdot \max_{i=1,\dots,r} |f_i(x_n) - f_i(x)|$ auch $f(x_n) \rightarrow f(x)$ in $\mathbb{R}^r$.

Definition 2.5. Seien (X, d_X) und (Y, d_Y) metrische Räume. Eine Abbildung

$$f: X \rightarrow Y$$

heißt **stetig im Punkt** ξ von X , wenn für jede Folge x_n , die in (X, d_X) gegen ξ konvergiert, die Bildfolge $f(x_n)$ in (Y, d_Y) gegen $f(\xi)$ konvergiert.

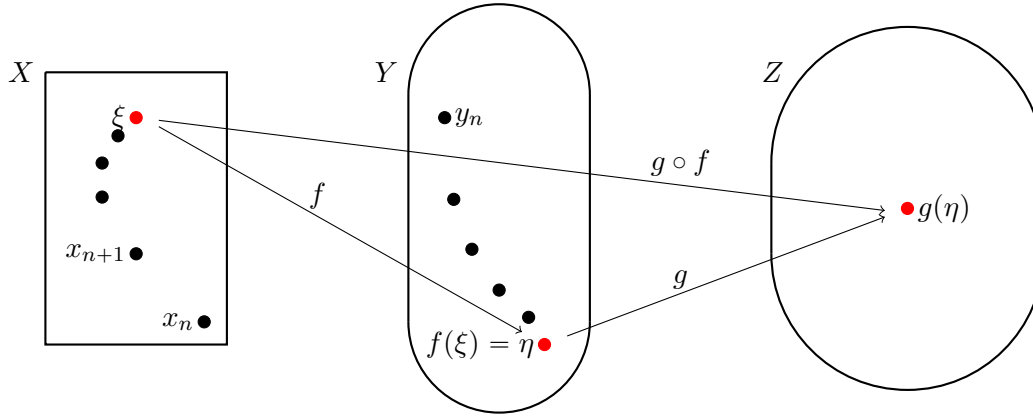
Offensichtlich ist $f: (X, d_X) \rightarrow (Y, d_Y)$ stetig genau dann, wenn f in jedem Punkt ξ von X stetig ist.

Satz 2.6. Sei $f: (X, d_X) \rightarrow (Y, d_Y)$ stetig im Punkt $\xi \in X$ und sei $g: (Y, d_Y) \rightarrow (Z, d_Z)$ stetig im Punkt $\eta = f(\xi) \in Y$. Dann ist die Komposition

$$(g \circ f): (X, d_X) \rightarrow (Z, d_Z)$$

stetig im Punkt $\xi \in X$.

Beweis. Zunächst eine Veranschaulichung



Nach Voraussetzung ist f stetig, also gilt für gegen ξ konvergente Folgen x_n

$$y_n := f(x_n) \rightarrow \eta = f(\xi).$$

Da g stetig ist, gilt wegen der Konvergenz von y_n gegen η

$$g(y_n) = (g \circ f)(x_n) \rightarrow g(\eta) = (g \circ f)(\xi).$$

Es folgt $(g \circ f)(x_n) \rightarrow (g \circ f)(\xi)$, und damit ist $g \circ f$ stetig im Punkt ξ . □

Korollar 2.7. Die Komposition stetiger Abbildungen ist stetig.

2.2 Eigenschaften stetiger Funktionen

Lemma 2.8. Sei $f: (X, d_X) \rightarrow (Y, d_Y)$ eine stetige Abbildung. Ist A abgeschlossen in (Y, d_Y) , dann ist das Urbild $f^{-1}(A)$ abgeschlossen in (X, d_X) .

Beweis. Sei $x_n \rightarrow x$ eine in (X, d_X) konvergente Folge mit $x_n \in f^{-1}(A)$. Zu zeigen ist $x \in f^{-1}(A)$. Da f stetig ist, konvergiert die Bildfolge $y_n = f(x_n) \rightarrow y = f(x)$. Wegen $y_n = f(x_n) \in A$, und da A abgeschlossen ist, folgt $y \in A$ und damit $x \in f^{-1}(y) \subseteq f^{-1}(A)$. □

Lemma 2.9. Sei $f: (X, d_X) \rightarrow (Y, d_Y)$ eine stetige Abbildung. Ist (X, d_X) folgenkompakt, dann ist auch das Bild $f(X)$ versehen mit der Einschränkung der Metrik d_Y folgenkompakt.

2 Stetige Abbildungen

Beweis. Sei y_n eine Folge in $f(X)$. Dann gilt $y_n = f(x_n)$ für eine Urbildfolge $x_n \in X$. Nach Annahme gibt es eine in (X, d_X) konvergente Teilfolge $\tilde{x}_n \rightarrow \tilde{x}$. Aus der Stetigkeit von f folgt, daß die Teilbildfolge $\tilde{y}_n = f(\tilde{x}_n)$ in $(f(X), d_Y)$ gegen $f(\tilde{x})$ konvergiert. $\square$

Satz 2.10. Ist (X, d_X) folgenkompakt, dann hat jede stetige reellwertige Funktion

$$f: (X, d_X) \rightarrow \mathbb{R}$$

ein beschränktes Bild. Ist X nicht leer, werden das Maximum und das Minimum von f auf X als Funktionswerte angenommen.

Beweis. Das Bild $f(X)$ ist folgenkompakt in $\mathbb{R}$ bezüglich der Euklidischen Metrik wegen Lemma 2.9. Nach Satz 1.22 ist dann $f(X)$ beschränkt und abgeschlossen. Letzteres zeigt wegen dem Zusatz von Satz 1.32

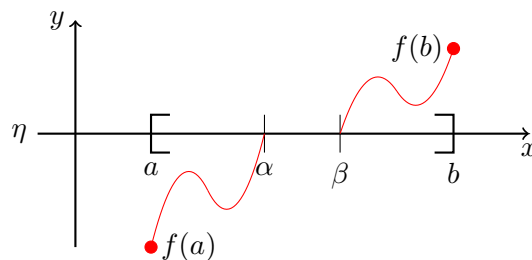
$$\sup(f(X)) = \max(f(X))$$

und analog $\inf(f(X)) = \min(f(X))$. Daraus folgt die Behauptung. $\square$

2.3 Der Zwischenwertsatz

Satz 2.11. Sei $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion mit $f(a) \leq f(b)$. Dann existiert für jedes $\eta \in [f(a), f(b)]$ ein $\alpha \in [a, b]$ mit $f(\alpha) = \eta$.

Beweis. Betrachte die nichtleere beschränkte Menge $A := \{x \in [a, b] \mid f(x) \leq \eta\}$. Sei $\alpha := \sup(A)$ und oBdA $\alpha \neq b$. Dann gilt $x > \alpha \implies f(x) > \eta$. Aus der Stetigkeit von f folgt $f(\alpha) \geq \eta$, da eine monoton fallend gegen α konvergente Folge von Punkten $x \in (\alpha, b]$ existiert mit $f(x) > \eta$. Andererseits gibt es eine monoton steigende Folge von Punkten aus A , welche gegen das Supremum α von A konvergiert (Zusatz 1.32). Aus Stetigkeitsgründen und der Definition von A folgt damit $f(\alpha) \leq \eta$. Beides zusammen ergibt $f(\alpha) = \eta$.



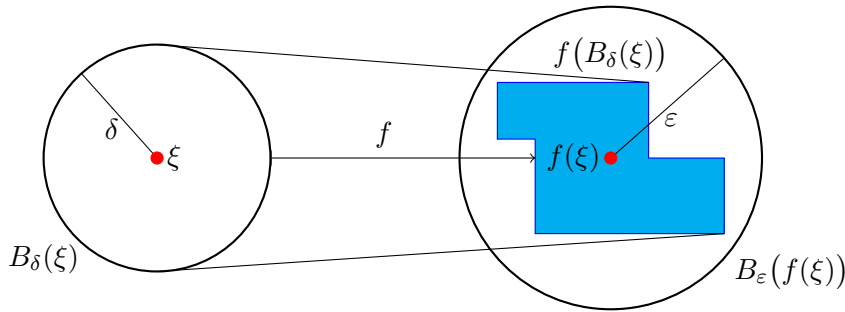
$\square$

2.4 Das ε - δ -Kriterium

Satz 2.12. Gegeben sei eine Funktion $f: (X, d_X) \rightarrow (Y, d_Y)$ und ein $\xi \in X$. Dann ist f genau dann stetig in ξ , wenn zu jedem $\varepsilon > 0$ ein $\delta = \delta(\varepsilon) > 0$ existiert, so daß gilt¹

$$d_X(x, \xi) < \delta \quad \Rightarrow \quad d_Y(f(x), f(\xi)) < \varepsilon.$$

Veranschaulichung. Betrachte die Kugeln $B_\delta(x) = \{x \in X \mid d_X(x, \xi) < \delta\}$. Für beliebiges $\varepsilon > 0$ soll es ein $\delta > 0$ geben, so daß zu jedem $x \in X$, dessen Abstand zu ξ kleiner als δ ist, der Abstand von $f(x)$ zu $f(\xi)$ kleiner als ε ist. f soll also die Kugel $B_\delta(\xi)$ um ξ vom Radius $< \delta$ in die Kugel $B_\varepsilon(f(\xi))$ um $f(\xi)$ vom Radius $< \varepsilon$ abbilden



Beweis. Zunächst zeigen wir, daß das ε - δ -Kriterium die Stetigkeit impliziert. Ist x_n eine Folge in X mit $\lim_{n \rightarrow \infty} (x_n) = \xi$, dann existiert zu jedem $\tilde{\varepsilon} > 0$ ein $N = N(\tilde{\varepsilon})$ mit $d_X(x_n, \xi) < \tilde{\varepsilon}$ für $n \geq N$. Für $d_X(x_n, \xi) < \delta = \delta(\varepsilon)$ gilt nach Annahme

$$d_Y(f(x_n), f(\xi)) < \varepsilon.$$

Wählt man $\tilde{\varepsilon}$ gleich δ , folgt $d_Y(f(x_n), f(\xi)) < \varepsilon$ für $n \geq N(\delta)$. Das heisst $f(x_n) \rightarrow f(\xi)$.

Zum Beweis der Gegenrichtung nehmen wir an f sei stetig im Punkt ξ und das ε - δ -Kriterium wäre nicht erfüllt im Punkt ξ . Dann würde gelten

$$\exists \varepsilon_0 > 0 \quad \forall \delta > 0 \quad \exists x \in X \quad (d_X(x, \xi) < \delta \text{ und } d_Y(f(x), f(\xi)) \geq \varepsilon_0).$$

Wähle nun $\delta = 1/n$. Dann existiert ein x_n mit $d_X(x_n, \xi) < 1/n$ und $d_Y(f(x_n), f(\xi)) \geq \varepsilon_0$. Aus $d_X(x_n, \xi) < 1/n$ folgt $x_n \rightarrow \xi$. Die Stetigkeit von f im Punkt ξ zeigt $f(x_n) \rightarrow f(\xi)$ im Widerspruch zu $d_Y(f(x_n), f(\xi)) \geq \varepsilon_0$. $\square$

Folgerung. Eine Funktion $f: (X, d_X) \rightarrow (Y, d_Y)$ ist stetig genau dann, wenn gilt

$$\boxed{\forall \xi \in X \quad \forall \varepsilon > 0 \quad \exists \delta > 0 \quad \forall x \in X \quad (d_X(x, \xi) < \delta \Rightarrow d_Y(f(x), f(\xi)) < \varepsilon)}.$$

¹Intuitiv besagt dies: Kleine Änderungen verursachen unter f nur kleine Wirkungen. Ändert man ξ nur wenig ab, dann variiert auch $f(\xi)$ nur wenig.

2.5 Gleichmässige Stetigkeit

Definition 2.13. Eine Abbildung $f: (X, d_X) \rightarrow (Y, d_Y)$ heißt **gleichmäßig stetig** auf (X, d_X) , wenn zu jedem $\varepsilon > 0$ ein $\delta = \delta(\varepsilon) > 0$ existiert, so daß für alle $\xi, x \in X$ gilt:

$$d_X(x, \xi) < \delta \quad \Rightarrow \quad d_Y(f(x), f(\xi)) < \varepsilon.$$

oder

$$\boxed{\forall \varepsilon > 0 \exists \delta > 0 \forall \xi \in X \forall x \in X (d_X(x, \xi) < \delta \Rightarrow d_Y(f(x), f(\xi)) < \varepsilon)}.$$

Satz 2.14 (Satz von Heine). Ist (X, d_X) folgenkompakt, dann gilt: Jede stetige Funktion auf (X, d_X) ist gleichmäßig stetig.

Beweis. Wäre die Aussage falsch, würde gelten²

$$\exists \varepsilon_0 > 0 \forall \delta > 0 \exists \xi \in X \exists x \in X (d_X(x, \xi) < \delta \text{ und } d_Y(f(x), f(\xi)) \geq \varepsilon_0).$$

Fixiere ein solches $\varepsilon_0 > 0$. Für alle natürlichen Zahlen $n \geq 1$ und $\delta := n^{-1}$ existieren daher $\xi_n, x_n \in X$ mit $d_X(\xi_n, x_n) < n^{-1}$ und $d_Y(f(\xi_n), f(x_n)) \geq \varepsilon_0$. Bei sorgfältiger Auswahl von Teilfolgen kann man dann durch Übergang zu Teilfolgen zusätzlich die Konvergenz annehmen

$$\xi_n \xrightarrow{n \rightarrow \infty} \xi, \quad x_n \xrightarrow{n \rightarrow \infty} x.$$

(Man geht dazu zu einer konvergenten Teilfolge $\tilde{\xi}_n$ über, und streicht die entsprechenden Folgenglieder auch in der Folge x_n . In dieser Teilfolge von x_n geht man nochmals durch durch Streichen von Folgengliedern zu einer konvergenten Teilfolge $\tilde{x}_n$ über. Die entsprechenden Glieder streicht man auch aus der Teilfolge $\tilde{\xi}_n$. Wir nennen diese Teilfolgen wieder ξ_n bzw. x_n , und $d_X(\xi_n, x_n) < n^{-1}$ gilt dann auch für diese neuen Folgen.) Wegen $d_X(\xi_n, x_n) < n^{-1}$ und

$$0 \leq d_X(x, \xi) \leq \underbrace{d_X(x, x_n)}_{< \frac{1}{3}\varepsilon} + \underbrace{d_X(x_n, \xi_n)}_{< n^{-1} < \frac{1}{3}\varepsilon} + \underbrace{d_X(\xi_n, \xi)}_{< \frac{1}{3}\varepsilon} < \frac{1}{3}\varepsilon + \frac{1}{3}\varepsilon + \frac{1}{3}\varepsilon = \varepsilon$$

für $n \geq N(\varepsilon)$ und gegebenes $\varepsilon > 0$ folgt $d_X(x, \xi) = 0$ im Limes $n \rightarrow \infty$. Also $x = \xi$. Damit folgt wegen $f(x) = f(\xi)$ aus der Dreiecksungleichung

$$0 < \varepsilon_0 \leq d_Y(f(\xi_n), f(x_n)) \leq \underbrace{d_Y(f(\xi_n), f(\xi))}_{< \frac{1}{2}\varepsilon_0} + \underbrace{d_Y(f(x_n), f(x))}_{< \frac{1}{2}\varepsilon_0} < \varepsilon_0.$$

Ein Widerspruch! [Wegen der Stetigkeit von f ist die rechte Seite $< \varepsilon_0$, falls $n \geq N_1(\frac{1}{2}\varepsilon_0)$ resp. $n \geq N_2(\frac{1}{2}\varepsilon_0)$ wegen der Konvergenz der Folgen $f(\xi_n) \rightarrow f(\xi)$ und $f(x_n) \rightarrow f(x)$.] $\square$

²Negiert man eine Aussage, vertauschen sich beim ‘Vorbeiziehen’ der Negation die Quantoren $\forall$ und $\exists$. Die Negierung der Implikation $(A \Rightarrow B)$ liefert die Aussage $(A \text{ gilt und ebenso die Negierung von } B)$.

2.6 Reellwertige stetige Funktionen

Definition 2.15. Für einen metrischen Raum (X, d_X) bezeichne

$$C(X) := \{f: X \rightarrow \mathbb{R} \mid f \text{ ist stetig auf } (X, d_X)\}$$

den Raum der stetigen reellwertigen Funktionen auf X .

Lemma 2.16. Seien $f, g \in C(X)$ und $\lambda \in \mathbb{R}$. Dann sind auch $f + g, \lambda \cdot g, f \cdot g$ stetig auf X , d.h. $C(X)$ bildet eine $\mathbb{R}$ -Algebra.

Beweis. Wir beweisen das Lemma lediglich für das Produkt $f \cdot g$, da die anderen Rechnungen analog durchführbar sind. Sei $\xi \in X$ und $\varepsilon > 0$ gegeben. Dann gilt

$$\begin{aligned} |f(x)g(x) - f(\xi)g(\xi)| &\leq |f(x)g(x) - f(x)g(\xi) + f(x)g(\xi) - f(\xi)g(\xi)| \\ &\leq \underbrace{|f(x)|}_{\leq c_1} \cdot \underbrace{|g(x) - g(\xi)|}_{< \varepsilon/2c_1} + \underbrace{|g(\xi)|}_{\leq c_2} \cdot \underbrace{|f(\xi) - f(x)|}_{< \varepsilon/2c_2} < \varepsilon \end{aligned}$$

für Konstanten $c_1, c_2 > 0$ [denn $|g(x) - g(\xi)| < \varepsilon/2c_1$ gilt für $d_X(x, \xi) < \delta_1$; $|f(\xi) - f(x)| < \varepsilon/2c_2$ gilt für $d_X(x, \xi) < \delta_2$; $|f(x) - f(\xi)| < 1$ gilt für $d_X(x, \xi) < \delta_3$. Für $d_X(x, \xi) < \delta := \min(\delta_1, \delta_2, \delta_3)$ ist damit also $|f(x)| \leq 1 + |f(\xi)| =: c_1$.] Alles zusammen zeigt, daß $(f \cdot g)(x) = f(x)g(x)$ stetig im Punkt ξ ist. $\square$

Korollar 2.17. Polynome sind stetige Funktionen auf $\mathbb{R}$.

Lemma 2.18. Sei $f: (X, d_X) \rightarrow \mathbb{R}$ stetig und $f(x) \neq 0$ für alle $x \in X$. Dann gilt:

$$\frac{1}{f(x)}: (X, d_X) \rightarrow \mathbb{R}$$

ist definiert und stetig auf (X, d_X) .

Beweis. Für gegebenes $\xi \in X$ und $\varepsilon > 0$ gilt

$$\left| \frac{1}{f(x)} - \frac{1}{f(\xi)} \right| = \left| \frac{f(\xi) - f(x)}{f(x)f(\xi)} \right|.$$

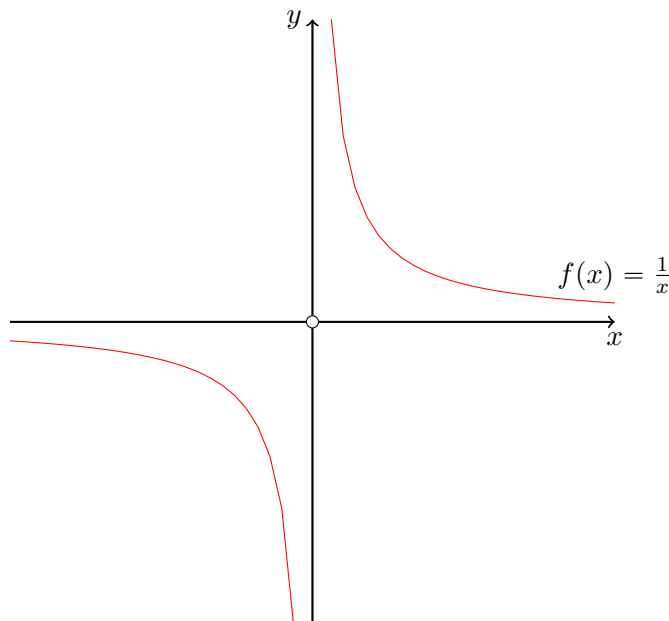
Aus $|f(x) - f(\xi)| < \frac{1}{3}|f(\xi)| =: \varepsilon_1 > 0$ für $d(X, \xi) < \delta_1$ (wegen der Stetigkeit von f) folgt $|\frac{2}{3}f(\xi)| < |f(x)|$ vermöge der unteren Dreiecksungleichung, d.h. $\left| \frac{1}{f(x)} \right| < \left| \frac{3}{2f(\xi)} \right|$. Also

$$\left| \frac{1}{f(x)} - \frac{1}{f(\xi)} \right| < \varepsilon,$$

falls $|f(x) - f(\xi)| < \frac{2}{3}|f(\xi)|^2 \cdot \varepsilon$, was gilt wenn $d_X(x, \xi)$ klein genug ist (Stetigkeit von f). $\square$

Beispiel 2.19. $f(x) = \frac{1}{x}$ erfüllt diese Bedingung für $X = \mathbb{R}^* = \mathbb{R} \setminus \{0\}$ und definiert damit eine stetige (allerdings nicht gleichmäßig stetige) Funktion auf $\mathbb{R}^*$

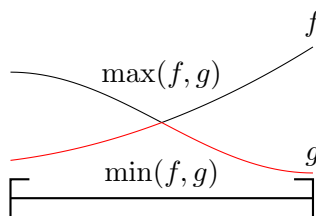
2 Stetige Abbildungen



Wir wollen uns nun mit dem Minimum und Maximum von Funktionen auseinandersetzen. Für Funktionen $f, g: (X, d_X) \rightarrow \mathbb{R}$ auf einem metrischen Raum (X, d_X) definieren wir

$$\max(f, g)(x) := \max(f(x), g(x))$$

und analog das Minimum $\min(f, g)(x)$.



Die stetigen Funktionen aus $C(X)$ definieren einen **Funktionsverband** auf X , d.h. $C(X)$ ist ein reeller Vektorraum von Funktionen auf X und es gilt

Satz 2.20. $\min(f, g)$ und $\max(f, g)$ sind in $C(X)$ für $f, g \in C(X)$.

Beweis. Es genügt, daß mit f auch $|f|$ (als Komposition der stetigen Abbildungen f und $|\cdot|$) stetig ist, denn $\max(f, g) = \frac{1}{2}(f + g) + \frac{1}{2}|f - g|$ und $\min(f, g) = -\max(-f, -g)$. $\square$

2.7 Gleichmässige Konvergenz

Definition 2.21. Eine Folge $(f_n)_{n \in \mathbb{N}}$ reellwertiger Funktionen $f_n: (X, d) \rightarrow \mathbb{R}$ **konvergiert punktweise** gegen eine Grenzfunktion $f: (X, d) \rightarrow \mathbb{R}$, wenn für jedes $x \in X$ gilt

$$\lim_{n \rightarrow \infty} f_n(x) = f(x).$$

Definition 2.22. Eine Folge $(f_n)_{n \in \mathbb{N}}$ von Funktionen $f_n: (X, d_X) \rightarrow \mathbb{R}$ **konvergiert gleichmässig** gegen eine Grenzfunktion $f: (X, d_X) \rightarrow \mathbb{R}$, wenn für jedes $\varepsilon > 0$ ein $N = N(\varepsilon) \in \mathbb{N}$ existiert, so daß für alle $n \geq N$ und alle $x \in X$ gilt³

$$|f_n(x) - f(x)| < \varepsilon.$$

Gleichmässige Konvergenz impliziert punktweise Konvergenz. Die Umkehrung gilt nicht. Wir geben später Beispiele im Zusammenhang mit dem Satz von Dini.

Die Supremumsnorm. Für beschränkte reellwertige Funktionen f auf X definiert man

$$\|f\| := \sup_{\xi \in X} |f(\xi)|.$$

Offensichtlich gilt $\|f\| = 0$ genau dann, wenn $f = 0$ gilt. Ebenso trivial ist die Eigenschaft $\|\lambda \cdot f\| = |\lambda| \cdot \|f\|$ für reelle Konstanten λ . Für zwei beschränkte reellwertige Funktionen f, g auf X ist aber auch $f + g$ beschränkt auf X und es gilt

$$\|f + g\| \leq \|f\| + \|g\|,$$

denn $|f(\xi) + g(\xi)| \leq |f(\xi)| + |g(\xi)| \leq \|f\| + \|g\|$ gilt für alle $\xi \in X$. Somit definiert $\|\cdot\|$ eine **Norm**, und $d(f, g) = \|f - g\|$ dann eine Metrik auf dem $\mathbb{R}$ -Vektorraum aller beschränkten Funktionen auf X .

Ist (X, d_X) ein folgenkompakter metrischer Raum, dann ist jede stetige Funktion beschränkt auf X . Für stetige Funktionen $f \in C(X)$ gilt sogar

$$\|f\| = \max_{x \in X} |f(x)|$$

und $d(f, g) = \|f - g\|$ sei die eingeschränkte Metrik auf $C(X)$. Offensichtlich gilt dann

Korollar 2.23. $(C(X), d)$ ist mit $d(f, g) := \|f - g\|$ ein metrischer Raum. Eine Folge von Funktionen $f_n \in C(X)$ konvergiert gegen $f \in C(X)$ in (X, d) genau dann, wenn $f_n(x)$ gleichmässig auf (X, d_X) gegen die Grenzfunktion $f(x)$ konvergiert.

³Das bedeutet $\|f - f_n\| < \varepsilon$ für die nachfolgend definierte Supremumsnorm.

2.8 Vollständigkeit von $C(X)$

Satz 2.24. Sei (X, d_X) folgenkompakt. Dann ist $(C(X), d)$ versehen mit der Supremums-Metrik $d(f, g) = \|f - g\|$ ein vollständiger metrischer Raum⁴.

Beweis. Sei f_n eine Cauchyfolge in $C(X)$. Dann gilt $\|f_n - f_m\| < \varepsilon$ für $n, m \geq N(\varepsilon)$. Wir konstruieren eine Grenzfunktion f . Für jeden fixierten Punkt $x \in X$ betrachten wir dazu die reelle Folge $y_n = f_n(x)$. Wegen $|y_n - y_m| \leq \|f_n - f_m\| < \varepsilon$ für $n, m \geq N(\varepsilon)$ ist y_n eine reelle Cauchyfolge. Ihr Limes $y_n \rightarrow y$ existiert und wir setzen $f(x) := y \in \mathbb{R}$.

1) Wir behaupten:

$$\|f - f_n\| < \varepsilon \quad \text{für} \quad n \geq N(\tfrac{1}{2}\varepsilon).$$

Ist nämlich $x \in X$ ein beliebiger Punkt, dann gibt es ein von x abhängiges $m_0 = m_0(\frac{1}{2}\varepsilon, x)$ mit $|f(x) - f_m(x)| < \frac{1}{2}\varepsilon$ für $m \geq m_0(\frac{1}{2}\varepsilon, x)$ wegen $f_n(x) \rightarrow f(x)$. Dies liefert für $n, m \geq N(\frac{1}{2}\varepsilon)$ und $m \geq m_0(\varepsilon, x)$

$$|f(x) - f_n(x)| \leq |f(x) - f_m(x)| + |f_m(x) - f_n(x)| < \tfrac{1}{2}\varepsilon + \underbrace{\|f_m - f_n\|}_{< \frac{1}{2}\varepsilon} \leq \varepsilon.$$

Das gewählte m ist aber beliebig, und kann daher (angepasst an das jeweilige x) beliebig groß gewählt werden! Es folgt damit wie behauptet $|f(x) - f_n(x)| < \varepsilon$ für $n \geq N(\frac{1}{2}\varepsilon)$.

2) Wir behaupten: $f \in C(X)$. Betrachte hierzu $x_1, x_2 \in X$. Dann gilt für geeignetes n und δ sowie $d_X(x_1, x_2) < \delta$

$$|f(x_1) - f(x_2)| \leq \underbrace{|f(x_1) - f_n(x_1)|}_{< \frac{1}{3}\varepsilon} + \underbrace{|f_n(x_1) - f_n(x_2)|}_{< \frac{1}{3}\varepsilon} + \underbrace{|f_n(x_2) - f(x_2)|}_{< \frac{1}{3}\varepsilon} < \varepsilon.$$

[Ist $d_X(x_1, x_2) < \delta = \delta(\frac{1}{3}\varepsilon, f_n)$, dann gilt $|f_n(x_1) - f_n(x_2)| < \frac{1}{3}\varepsilon$, da f_n stetig und damit gleichmäßig stetig auf (X, d_X) ist. Ausserdem haben wir für beliebiges x (insbesondere also für $x = x_1, x_2$) gezeigt $|f(x) - f_n(x)| < \varepsilon/3$ für $n \geq N(\frac{1}{2}\varepsilon/3)$. Wählt man daher ein $n \geq N(\varepsilon/6)$, dann folgt für alle x_1, x_2 mit $d_X(x_1, x_2) < \delta$, wobei δ jetzt nur noch von ε abhängt, die Ungleichung $|f(x_1) - f(x_2)| < \varepsilon$.] Das heißt f ist stetig auf X . Wegen $\|f - f_n\| < \varepsilon$ für $n \geq N(\frac{1}{2}\varepsilon)$ konvergiert die Funktionenfolge $f_n(x)$ gleichmäßig gegen $f(x)$. Dies zeigt die Cauchy-Vollständigkeit von $C(X)$. $\square$

2.9 Monotone Folgen stetiger Funktionen

Sei X eine Menge und $A \subseteq X$ eine Teilmenge. Die **charakteristische** Funktion $\chi_A : X \rightarrow \mathbb{R}$ ist definiert durch: $\chi_A(x) = 1$ für $x \in A$ und $\chi_A(x) = 0$ für $x \notin A$. Man sagt, eine Funktion $f : X \rightarrow \mathbb{R}$ hat **Träger** in A , falls für alle $x \notin A$ gilt $f(x) = 0$. Für jede Funktion $f : X \rightarrow \mathbb{R}$ hat die Funktion $\chi_A \cdot f$ Träger in A . Bemerkung: Der Raum $C_c(X) \subseteq C(X)$ aller Funktionen

⁴Dies überträgt sich verbatim auf den Raum $C(X, \mathbb{R}^N)$ aller Funktionen auf X mit Werten in $\mathbb{R}^N$.

$f \in C(X)$ mit kompakten Träger in (X, d) ist ein Verband [benutze dazu: Die Vereinigung $A = A_1 \cup A_2$ zweier folgenkompakter Mengen A_1, A_2 ist folgenkompakt (Schubfachschluss)].

Für eine abgeschlossene Teilmenge $A \subseteq X$ des metrischen Raumes (X, d) ist die in Beispiel 2.2(3) erklärte Abstandsfunktion $d(x, A) = \inf\{d(x, a) \mid a \in A\}$ eine stetige Funktion auf (X, d) mit $d(x, A) = 0 \iff x \in A$.

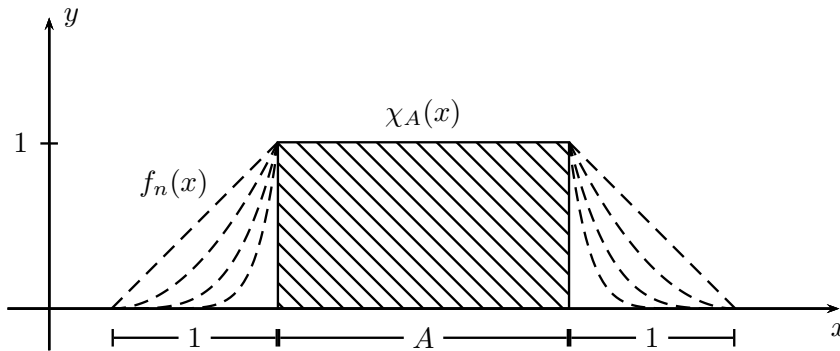
Beispiel 2.25. Sei $f(x) \geq 0$ eine nichtnegative stetige Funktionen auf X . Nach Satz 2.20, Lemma 2.16 und Beispiel 2.2 sind dann auch die Funktionen

$$f_n(x) = \max(0, 1 - d(x, A))^n \cdot f(x)$$

stetig auf X . Wegen $f(x) \geq 0$ definieren die $f_n(x) \geq 0$ eine monoton fallende Folge stetiger Funktionen auf X , welche daher gegen eine Grenzfunktion $g(x) \geq 0$ auf X konvergiert. Wegen $0 \leq \max(0, 1 - d(x, A)) < 1$ für $x \notin A$ und $0 \leq \max(0, 1 - d(x, A)) = 1$ für $x \in A$ konvergiert diese Folge (wegen Lemma 1.15) punktweise gegen die Grenzfunktion $g = \chi_A \cdot f$

$$C(X) \ni f_n(x) \searrow \chi_A(x) \cdot f(x) \quad , \quad A \text{ abgeschlossen in } X, f \geq 0 \text{ in } C(X) \quad .$$

Bemerkung. Ein gleichmäßiger Limes von stetigen Funktionen f_n auf einem kompakten Raum (X, d) definiert eine stetige Grenzfunktion (Satz 2.24). Das obige Beispiel 2.25 zeigt, daß die *analoge Aussage für monotone Limiten stetiger Funktionen f_n falsch ist*. Sei dazu X etwa das Intervall $[0, 5]$, $f = 1$ sowie $A = [a, b] \subset X$ für $1 < a < b < 4$. Dann ist der Grenzwert $\chi_A = \lim_n f_n$ keine stetige Funktion auf X , denn der Limes $\lim_n \chi_A(x_n) = 0$ für eine linksseitige Folge $x_n \rightarrow a$ (d.h. mit $x_n < a$) und ist verschieden von $\chi_A(a) = 1$.



Der folgende Satz von Dini zeigt, daß eine monotone Folge stetiger Funktion f_n auf einem kompakten metrischen Raum X genau dann eine stetige Grenzfunktion f besitzt, wenn die Folge f_n gleichmäßig auf X gegen f konvergiert.

Satz 2.26 (Satz von Dini). Sei (X, d) folgenkompakt. Sei $f_n \searrow f$ eine monoton fallende punktweise konvergente Folge von stetigen reellwertigen Funktionen f_n auf X . Ist die Grenzfunktion f stetig auf X , konvergieren die f_n gleichmäßig gegen f auf X .

2 Stetige Abbildungen

Beweis. Ersetzt man f_n durch $f_n - f$, kann oBdA $f = 0$ angenommen werden. Wir nehmen an die Folge f_n sei monoton fallend. Wäre die Konvergenz nicht gleichmässig, gäbe es ein $\varepsilon_0 > 0$ so daß für alle $n \in \mathbb{N}$ ein $x_n \in X$ existiert mit

$$0 < \varepsilon_0 \leq f_n(x_n) \quad (\forall n \in \mathbb{N}).$$

a) Da X folgenkompakt ist, können wir durch Übergang zu einer Teilfolge oBdA annehmen $x_n \rightarrow \xi$. Wegen $f_n(\xi) \rightarrow f(\xi) = 0$ gibt es dann ein n_0 mit

$$0 \leq f_{n_0}(\xi) < \varepsilon_0/2.$$

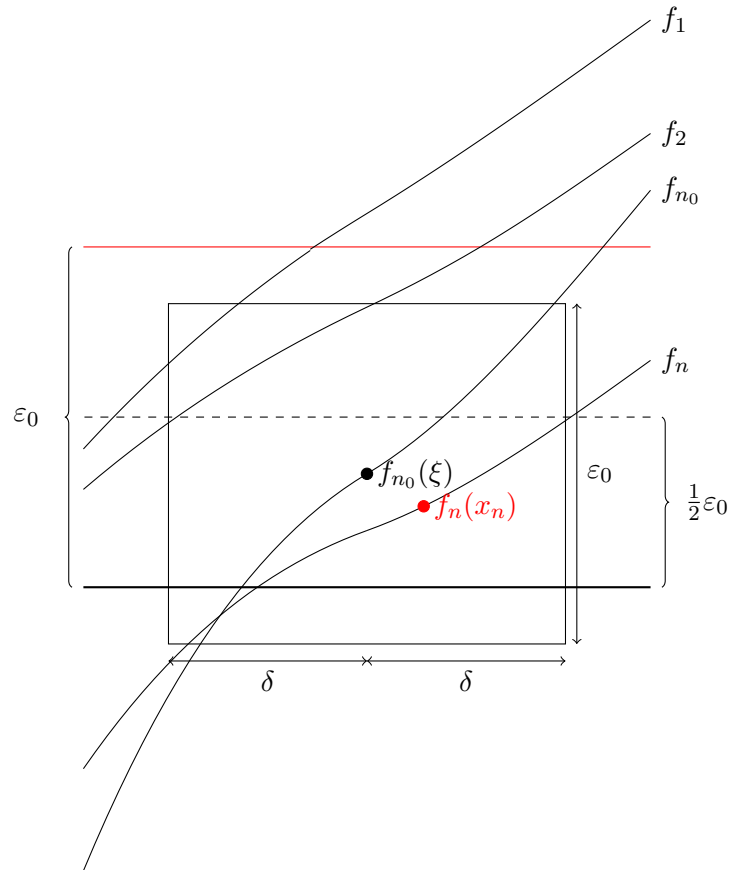
b) Wegen der Stetigkeit von $f_{n_0}(x)$ im Punkt ξ gilt $f_{n_0}(x_n) - f_{n_0}(\xi) < \varepsilon_0/2$ für alle $n \geq n_1$, bei geeigneter Wahl von $n_1 > n_0$ wegen $x_n \rightarrow \xi$. Also

$$0 \leq f_{n_0}(x_n) = (f_{n_0}(x_n) - f_{n_0}(\xi)) + f_{n_0}(\xi) < \varepsilon_0, \quad (\forall n \geq n_1).$$

c) Aus der Monotonie der Folge $f_m(x)$ ergibt sich dann

$$0 \leq f_m(x_n) \leq f_{n_0}(x_n) < \varepsilon_0, \quad (\forall m \geq n_1).$$

Für $m = n$ und beliebiges $n \geq n_1$ steht dies im Widerspruch zu $\varepsilon_0 \leq f_n(x_n)$. □



3 Integration

3.1 Vorbemerkungen

Es ist ein klassisches Problem für reellwertige Funktionen $f : X \rightarrow \mathbb{R}$ auf $X = \mathbb{R}$ die Fläche zwischen f und der x -Achse, das sogenannte Integral $\int_X f(x)dx$, zu bestimmen. Dies setzt voraus, daß dieser Flächenbegriff überhaupt sinnvoll für die Funktion f definiert werden kann; wenn er definierbar ist, ist zu klären ob diese Definition eindeutig ist. Dazu machen wir gewisse Einschränkungen an die Menge $B(X)$ der zugelassenen Funktionen f und diskutieren zuerst axiomatische Eigenschaften, die ein solches Integral besitzen sollte.

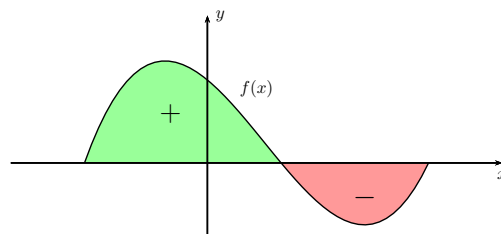
Wir nehmen dazu an, $B(X)$ sei ein $\mathbb{R}$ -Vektorraum reellwertiger Funktionen auf $X = \mathbb{R}$, d.h. es gelte $f + g \in B(X)$ und $\lambda \cdot f \in B(X)$ für alle $f, g \in B(X)$ und $\lambda \in \mathbb{R}$. Das gesuchte **Integral**

$$I(f) = \int_{\mathbb{R}} f(x)dx$$

fassen wir auf als eine Abbildung $I : B(X) \rightarrow \mathbb{R}$ auf, und erwarten folgende Eigenschaften

1. **$\mathbb{R}$ -Linearität:** $I(f + g) = I(f) + I(g)$ und $I(\lambda \cdot f) = \lambda \cdot I(f)$ für $\lambda \in \mathbb{R}$.
2. **Monotonie:** $f \leq g \implies I(f) \leq I(g)$
3. **Normierung:** Es gilt $\chi_{[0,1]} \in B(X)$ und $I(\chi_{[0,1]}) = 1$.
4. **Translationsinvarianz:** Für $f \in B(X)$, $x_0 \in \mathbb{R}$ ist das Translat $g(x) = f(x + x_0)$ in $B(X)$ und es gilt $I(f) = I(g)$.

Die erste Eigenschaft impliziert $I(0) = 0$. Man sollte berücksichtigen, daß im allgemeinen f keine positive Funktion sein wird. Hat f negative Werte, zerlegt man $f(x) = f_+(x) + f_-(x)$ in die **positive** Funktion $f_+ = \max(f, 0) \geq 0$ und die **negative** Funktion $f_- = \min(f, 0) \leq 0$ so daß gelten soll $I(f) = I(f_+) - I(-f_-)$.



3 Integration

Der Wert $I(f)$ ist im allgemeinen verschieden von $I(|f|) = I(f_+) + I(-f_-)$. In der Tat, aus $-f_-(x) \geq 0$ und der Monotonie folgt $I(-f_-) \geq 0$. Wegen der Linearität ist dann aber $I(f_-) = -I(-f_-)$ eine negative Zahl. Das heisst, die Fläche unter der x -Achse wird negativ sein. Bei dem gesuchten Begriff des Flächeninhaltes handelt es sich daher um den Begriff einer **orientierten Fläche**¹. Für Funktionen mit positiven Werten sollte das gesuchte Integral positive Werte besitzen. Dies erklärt Eigenschaft 2, denn $g - f \geq 0$ impliziert $I(g) - I(f) = I(g - f) \geq 0$ und damit die Monotonie $I(g) \geq I(f)$.

Daß die Eigenschaften 1 und 2 das Integral höchstens bis auf eine Normierungskonstante festlegen können, erklärt Axiom 3. Aber auch die ersten drei Eigenschaften können I nicht eindeutig festlegen, wie man sofort an einem Beispiel sieht. Für jeden beliebigen Punkt $\xi \in [0, 1]$ ist $f \mapsto I(f) := f(\xi)$ linear, monoton und erfüllt $I(\chi_{[0,1]}) = 1$, aber nicht Eigenschaft 4. Dies motiviert das letzte Axiom als natürliche und vor allem geometrisch plausible Annahme.

Hat f keinen kompakten Träger, kann man nicht erwarten, daß das Integral $I(f)$ definiert ist. Fixiert man ein kompaktes Intervall $A \subset \mathbb{R}$ mit $\chi_A \in B(X)$ und betrachtet nur Funktionen f im Unterraum $B_A(X) \subseteq B(X)$ der Funktionen mit Träger in A , liefert Eigenschaft 1 und 2 sowie die triviale Abschätzung $-\|f\| \cdot \chi_A \leq f \leq \|f\| \cdot \chi_A$ dann $-\|f\| \cdot I(\chi_A) \leq I(f) \leq \|f\| \cdot I(\chi_A)$, d.h. die **Boxungleichung**² für die **Supremumsnorm** $\|f\| = \sup_{x \in A} |f(x)|$ von f auf A

$$\boxed{|I(f)| \leq I(\chi_A) \cdot \|f\|} \quad \text{für alle } f \in B_A(X).$$

Es folgt für $X = \mathbb{R}$

Lemma 3.1. Sind alle Funktionen in $B_A(X)$ beschränkt, definiert die Einschränkung eines Integrals $I : B(X) \rightarrow \mathbb{R}$ mit den obigen Eigenschaften 1.-4. eine **stetige** $\mathbb{R}$ -lineare Abbildung

$$I : B_A(X) \longrightarrow \mathbb{R},$$

wenn $B_A(X)$ mit der Supremums-Metrik $d(f, g) = \|f - g\|$ versehen wird.

3.2 Treppenfunktionen

Definition 3.2. Ein **Verband** $B(X)$ auf X ist ein Raum von Funktionen $f : X \rightarrow \mathbb{R}$ mit folgender Eigenschaft: Für $f, g \in B(X)$ und $\lambda \in \mathbb{R}$ sind auch $f + g$, $\lambda \cdot f$ und $|f|$ in $B(X)$. D.h. $B(X)$ ist ein $\mathbb{R}$ -Vektorraum und es gilt $f, g \in B(X) \implies \min(f, g) \in B(X)$ und $\max(f, g) \in B(X)$ (siehe den Beweis von Satz 2.20).

Sei $B(X)$ ein Verband auf X . Dann definiert für jede Teilmenge $A \subseteq X$ der Unterraum $B_A(X) \subseteq B(X)$ aller Funktionen $f \in B(X)$, deren Träger in $A \subseteq X$ liegt, erneut einen Verband auf X .

¹Was wäre die natürliche Orientierung der Berandung?

²Ist f nicht beschränkt, gilt $\|f\| = +\infty$. Die Box-Ungleichung bleibt richtig in $\mathbb{R}^+$ unter der Annahme $I(\chi_A) > 0$.

Definition 3.3. Eine reellwertige Funktion f auf $X = \mathbb{R}$ nennt man **Treppenfunktion**, wenn es endliche viele reelle Zahlen $\xi_0 \leq \xi_1 \leq \dots \leq \xi_m$ gibt mit: 1) Die Funktion hat Träger im Intervall $[\xi_0, \xi_m]$ und 2) Die Einschränkung $f|_{(\xi_{i-1}, \xi_i)}$ von f auf die offenen Teilintervalle (ξ_{i-1}, ξ_i) ist eine konstante Funktion für alle $i = 1, \dots, m$.

Jede Treppenfunktion ist eine $\mathbb{R}$ -Linearkombination von charakteristischen Funktionen von abgeschlossenen Intervallen (*)

$$f(x) = \sum_{i=1}^m c_i \cdot \chi_{[\xi_{i-1}, \xi_i]}(x) + \sum_{j=0}^m d_j \cdot \chi_{[\xi_j, \xi_j]}(x)$$

für gewisse reelle Konstanten c_i, d_j . [Beachte $f(\xi_i) = c_i + c_{i+1} + d_i$]. Die Umkehrung gilt auch. Jede solche Linearkombination ist eine Treppenfunktion. Die Menge der Stützstellen $\xi_0, \dots, \xi_m$ einer gegebenen Treppenfunktion f ist dahingehend willkürlich, daß man $\{\xi_0, \dots, \xi_m\}$ durch eine beliebige grössere endliche Menge von Stützstellen ersetzen kann.

Lemma 3.4. Der Raum $T(\mathbb{R})$ aller Treppenfunktionen auf $X = \mathbb{R}$ ist ein $\mathbb{R}$ -Vektorraum³ und wird erzeugt von den charakteristischen Funktionen der kompakten Intervalle $A \subset \mathbb{R}$, und $T(\mathbb{R})$ definiert einen Verband auf $\mathbb{R}$.

Beweis. Offensichtlich ist jedes skalare Vielfache einer Treppenfunktionen wieder eine Treppenfunktion. Seien $f, g \in T(\mathbb{R})$ Treppenfunktion. Durch Vergrössern der Menge der Stützstellen kann man oBdA annehmen, daß sowohl f als auch g Träger in $[\xi_0, \xi_m]$ haben und konstant sind auf den Teilintervallen (ξ_{i-1}, ξ_i) für $i = 1, \dots, m$. Dann ist natürlich klar, daß auch $f + g$ eine Treppenfunktion ist mit Träger in $[\xi_0, \xi_m]$ und Stützstellen $\xi_0, \dots, \xi_m$. Dies zeigt, $T(\mathbb{R})$ ist ein $\mathbb{R}$ -Vektorraum. Wegen (*) bilden die charakteristischen Funktionen χ_A aller kompakten Intervalle $A \subset \mathbb{R}$ ein Erzeugendensystem des Vektorraums $T(\mathbb{R})$. Um zu sehen daß $T(\mathbb{R})$ ein Verband ist, beachte daß für Treppenfunktionen wie in (*) gilt

$$|f(x)| = \sum_{i=1}^m |c_i| \cdot \chi_{[\xi_{i-1}, \xi_i]}(x) + \sum_{j=0}^m b_j \cdot \chi_{[\xi_j, \xi_j]}(x)$$

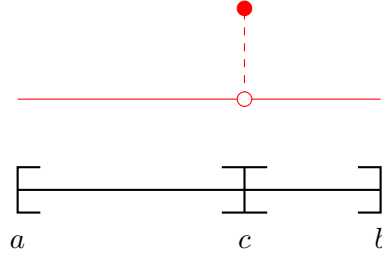
für die Konstanten $b_j := |c_j + c_{j+1} + d_j| - |c_j| - |c_{j+1}|$. □

Achtung. Die charakteristischen Funktionen χ_A bilden zwar ein Erzeugendensystem von $T(\mathbb{R})$, aber keine Basis. Beispiele für lineare Abhängigkeiten zwischen den erzeugenden charakteristischen Funktionen sind etwa

$$\chi_{[a,b]} = \chi_{[a,c]} + \chi_{[c,b]} - \chi_{[c,c]}.$$

für $a \leq c \leq b$. Beachte auch $\chi_{(c,d)} = \chi_{[c,d]} - \chi_{[c,c]} - \chi_{[d,d]}$ und $\chi_{(a,b)} = \chi_{[a,b]} - \chi_{[a,a]}$.

³Wie man leicht sieht ist $T(\mathbb{R})$ unendlich dimensional.



Die Eigenschaften 1.-4. des vorigen Abschnitts legen ein eindeutiges Integral fest, d.h.

Lemma 3.5. *Es gibt eine eindeutig bestimmte $\mathbb{R}$ -lineare monotone translationsinvariante und normierte Abbildung $I : T(\mathbb{R}) \rightarrow \mathbb{R}$ auf dem Verband der Treppenfunktion $T(\mathbb{R})$.*

Beweis. $T(\mathbb{R})$ wird als Vektorraum von den Funktionen χ_A , $A = [a, b]$ erzeugt. Es genügt daher zum Beweis der *Eindeutigkeit* zu zeigen

$$I(\chi_A) = |b - a|, \quad A = [a, b].$$

OBdA $a = 0$ (wegen der Translationsinvarianz von I). Zu zeigen ist dann $I(\chi_{[0,b]}) = b$. Für $b = 1$ ist dies klar. Beachte $0 \leq \sum_{i=1}^n \chi_{[\frac{i-1}{n}, \frac{i}{n}]} \leq \chi_{[0,1]}$. Aus der Translationsinvarianz $I(\chi_{[\frac{i-1}{n}, \frac{i}{n}]}) = I(\chi_{[0,1/n]})$ und der Monotonie folgt daher $0 \leq I(\chi_{[0,1/n]}) \leq 1/n$ für alle n , also $I(\chi_{[a,a]}) = I(\chi_{[0,0]}) = 0$ für alle a . Aus $\sum_{i=1}^n \chi_{[\frac{i-1}{n}, \frac{i}{n}]} = \chi_{[0,1]} + \sum_{i=1}^{n-1} \chi_{[\frac{i}{n}, \frac{i+1}{n}]}$ und der Translationsinvarianz folgt damit $I(\chi_{[0, \frac{1}{n}]}) = 1/n$. Analog zeigt man $I(\chi_{[0,b]}) = b$ für alle rationalen Zahlen $b = \frac{m}{n} \geq 0$. Jedes reelle $b > 0$ kann durch rationale Zahlen ν, μ mit $\nu < b < \mu$ beliebig gut angenähert werden (Dezimalbruchentwicklung von b). Die Monotonie von I liefert $\nu \leq I(\chi_{[0,b]}) \leq \mu$, und im Limes $\nu \rightarrow b, \mu \rightarrow b$ folgt $I(\chi_{[0,b]}) = b$.

Die *Existenz* von I zeigt man durch den Ansatz

$$I(f) = \sum_{i=1}^m c_i \cdot |\xi_i - \xi_{i-1}|$$

für $f(x) = \sum_{i=1}^m c_i \cdot \chi_{[\xi_{i-1}, \xi_i]}(x) + \sum_{j=0}^m d_j \cdot \chi_{[\xi_j, \xi_j]}(x)$ in $T(\mathbb{R})$. Offensichtlich ändert sich dieser Wert nicht bei Hinzunahme von weiteren Stützstellen ξ_i . Damit ist $I(f)$ wohldefiniert und hängt nur von der Funktion $f \in T(\mathbb{R})$ ab. Durch Hinzunahme von Stützstellen zeigt man leicht die behaupteten Eigenschaften 1.-4. $\square$

3.3 Das reelle Standardintegral

Definition 3.6. *Eine Funktion $f: \mathbb{R} \rightarrow \mathbb{R}$ heisst **stückweise stetig**, wenn es endlich viele reelle Stützpunkte $\xi_0 \leq \xi_1 \leq \xi_2 \leq \dots \leq \xi_m$ gibt mit: 1) f hat Träger im Intervall $[\xi_0, \xi_m]$ und 2) die Einschränkungen $f|_{(\xi_{i-1}, \xi_i)}$ von f auf die offenen Teilintervalle (ξ_{i-1}, ξ_i) lassen sich zu stetigen Funktionen $f_i: A_i \rightarrow \mathbb{R}$ der abgeschlossenen Intervalle $A_i = [\xi_{i-1}, \xi_i]$ fortsetzen.*

Können die Funktionen f_i konstant gewählt werden, ist f eine Treppenfunktion. Also ist der Raum der Treppenfunktion $T(\mathbb{R})$ ein Unterraum des Raumes $CT(\mathbb{R})$ aller stückweise stetigen Funktionen auf $\mathbb{R}$. Wie im Beweis von Lemma 3.4 zeigt man, daß $CT(\mathbb{R})$ einen Verband auf $\mathbb{R}$ definiert. Dieser Verband enthält ausser $T(\mathbb{R})$ auch den Verband $C_c(\mathbb{R})$ aller stetigen Funktionen auf $\mathbb{R}$ mit kompaktem Träger. Wir verallgemeinern nun Lemma 3.5 auf stückweise stetige Funktionen.

Satz 3.7. *Es gibt eine eindeutig bestimmte $\mathbb{R}$ -lineare monotone translationsinvariante und normierte Abbildung*

$$I : CT(\mathbb{R}) \longrightarrow \mathbb{R}$$

auf dem Verband $CT(\mathbb{R})$ der stückweise stetigen Funktionen auf $\mathbb{R}$.

Dieses I nennen wir das **Euklidische Standardintegral** und schreiben dafür $I(f) = \int_{\mathbb{R}} f(x) dx$. Da für jede Treppenfunktion $f \in CT(\mathbb{R})$ auch die Funktion $\chi_{[a,b]} \cdot f$ in $CT(\mathbb{R})$ ist, benutzen wir folgende Schreibweise

$$\int_{[a,b]} f(x) dx := \int_{\mathbb{R}} \chi_{[a,b]}(x) \cdot f(x) dx .$$

Ganz wesentlich für den Beweis von Satz 3.7 ist das folgende

Lemma 3.8. *Für jedes $f \in CT(\mathbb{R})$ und jede reelle Zahl $\varepsilon > 0$ existiert ein $g \in T(\mathbb{R})$ mit $\|f - g\| < \varepsilon$ (Supremumsnorm auf $CT(\mathbb{R})$). Hat f Träger im Intervall A , dann oBdA auch g .*

Beweis. Man reduziert dies sofort auf den Fall $f \in C(A)$ und $g \in T_A(\mathbb{R})$, für ein Intervall $A = [a, b]$. Es existiert dann ein $\delta > 0$ mit $|f(x) - f(y)| < \varepsilon$ für alle $x, y \in A$, $|x - y| < \delta$ (nach Satz 2.14). Zerlege A äquidistant in Teilintervalle $[\xi_{i-1}, \xi_i]$ der Länge $|\xi_i - \xi_{i-1}| < \delta$. Setze $g(x) = f(\xi_i)$ für $x \in (\xi_{i-1}, \xi_i]$ und $g(\xi_0) = f(\xi_0)$. Dann ist $g \in T_A(\mathbb{R})$ und $\|f - g\| < \varepsilon$. $\square$

Beweis von Satz 3.7. Für die Existenz von $I(f)$ benutzt man die Methode des Beweises von Satz 2.24: Für gegebenes $f \in CT_A(\mathbb{R})$ wähle eine Folge $g_n \in T_A(\mathbb{R})$ mit $\|f - g_n\| \rightarrow 0$ vermöge Lemma 3.8. Bezüglich der Supremumsnorm ist dann g_n eine Cauchyfolge in $T_A(\mathbb{R})$. Nach Lemma 3.1 ist damit $I(g_n)$ (definiert in Lemma 3.5) eine reelle Cauchyfolge. Diese konvergiert gegen einen Grenzwert in $\mathbb{R}$ und wir setzen

$$I(f) := \lim_n I(g_n) .$$

Durch Mischen von Folgen in $T(\mathbb{R})$ zeigt man nun, daß $I(f)$ nur von f und nicht von der Wahl der Hilfsfolge $g_n \in T(\mathbb{R})$ abhängt. Man zeigt dann leicht die Linearität von I und die Translationsinvarianz mit Hilfe von Lemma 3.5. Ist $f \geq 0$, kann man oBdA die $g_n \in T(\mathbb{R})$ durch die Treppenfunktionen $\max(g_n, 0)$ ersetzen und damit $g_n \geq 0$ und $I(g_n) \geq 0$ annehmen. Dies beweist $I(f) \geq 0$ für $f \geq 0$. Die Normierungseigenschaft ist trivial.

Eindeutigkeit. Seien I_1, I_2 verschiedene Lösungen. Dann ist $J = I_1 - I_2$ eine stetige $\mathbb{R}$ -lineare Abbildung auf $CT(\mathbb{R})$ und es gilt $J(g) = 0$ für alle $g \in T(\mathbb{R})$ wegen Lemma 3.5. Nach Definition gilt wegen $g_n \in T(\mathbb{R})$ daher $J(f) = \lim_n J(g_n) = 0$ für alle $f \in CT(\mathbb{R})$. $\square$

3.4 Eigenschaften des Standardintegrals

Man hat folgende elementare **Substitutionsregel** für das Standardintegral

Lemma 3.9. Für reelle a, b und $\lambda > 0$ und Funktionen $f \in CT_{[a,b]}(\mathbb{R}) \subset CT(\mathbb{R})$ gilt

$$\frac{1}{\lambda} \int_{[\lambda a, \lambda b]} f\left(\frac{x}{\lambda}\right) dx = \int_{[a,b]} f(x) dx,$$

Beweis. Für $f \in CT(\mathbb{R})$ ist $h_f(x) := \frac{1}{\lambda} f\left(\frac{x}{\lambda}\right)$ in $CT(\mathbb{R})$, und daher $J(f) := \int_{\mathbb{R}} h_f(x) dx$ eine $\mathbb{R}$ -lineare monotone und translationsinvariante normierte Abbildung $CT(\mathbb{R}) \rightarrow \mathbb{R}$. [Wegen $\int_{\mathbb{R}} \chi_{[0,1]}(\frac{x}{\lambda}) dx = \int_{\mathbb{R}} \chi_{[0,\lambda]}(x) dx = \lambda$ ist J normiert]. Aus Satz 3.7 folgt $J(f) = \int_{\mathbb{R}} f(x) dx$. Es folgt $\frac{1}{\lambda} \int_{\mathbb{R}} f\left(\frac{x}{\lambda}\right) dx = \int_{\mathbb{R}} f(x) dx$ für alle $f \in CT(\mathbb{R})$. $\square$

Lemma 3.10. Für alle $f \in CT(\mathbb{R})$ und alle reellen Zahlen $a \leq c \leq b$ gilt

$$\int_{[a,b]} f(x) dx = \int_{[a,c]} f(x) dx + \int_{[c,b]} f(x) dx.$$

Beweis. Benutze $\chi_{[a,b]} = \chi_{[a,c]} + \chi_{[c,b]} - \chi_{[c,c]}$ und $\int_{[c,c]} f(x) dx = f(c)I(\chi_{[c,c]}) = 0$. $\square$

3.5 Der Logarithmus

Für $0 < a \leq b$ ist $f(t) = \frac{1}{t}$ stetig auf $[a, b]$. Wegen $C([a, b]) \hookrightarrow CT(\mathbb{R})$ ist daher nach Satz 3.7 das Integral $\int_{[a,b]} \frac{dt}{t}$ erklärt. Dies definiert für $x \geq 1$ den natürlichen **Logarithmus**

$$\log(x) := \int_{[1,x]} \frac{dt}{t}.$$

Aus Lemma 3.9 für $\lambda = y$ folgt $\int_{[y,xy]} \frac{dt}{t} = \int_{[1,x]} \frac{dt}{t}$ für $x, y \geq 1$. Wegen Lemma 3.10 gilt

$$\int_{[1,y]} \frac{dt}{t} + \int_{[y,xy]} \frac{dt}{t} = \int_{[1,xy]} \frac{dt}{t},$$

und zusammen genommen erhält man die **Funktionalgleichung**

$$\log(x) + \log(y) = \log(xy).$$

Beachte $\log(1) = 0$. Setzt man $\log(x) := -\log(1/x)$ für $x \in (0, 1)$, folgt diese Formel dann rein formal für alle $x, y > 0$. Es gilt **Monotonie**: $\log(b) - \log(a) = \log(b/a) > 0$ für alle $b > a > 0$ [wegen $b/a > 1$ und $\log(x) > 0$ für $x > 1$; beachte $\log(x) \geq l([1, x]) \cdot \frac{1}{x} = \frac{x-1}{x}$ wegen $\min_{t \in [1,x]} \frac{1}{t} = \frac{1}{x}$ und der Monotonie des Integrals $\int_{[1,x]} \frac{dt}{t} \geq \int_{[1,x]} \frac{dt}{x}$.] Dies zeigt

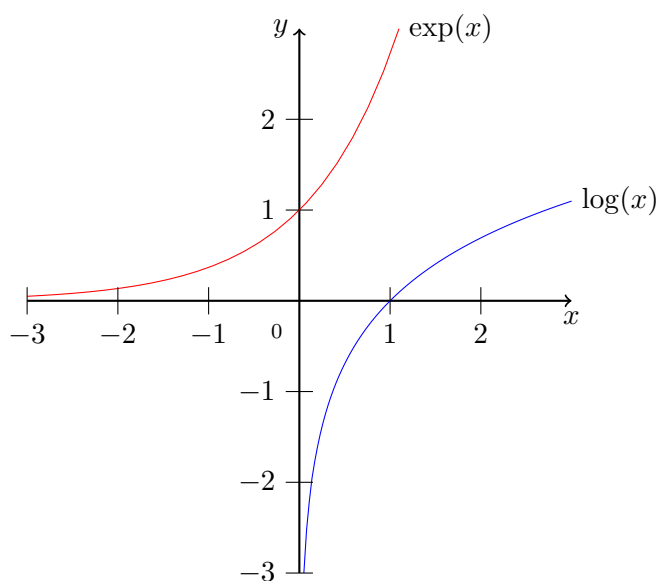
Satz 3.11. *Der Logarithmus $\log: \mathbb{R}_{>0} \rightarrow \mathbb{R}$ ist eine streng monoton wachsende Funktion, und definiert einen Gruppenhomomorphismus von der multiplikativen Gruppe in die additive Gruppe*

$$\log: (\mathbb{R}_{>0}, \cdot) \rightarrow (\mathbb{R}, +).$$

Es ist nun leicht zu sehen, daß der Logarithmus sogar einen Isomorphismus definiert

$$(\mathbb{R}_{>0}, \cdot) \cong (\mathbb{R}, +).$$

Die Injektivität folgt sofort aus der strengen Monotonie. Für die Surjektivität genügt, daß jede Zahl ≥ 0 im Bild liegt. Wegen $\log(2^n) = n \cdot \log(2)$ und $\log(2) > 0$ nimmt $\log(x)$ beliebig große Werte an. Andererseits gilt $\log(x) - \log(y) \leq \frac{1}{a}|x - y|$ für $x, y \in [a, b]$ (Boxungleichung) wegen $\max_{t \in [a, b]} (\frac{1}{t}) = \frac{1}{a}$. Also ist $\log(x)$ eine Lipschitz-stetige Funktion auf $[1, b]$. Aus dem Zwischenwertsatz folgt dann, daß jede reelle Zahl ≥ 0 im Bild des Logarithmus liegt.



Die dadurch eindeutig bestimmte monotone Umkehrfunktion des Logarithmus ist die sogenannte **Exponentialfunktion**, die einen bijektiven Gruppenhomomorphismus

$$\exp: (\mathbb{R}, +) \rightarrow (\mathbb{R}_{>0}, \cdot)$$

definiert mit $\exp(0) = 1$. Für alle reellen Zahlen x, y gilt daher die **Funktionalgleichung**

$$\boxed{\exp(x + y) = \exp(x) \cdot \exp(y)}.$$

Notation. Für beliebige reelle Zahlen α und $x > 0$ ist daher die α -Potenz wohldefiniert

$$\boxed{x^\alpha := \exp(\alpha \cdot \log(x))}.$$

Es gilt $(xy)^\alpha = x^\alpha \cdot y^\alpha$ und $x^{\alpha+\beta} = x^\alpha \cdot x^\beta$ für alle reellen Zahlen $x > 0$ und $y > 0$. Beachte, daß für natürliche Zahlen α der Wert x^α die übliche Potenz von x ist wegen der Funktionalgleichung der Exponentialfunktion.

3.6 Das mehrdimensionale Integral $\int_{\mathbb{R}^n} f(x) dx$

Wir verallgemeinern jetzt das Standardintegrals auf den n -dimensionalen Fall. Betrachten wir zunächst den Raum $T(\mathbb{R}^n)$ der Treppenfunktionen.

Definition. Wir definieren $T(\mathbb{R}^n)$ induktiv, d.h. wir nehmen an die $T(\mathbb{R}^m)$ seien für $m < n$ bereits definierte Verbände auf $\mathbb{R}^m$. Per Definition liegt dann eine Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$ in $T(\mathbb{R}^n)$, wenn $f(x)$ sich für $x = (x_1, \dots, x_n)$ in der Form schreiben lässt

$$f(x) = \sum_{i=1}^m c_i(x_1, \dots, x_{n-1}) \cdot \chi_{[\xi_{i-1}, \xi_i]}(x_n) + \sum_{j=0}^m d_j(x_1, \dots, x_{n-1}) \cdot \chi_{[\xi_j, \xi_{j+1}]}(x_n)$$

für gewisse Funktionen $c_i, d_j \in T(\mathbb{R}^{n-1})$ und gewisse reelle Zahlen $\xi_0 \leq \dots \leq \xi_m$. Wie in Lemma 3.4 folgt daraus sofort, daß $T(\mathbb{R}^n)$ einen Verband auf $\mathbb{R}^n$ definiert.

Quader. Sei $A = [a_1, b_1] \times \dots \times [a_n, b_n]$ ein beschränkter **Euklidischer Quader** im $\mathbb{R}^n$. Wir nehmen $a_i \leq b_i$ an, lassen aber **degenerierte Quader** mit der Eigenschaft $a_i = b_i$ zu. Durch Induktion nach n zeigt man dann leicht: $\chi_A \in T(\mathbb{R}^n)$, und $T(\mathbb{R}^n)$ wird als Vektorraum von den charakteristischen Funktionen χ_A der kompakten Quader $A \subset \mathbb{R}^n$ erzeugt (Übungsaufgabe). Dies liefert sofort die Eindeutigkeitsaussage im nächsten

Lemma 3.12. *Es gibt eine eindeutig bestimmte $\mathbb{R}$ -lineare monotone Abbildung*

$$I : T(\mathbb{R}^n) \rightarrow \mathbb{R}$$

mit der Eigenschaft: Es gilt $I(\chi_A) = \prod_{i=1}^n |b_i - a_i|$ für kompakte Quader $A = \prod_{i=1}^n [a_i, b_i]$.

Beweis. Die Existenz von I folgt induktiv aus der Existenz einer $\mathbb{R}$ -linearen monotonen Abbildung $T(\mathbb{R}^n) \rightarrow T(\mathbb{R}^{n-1})$, welche χ_A auf $|b_n - a_n| \cdot \chi_{A'}$, $A' = \prod_{i=1}^{n-1} [a_i, b_i]$ abbildet. Die gesuchte relative Abbildung ist

$$f(x) \mapsto \int_{\mathbb{R}} f(x_1, \dots, x_{n-1}, t) dt.$$

Sie ist offensichtlich $\mathbb{R}$ -linear und monoton. Daß sie $T(\mathbb{R}^n)$ in $T(\mathbb{R}^{n-1})$ abbildet, ist klar. Es genügt daher für die Erzeuger $\chi_A(x) = \prod_{i=1}^n \chi_{[a_i, b_i]}(x_i)$ des Vektorraums $T(\mathbb{R}^n)$ die einfache Formel

$$\int_{\mathbb{R}} \chi_A(x_1, \dots, x_{n-1}, t) dt = \chi_{A'}(x_1, \dots, x_{n-1}) \int_{\mathbb{R}} \chi_{[a_n, b_n]}(t) dt = \chi_{A'}(x_1, \dots, x_{n-1}) \cdot |b_n - a_n|.$$

□

Bemerkung. Lemma 3.8 überträgt sich auf $\mathbb{R}^n$. Sei $f \in C(A)$ und $A \subset \mathbb{R}^n$ ein kompakter Quader und $\varepsilon > 0$. Dann existiert⁴ ein $g \in T_A(\mathbb{R})$ mit $\|f - g\| < \varepsilon$ (Supremumsnorm). Für den von den Funktionen $\chi_A(x) \cdot f(x)$, für kompakte Quader A und $f \in C(\mathbb{R}^n)$, aufgespannten $\mathbb{R}$ -Vektorraum $CT(\mathbb{R}^n)$ folgt daher (wie im Beweis von Satz 3.7)

⁴Hinweis: Zerlege dazu A äquidistant in halboffene 'Quader' der Gestalt $\prod_{\nu=1}^n (a_\nu, b_\nu]$; verwende am Rand Durchschnitte der halboffenen Quader mit A . Oder man benutzt das allgemeinere Resultat 8.12.

Satz 3.13. Das Integral $I : T(\mathbb{R}^n) \rightarrow \mathbb{R}$ aus Lemma 3.12 setzt sich auf eindeutige Weise zu einer $\mathbb{R}$ -linearen monotonen Abbildung $I : CT(\mathbb{R}^n) \rightarrow \mathbb{R}$ fort.

Für den Raum der stetigen Funktionen mit kompaktem Träger gilt $C_c(\mathbb{R}^n) = \bigcup_A C_A(\mathbb{R}^n)$, die Vereinigung läuft über die kompakten Quader $A \subset \mathbb{R}^n$. Also $C_c(\mathbb{R}^n) \subset CT(\mathbb{R}^n)$, und I definiert daher eine monotone $\mathbb{R}$ -lineare Abbildung

$$C_c(\mathbb{R}^n) \longrightarrow \mathbb{R}$$

$$I(f) =: \int_{\mathbb{R}^n} f(x) dx ,$$

das **Euklidische Standardintegral** auf $\mathbb{R}^n$. Wir setzen wieder $\int_A f(x) dx := \int_{\mathbb{R}^n} \chi_A(x) f(x) dx$ für Quader A ; und man schreibt oft auch $dx_1 dx_2 \cdots dx_n$ anstelle von dx .

Bemerkung. Für *degenerierte* kompakte Quader $A \subset \mathbb{R}^n$ gilt

$$\boxed{A \text{ degeneriert und } f \in C_c(\mathbb{R}^n) \implies \int_A f(x) dx = 0} .$$

In der Tat ist f beschränkt auf A . Wegen $I(\chi_A) = 0$ folgt diese Behauptung daher aus der Boxungleichung (Abschnitt 3.1).

Translationsinvarianz. Unmittelbar aus Lemma 3.12 folgt für alle $f \in CT(\mathbb{R}^n)$

$$\boxed{\int_{\mathbb{R}^n} f(x + x_0) dx = \int_{\mathbb{R}^n} f(x) dx} .$$

3.7 Monotone Hüllen

Wir lösen uns jetzt vom konkreten Fall $X = \mathbb{R}^n$. Sei X eine beliebige Menge, später meist aber ein metrischer Raum. Wir betrachten auf X jetzt beliebige Funktionen mit Werten in

$$\mathbb{R}^+ = \mathbb{R} \cup \infty \quad \text{oder} \quad \mathbb{R}^- = \mathbb{R} \cup -\infty .$$

Im Gegensatz zu reellwertigen Funktionen auf X , die in natürlicher Weise einen $\mathbb{R}$ -Vektorraum bilden, ist der Raum aller $\mathbb{R}^+$ -wertigen Funktionen auf X nur noch unter Addition und unter Multiplikation mit positiven Skalaren $\lambda \geq 0$ abgeschlossen (Konvention: $0 \cdot \infty = \infty \cdot 0 = 0$).

Definition 3.14. Eine Teilmenge $B(X)$ der $\mathbb{R}^+$ -(oder $\mathbb{R}^-$)-wertigen Funktionen auf X heisst **Halbverband**, wenn $B(X)$ unter den Bildungen

$$\min(f(x), g(x)) \quad \text{und} \quad \max(f(x), g(x))$$

sowie unter Addition und Multiplikation mit reellen Skalaren $\lambda \geq 0$ abgeschlossen ist.

Das einfachste Beispiel für einen Halbverband ist $\mathbb{R}^+$ selbst; hier sei X einelementig. Jeder Verband ist ein Halbverband. Ist $B(X)$ ein Halbverband mit Werten in $\mathbb{R}^\pm$, dann ist $-B(X)$ ein Halbverband mit Werten in $\mathbb{R}^\mp$ wegen $\max(-f, -g) = -\min(f, g)$ etc.

Lemma 3.15. Ein Halbverband $B(X)$ ist ein Verband $\iff B(X) = -B(X)$.

Beweis. Wir zeigen $\Leftarrow$: Alle $f \in B(X)$ haben wegen $B(X) = -B(X)$ Werte in $\mathbb{R}$. Wegen $B(X) = -B(X)$ und den Halbverbandsaxiomen ist $B(X)$ ein $\mathbb{R}$ -Vektorraum. Für $f \in B(X)$ gilt $f = f^+ + f^-$ für $f^+ = \max(f, 0) \geq 0$ und $f^- = \min(f, 0) \leq 0$ in $B(X)$. Nach Annahme gilt $-f^- \in B(X)$; es folgt $|f| = f^+ + (-f^-) \in B(X)$. Also ist $B(X)$ ein Verband. $\square$

Die **monotone Hülle** $B^+(X)$ eines Halbverbandes $B(X)$ mit Werten in $\mathbb{R}^+$ besteht aus allen Funktionen $f : X \rightarrow \mathbb{R}^+$ für die eine punktweise monotone aufsteigende Folge $f_n \in B(X)$ existiert, d.h. $f_1(x) \leq f_2(x) \leq f_3(x) \leq \dots$ für alle $x \in X$, mit $f(x) = \sup\{f_n(x) \mid n \in \mathbb{N}\}$ oder kurz

$$f = \sup_n f_n = \sup f_n.$$

Wir schreiben symbolisch $f_n \nearrow f$, wenn f in diesem Sinne als punktweise ‘monotoner Limes’ von Funktionen f_n definiert ist.

Die **monotone Hülle** $B^-(X)$. Für einen Halbverband $B(X)$ mit Werten in $\mathbb{R}^-$, definiert man analog $f_n \searrow f$ und $B^-(X)$ mittels monoton fallender Limiten. Funktionen in $B^-(X)$ haben dann Werte in $\mathbb{R}^- = \mathbb{R} \cup -\infty$. Alle Aussagen sind analog, so daß wir uns oft auf den Fall von Halbverbänden mit Werten in $\mathbb{R}^+$ beschränken. Für einen Verband $B(X)$ gilt offensichtlich

$$B^-(X) = -B^+(X).$$

Lemma 3.16. Ist $B(X)$ ein Halbverband mit Werten in $\mathbb{R}^+$, dann ist auch $B^+(X) \supseteq B(X)$ ein Halbverband mit Werten in $\mathbb{R}^+$. Der Halbverband $B^+(X)$ ist unter monoton wachsenden Limiten abgeschlossen: Aus $f_i \nearrow f$ für $f_i \in B^+(X)$ folgt $f \in B^+(X)$. Das heisst

$$B^{++}(X) = B^+(X).$$

Beweis. $B(X) \subseteq B^+(X)$, denn für die konstante Folge $f_n = f \in B(X)$ gilt $f_n \nearrow f$.

Für $f_n, g_n \in B(X)$ mit $f_n \nearrow f, g_n \nearrow g$ und $\lambda > 0$ gilt $f_n + g_n \nearrow f + g$ und $\lambda f_n \nearrow \lambda f$. Zum Nachweis von $\max(f_n, g_n) \nearrow \max(f, g)$ beachte

$$\max(f_n, g_n) \leq \max(f_{n+1}, g_n) \leq \max(f_{n+1}, g_{n+1}),$$

da $h \leq g$ die Ungleichungen $\max(h, f) \leq \max(g, f)$ (und $\min(h, f) \leq \min(g, f)$) impliziert. Für $\min(f_n, g_n) \nearrow \min(f, g)$ benutze $\min(f_n, g_n) \leq \min(f_{n+1}, g_n) \leq \min(f_{n+1}, g_{n+1})$.

Für den Beweis der letzten Aussage wähle $f_{ij} \nearrow f_i$ mit $f_{ij} \in B(X)$. Wegen $f_{ij} \leq f_{i,j+1}$ für alle $i + j = n$ gilt dann $F_n := \max_{i+j=n}(f_{ij}) \nearrow f$ sowie $F_n \in B(X)$, also $f \in B^+(X)$. $\square$

3.8 Abstrakte Integrale

Satz 3.17. Ist $B(X) = C_c(X)$ der Verband der stetigen Funktionen mit kompaktem Träger auf einem metrischen Raum (X, d) und ist $I : B(X) \rightarrow \mathbb{R}$ monoton und $\mathbb{R}$ -linear, dann ist I **halbstetig**, d.h. für eine monotone Folge $f_n \in B(X)$ mit Grenzfunktion $f \in B(X)$ konvergiert $\lim_n I(f_n)$ monoton gegen $I(f)$.

Beweis. Ersetzt man f_n durch $f_n - f$, kann man oBdA $f = 0$ annehmen. Ersetzt man f_n durch $-f_n$, ist oBdA f_n monoton fallend. Der Träger aller f_n ist damit im kompakten Träger von f_1 enthalten. Aus dem *Satz von Dini* (Satz 2.26) folgt daher die gleichmässige Konvergenz der Folge f_n auf X . Aus der Monotonie von I und Lemma 3.1 folgt damit die Behauptung. $\square$

Sei $B(X)$ ein Halbverband (oBdA) mit Werten in $\mathbb{R}^+$. Für eine Abbildung I von $B(X)$ in einen anderen Halbverband mit Werten in $\mathbb{R}^+$ betrachten wir folgende Eigenschaften

1. **Semilinearität:** $I(f + g) = I(f) + I(g)$ sowie $I(\lambda \cdot f) = \lambda \cdot I(f)$ für reelles $\lambda \geq 0$.
2. **Monotonie:** $f \leq g \implies I(f) \leq I(g)$
3. **Halbstetigkeit:** $I(f_n) \nearrow I(f)$ für $f_n \nearrow f$ und $f, f_n \in B(X)$.

Definition 3.18. Ein **abstraktes Integral** auf $B(X)$ ist eine semilineare monotone und halbstetige Abbildung $I : B(X) \rightarrow \mathbb{R}^+$ (d.h. I erfüllt die obigen Eigenschaften 1, 2 und 3).

Beispiel 3.19. Das Euklidische Standardintegral definiert auf $C_c(\mathbb{R}^n)$ ein abstraktes Integral (Satz 3.17).

Integrale auf Verbänden. Aus 1) für $\lambda = 0$ folgt $0 \in B(X)$ und aus 1) Additivität folgt $I(0) = 0$. Ist $B(X)$ sogar ein Verband, gilt $f \in B(X) \implies -f \in B(X)$. Aus 1) folgt dann $I(-f) = -I(f)$ und insbesondere ist $I(f) < \infty$ reell. Die Semilinearität von I ist damit auf Verbänden äquivalent zur **$\mathbb{R}$ -Linearität** der Abbildung I . Ist $B(X)$ ein Verband, folgt die Monotonie (Eigenschaft 2) bereits aus der schwächeren Bedingung $f \geq 0 \implies I(f) \geq 0$, indem man die Hilfsfunktion $g - f \geq 0$ betrachtet.

Satz 3.20 (Fortsetzungssatz). Sei $B(X)$ ein Halbverband mit Werten in $\mathbb{R}^+$ und sei I ein abstraktes Integral⁵ auf $B(X)$. Dann setzt sich I auf eindeutige Weise zu einem abstrakten Integral I^+

$$I^+ : B^+(X) \rightarrow \mathbb{R}^+$$

der monotonen Hülle $B^+(X)$ fort. Analoges gilt, wenn man die $+$ durch $-$ -Zeichen ersetzt.

Beweis. Existenz. Für $g_n \nearrow g$ und $g_n \in B$ setzen wir

$$I^+(g) := \sup_n I(g_n),$$

müssen aber zeigen, dass dieser Grenzwert $I^+(g)$ nicht von der Wahl der Folge $f_n \nearrow g$ mit $f_n \in B$ abhängt. Mischen der Folgen wie im Beweis von Satz 3.7 zerstört die Monotonie der Folgen. Wir werden daher die Halbstetigkeit von I in Form von Lemma 3.21 ins Spiel bringen! Dieses Lemma zeigt $I(f_n) \leq \sup I(g_n)$ und damit $\sup I(f_n) \leq \sup I(g_n)$ im Limes. Vertauscht man die Rollen von f_n und g_n , folgt wie gewünscht $\sup I(f_n) = \sup I(g_n)$.

⁵oder allgemeiner: Eine beliebige semilineare monotone halbstetige Abbildung $I : B(X) \rightarrow \tilde{B}(Y)$ zwischen zwei Halbverbänden mit Werten (jeweils) in $\mathbb{R}^+$ setzt sich eindeutig auf die monotonen Hüllen fort.

3 Integration

Eindeutigkeit. Für $g \in B^+(X) = B^+$ wähle $g_n \in B(X) = B$ mit $g_n \nearrow g$. Ein abstraktes Integral I^+ auf B^+ ist als Fortsetzung von I durch $I^+(g) = \sup I^+(g_n) = \sup I(g_n)$ (wegen der Halbstetigkeit) eindeutig festgelegt.

Semi-Linearität. Diese folgt durch Limesbildung aus der Semilinearität von I .

Monotonie: $f \leq g \implies I^+(f) \leq I^+(g)$. [Für $f, g \in B^+$ wähle $f_n \nearrow f$ mit $f_n \in B$; aus Lemma 3.21 folgt wieder $I(f_n) \leq I^+(g)$ und damit $I^+(f) \leq I^+(g)$ im Limes $n \rightarrow \infty$].

Halbstetigkeit. Für $f_n \in B^+$ mit $f_n \nearrow f \in (B^+)^+ = B^+$ gilt $\sup I^+(f_n) = I^+(f)$. Benutze dazu die Hilfsfolge $F_n \nearrow f$ definiert durch $F_n = \max_{i+j=n} f_{ij} \in B$ wie in Lemma 3.16. Wegen $F_n \leq f_n \leq f$ gilt $I^+(f) := \sup I(F_n) \leq \sup I^+(f_n) \leq I^+(f)$. Also $\sup I^+(f_n) = I^+(f)$ (Halbstetigkeit). $\square$

Lemma 3.21. Sei $g_1 \leq g_2 \leq g_3 \leq \dots$ eine monoton wachsende Folge von Funktionen g_n eines Halbverbandes $B(X)$ mit Werten in $\mathbb{R}^+$ und I ein abstraktes Integral auf $B(X)$. Dann gilt: $B(X) \ni f \leq \sup_n g_n \implies I(f) \leq \sup_n I(g_n)$.

Beweis. Aus $f \leq \sup_n g_n$ folgt $f_n \nearrow f \in B(X)$ für $f_n := \min(f, g_n) \in B(X)$, also $I(f) = \sup I(f_n)$ wegen der Halbstetigkeit von I . Wegen $f_n \leq g_n$ folgt aus der Monotonie $I(f_n) \leq I(g_n)$. Dies zeigt $\sup_n I(f_n) \leq \sup_n I(g_n)$. $\square$

Folgerung. Das Euklidische Standardintegral

$$I = \int_{\mathbb{R}^n} : C_c(\mathbb{R}^n) \rightarrow \mathbb{R}$$

läßt sich zu einem abstrakten Integral $I^- : C_c^-(\mathbb{R}^n) \rightarrow \mathbb{R}$ fortsetzen.

Die monotone Hülle $C_c^-(\mathbb{R}^n)$ aller monoton fallenden Folgen von Funktion f in $C_c(\mathbb{R}^n)$ enthält wegen Beispiel 2.25 alle Funktionen der Gestalt $\chi_A(x) \cdot f(x)$, für beliebige abgeschlossene Teilmengen A von $\mathbb{R}^n$ und beliebige nichtnegative Funktionen $f \in C_c(\mathbb{R}^n)$. Wir benutzen dies später bei dem Beweis der allgemeinen **Substitutionsregel** (Satz 4.27).

4 Differentiation

Eine Teilmenge X von $\mathbb{R}^n$ heisst **zulässig**, wenn für jeden Punkt ξ aus X ein abgeschlossener nicht degenerierter Quader Q existiert mit $\xi \in Q \subset X$.

Eine Teilmenge X von $\mathbb{R}^n$ (oder allgemeiner eine Teilmenge X eines metrischen Raumes) heisst **offen**, wenn für jeden Punkt $\xi \in X$ eine Kugel $B_r(\xi) = \{x \mid d(x, \xi) < r\}$ vom Radius $r > 0$ existiert mit $\xi \in B_r(\xi) \subset X$.

Jede offene Teilmenge von $\mathbb{R}^n$ ist zulässig. Der Durchschnitt von zwei offenen Mengen ist wieder offen. Dies gilt jedoch nicht für zulässige Mengen.

4.1 Das Landausymbol

Wir erklären in diesem Abschnitt, was es bedeuten soll, daß eine Funktion f *schneller in einem Punkt ξ gegen Null geht als jede von Null verschiedene lineare Funktion*. Man schreibt in diesem Fall nach Landau: $f(x) = o(x - \xi)$. Dies soll mit Lemma 4.1 präzise gemacht werden.

Sei dazu X eine Teilmenge des Euklidischen Raumes $\mathbb{R}^n$ und $\xi \in X$ ein gegebener Punkt. Für eine Funktion

$$f : X \rightarrow \mathbb{R}^m$$

schreiben wir

$$\boxed{f(x) = o(x - \xi)},$$

wenn eine Funktion $H : X \rightarrow \mathbb{R}^m$ existiert, die stetig im Punkt ξ ist mit $H(\xi) = 0$, so daß gilt

$$\boxed{f(x) = \|x - \xi\|_{\mathbb{R}^n} \cdot H(x)}.$$

Bemerkung. Für Funktionen $f_1(x)$ und $f_2(x)$ auf X mit Werten in $\mathbb{R}^m$ sieht man sofort: $f_i(x) = o(x - \xi) \implies \alpha \cdot f_1(x) + \beta \cdot f_2(x) = o(x - \xi)$ für alle $\alpha, \beta \in \mathbb{R}$.

Lemma 4.1. Sei $X \subset \mathbb{R}^n$ eine zulässige Teilmenge. Sei $L : \mathbb{R}^n \rightarrow \mathbb{R}^m$ eine $\mathbb{R}$ -lineare Abbildung und ξ aus X . Gilt $L(x - \xi) = o(x - \xi)$ auf X , dann ist $L = 0$.

Beweis. OBdA ist X ein Quader und durch eine Translation ist oBdA $\xi = 0$. Für $x \in X$ gilt

$$L(x) = \|x\| \cdot H(x)$$

4 Differentiation

sowie: $\forall \varepsilon > 0 \exists \delta = \delta(\varepsilon) > 0$ mit $\|x\| < \delta \implies \|H(x)\| < \varepsilon$. Mit x liegt auch x/n im Quader X , für alle $n \geq 1$ in $\mathbb{N}$. Da L linear ist, gilt $L(x/n) = L(x)/n$. Ist $L(x) \neq 0$, folgt daraus $H(x) = H(x/n)$ wegen $\|x/n\| = \|x\|/n$. Für große n gilt $\|x/n\| < \delta(\varepsilon)$ und damit $\|H(x/n)\| < \varepsilon$, also $0 \leq \|H(x)\| < \varepsilon$. Da dies bei festem $x \in X$ für beliebiges $\varepsilon > 0$ gilt, folgt $\|H(x)\| = 0$ und damit $H(x) = 0$ (für alle $x \in X$). Dies zeigt $L(x) = 0$ für alle $x \in X$.

Da der Quader X eine Basis des Vektorraums $\mathbb{R}^n$ enthält, folgt daraus $L = 0$. $\square$

4.2 Differenzierbarkeit

Sei $X \subset \mathbb{R}^n$ eine zulässige Teilmenge, und sei

$$f : X \rightarrow \mathbb{R}^m$$

eine Funktion. Unter dieser Annahme machen wir nun die folgende

Definition 4.2. f heisst **differenzierbar** im Punkt $\xi \in X$, wenn es eine $\mathbb{R}$ -lineare Abbildung $L : \mathbb{R}^n \rightarrow \mathbb{R}^m$ gibt, so dass gilt (*)

$$f(x) - f(\xi) - L(x - \xi) = o(x - \xi).$$

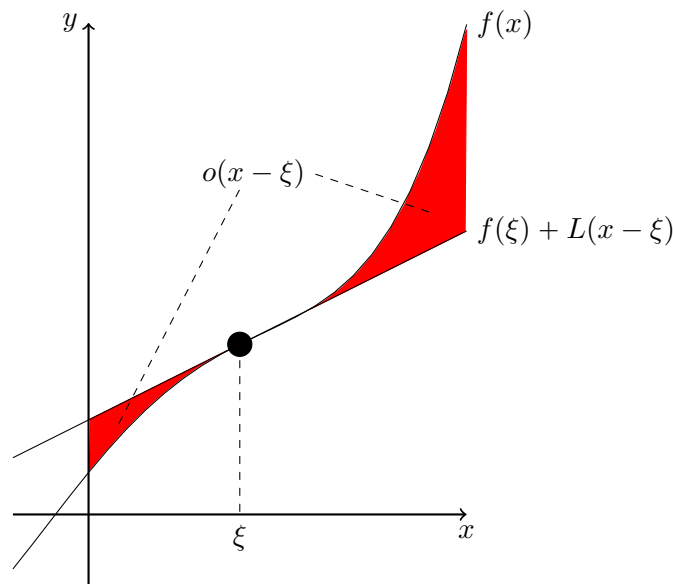
Gilt (*), dann ist die (stetige) lineare Abbildung

$$L : \mathbb{R}^n \rightarrow \mathbb{R}^m$$

eindeutig bestimmt durch f und ξ , und man nennt L das Differential der Abbildung f im Punkt ξ und schreibt

$$L = Df(\xi).$$

Ableitung/Tangente im eindimensionalen Fall:



Beweis. Angenommen zwei lineare Abbildungen L_1, L_2 erfüllen Eigenschaft (*). Dann hätte die Differenz $L = L_1 - L_2$ die Eigenschaft $L(x - \xi) = o(x - \xi)$ auf X . Nach Lemma 4.1 folgt daraus $L = 0$. $\square$

Bemerkung. Im Fall $X \subseteq \mathbb{R}$ und $f : X \rightarrow \mathbb{R}$, d.h. im Fall der Dimensionen $n = m = 1$, schreibt man $Df(\xi)(x) = m \cdot x$ und benutzt für die *Tangentensteigung* m der Funktion $f(x)$ im Punkt ξ folgende synonyme Notationen

$$m = \dot{f}(\xi) = f'(\xi) = \frac{df(x)}{dx} \Big|_{x=\xi} = \frac{df(\xi)}{dx} = \frac{d}{dx} f(\xi) = \frac{\partial f(\xi)}{\partial x} = \frac{\partial f}{\partial x}(\xi) = \partial_x f(\xi).$$

Definition 4.3. Eine Funktion $f : X \rightarrow \mathbb{R}^m$ heißt differenzierbar, wenn sie in jedem Punkt $\xi \in X$ differenzierbar ist.

Beispiel 4.4. Seien $(\mathbb{R}^n, \|\cdot\|), (\mathbb{R}^m, \|\cdot\|)$ Euklidische Räume und $c \in \mathbb{R}^m$ eine Konstante und $L : \mathbb{R}^n \rightarrow \mathbb{R}^m$ eine $\mathbb{R}$ -lineare Abbildung. Dann ist die **affin lineare** Abbildung $f(x) = c + L(x)$ differenzierbar, und hat in jedem Punkt $\xi \in \mathbb{R}^n$ die Ableitung

$$Df(\xi) = L,$$

denn $H(x) := f(x) - f(\xi) - L(x - \xi)$ ist Null, insbesondere gilt also $H(x) = o(x - \xi)$.

Eine Reduktion. Eine Abbildung $f : X \rightarrow \mathbb{R}^m$ mit Werten in dem Euklidischen Vektorraum $\mathbb{R}^m$ wird beschrieben durch m reellwertige Abbildungen $f_i : X \rightarrow \mathbb{R}$ für $i = 1, \dots, m$ (die sogenannten Komponenten f_i von f)

$$f(x) = \begin{pmatrix} f_1(x) \\ f_2(x) \\ \vdots \\ f_{m-1}(x) \\ f_m(x) \end{pmatrix}$$

Lemma 4.5. $f(x)$ ist differenzierbar im Punkt ξ genau dann, wenn jede der Komponenten $f_1(x), \dots, f_m(x)$ differenzierbar ist im Punkt ξ .

Beweis. Die durch $f(x) - f(\xi) - L(x - \xi) = \|x - \xi\| \cdot H(x)$ für $x \neq \xi$ definierte vektorwertige Funktion $H(x)$ konvergiert gegen Null für $x \rightarrow \xi$ genau dann, wenn die Komponenten $H_1(x), \dots, H_m(x)$ von $H(x)$ gegen Null konvergieren für $x \rightarrow \xi$. Analog ist L (stetig und) $\mathbb{R}$ -linear genau dann, wenn alle Komponenten $L_1, \dots, L_m$ (stetige) $\mathbb{R}$ -Linearformen sind. $\square$

Lemma 4.6. Ist f differenzierbar im Punkt ξ , dann ist f stetig im Punkt ξ .

Beweis. Sowohl $d(x, \xi) = \|x - \xi\|$ als auch $H(x)$ sind stetig im Punkt ξ . Summen und Produkte in ξ stetiger reellwertiger Funktionen sind stetig in ξ . Daher sind die Komponenten der Funktion $f(x) = f(\xi) + L(x - \xi) + \|x - \xi\| \cdot H(x)$ stetig im Punkt ξ . Dasselbe gilt daher auch für $f(x)$. $\square$

4 Differentiation

Lemma 4.7 (Kettenregel). Seien $X \subseteq \mathbb{R}^n, Y \subseteq \mathbb{R}^m$ nicht degenerierte Quader (oder allgemeiner zulässige Mengen) und $Z = \mathbb{R}^k$. Gegeben seien Abbildungen f, g sowie $\xi \in X$ mit

$$X \xrightarrow{f} Y \xrightarrow{g} Z$$

$$\xi \xrightarrow{f} \eta = f(\xi)$$

Dann gilt: Ist f differenzierbar im Punkt ξ und ist g differenzierbar im Punkt $\eta = f(\xi)$, dann ist die Zusammensetzung $g \circ f$ differenzierbar im Punkt ξ und es gilt für

$$\mathbb{R}^n \xrightarrow{Df(\xi)} \mathbb{R}^m \xrightarrow{Dg(\eta)} \mathbb{R}^k$$

die Kettenregel: $\boxed{D(g \circ f)(\xi) = Dg(\eta) \circ Df(\xi)}$.

Beweis. Die Differenzierbarkeit von f im Punkt ξ zeigt

$$(**) \quad f(x) = f(\xi) + Df(\xi)(x - \xi) + H(x) \cdot \|x - \xi\|$$

für eine Funktion $H(x)$ mit $\lim_{x \rightarrow \xi} \|H(x)\| = 0$. Analog gilt

$$(*) \quad g(y) = g(\eta) + Dg(\eta)(y - \eta) + \tilde{H}(y) \cdot \|y - \eta\|$$

für eine Funktion $\tilde{H}(y)$ mit $\lim_{y \rightarrow \eta} \|\tilde{H}(y)\| = 0$. Setzt man $(**)$ in $(*)$ ein, ergibt sich

$$(g \circ f)(x) = g(f(x)) \stackrel{(*)}{=} g(\eta) + Dg(\eta)(f(x) - \eta) + \tilde{H}(f(x)) \cdot \|f(x) - \eta\|$$

$$\stackrel{(**)}{=} g(\eta) + Dg(\eta) \left(Df(\xi)(x - \xi) + H(x) \cdot \|x - \xi\| \right) + \tilde{H}(f(x)) \cdot \|Df(\xi)(x - \xi) + H(x) \cdot \|x - \xi\|\|$$

$$= (g \circ f)(\xi) + \left(Dg(\eta) \circ Df(\xi) \right) (x - \xi) + H_1(x) \cdot \|x - \xi\|$$

für

$$H_1(x) = Dg(\eta)(H(x)) + \tilde{H}(f(x)) \cdot \frac{\|Df(\xi)(x - \xi) + H(x) \cdot \|x - \xi\|\|}{\|x - \xi\|}.$$

Zweimaliges Anwenden der Dreiecksungleichung sowie die Ungleichung $\|L(v)\| \leq \|L\|\|v\|$ für $\mathbb{R}$ -lineare Abbildungen L zeigt

$$\|H_1(x)\| \leq \|Dg(\eta)\| \cdot \|H(x)\| + \|\tilde{H}(f(x))\| \cdot (\|Df(\xi)\| + \|H(x)\|).$$

f ist nach Lemma 4.6 stetig im Punkt ξ , also $\lim_{x \rightarrow \xi} \|\tilde{H}(f(x))\| = \lim_{y \rightarrow \eta} \|\tilde{H}(y)\| = 0$ für $y = f(x)$. Wegen $\lim_{x \rightarrow \xi} \|H(x)\| = 0$ folgt insgesamt $\lim_{x \rightarrow \xi} \|H_1(x)\| = 0$. Also ist $g \circ f$ differenzierbar im Punkt ξ mit dem Differential $L = Dg(\eta) \circ Df(\xi)$. $\square$

4.3 Die Jacobi-Matrix

Sei $X \subseteq \mathbb{R}^n$ eine zulässige Teilmenge im Euklidischen Raum und

$$f: X \rightarrow \mathbb{R}^m$$

eine im Punkt $\xi \in X$ differenzierbare Abbildung.

Für $t \in (-\varepsilon, 0]$ oder $[0, \varepsilon)$ liegt $i_\nu(t) := \xi + t \cdot e_\nu$ in X für genügend kleines $\varepsilon > 0$ (hierbei sei e_ν der ν -te Basisvektor des $\mathbb{R}^n$). Sei $p_\mu: \mathbb{R}^m \rightarrow \mathbb{R}$ die Projektion auf die μ -te Koordinate. Die Zusammensetzung

$$p_\mu \circ f \circ i_\nu(t) = f_\mu(\xi_1, \dots, \xi_{\nu-1}, \xi_\nu + t, \xi_{\nu+1}, \dots, \xi_n)$$

ist reellwertig und definiert auf einem zulässigen Intervall in $\mathbb{R}$; und ist nach der Kettenregel differenzierbar im Punkt $t = 0$ als Zusammensetzung von i_ν (affin linear), f und p_μ (linear). Die Kettenregel berechnet die Ableitung nach t im Punkt $t = 0$

$$\left. \frac{d}{dt} f_\mu(\xi_1, \dots, \xi_{\nu-1}, \xi_\nu + t, \xi_{\nu+1}, \dots, \xi_n) \right|_{t=0} = Dp_\mu(\eta) \circ Df(\xi) \circ Di_\nu(0).$$

Beispiel 4.4 berechnet zwei der Terme der rechten Seite. Bezeichne $J(f)(\xi)$ die der linearen Abbildung $D(f)(\xi): \mathbb{R}^n \rightarrow \mathbb{R}^m$ zugeordnete $m \times n$ -Matrix. Dies ergibt für die linke Seite, welche man die **partielle Ableitung** von f im Punkt ξ nennt, in Matrixschreibweise die Formel

$$\frac{\partial f_\mu(\xi)}{\partial x_\nu} = (0, \dots, 1, \dots, 0) \cdot Jf(\xi) \cdot \begin{pmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{pmatrix} = Jf(\xi)_{\mu\nu}.$$

Auf der rechten Seite steht dann der Matrixkoeffizient $Jf(\xi)_{\mu\nu}$ von $Jf(\xi)$ an der μ, ν -ten Stelle.

*Wir fassen zusammen: Differenzierbarkeit im Punkt ξ impliziert partielle Differenzierbarkeit im Punkt ξ , und die Ableitung $Df(\xi)$ wird als lineare Abbildung gegeben durch die Matrix $Jf(\xi)$ (**Jacobimatrix**) der partiellen Ableitungen für $\nu = 1, \dots, n$ und $\mu = 1, \dots, m$*

$$Df(\xi) \leftrightarrow Jf(\xi) = \left(\frac{\partial f_\mu(\xi)}{\partial x_\nu} \right).$$

Die **Kettenregel** (Lemma 4.7), via Matrixmultiplikation $J(g \circ f)(\xi)_{\lambda\nu} = \sum_{\mu=1}^m Jg(\eta)_{\lambda\mu} Jf(\xi)_{\mu\nu}$, liefert daher

$$\frac{\partial (g \circ f)_\lambda(\xi)}{\partial x_\nu} = \sum_{\mu=1}^m \frac{\partial g_\lambda(\eta)}{\partial y_\mu} \cdot \frac{\partial f_\mu(\xi)}{\partial x_\nu}.$$

Notation. Wir schreiben oft $\partial_\nu g$ anstelle von $\frac{\partial g}{\partial x_\nu}$ und manchmal $Df(\xi)$ anstatt $Jf(\xi)$.

4.4 Extremwerte

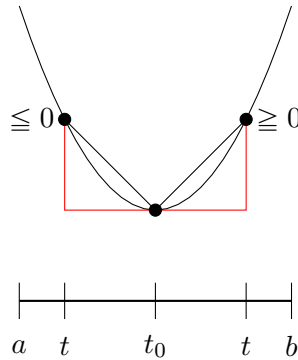
Jede differenzierbare Funktion $h : [a, b] \rightarrow \mathbb{R}$ ist eine stetige Funktion (Lemma 4.6) und nimmt damit auf dem Intervall $[a, b]$ Minimum und Maximum an (Satz 2.10). Sei

$$t_0 \in (a, b)$$

ein *innerer* Punkt, in dem h sein Minimum annimmt. Wählt man eine Folge $t \rightarrow t_0$, deren Glieder alle von t_0 verschieden sind, dann impliziert die Differenzierbarkeit von h im Punkt t_0

$$\frac{h(t) - h(t_0)}{t - t_0} \longrightarrow h'(t_0) + \lim_{t \rightarrow t_0} \left(H(t) \cdot \frac{|t - t_0|}{t - t_0} \right) = h'(t_0) .$$

wegen $Dh(t_0) = H'(t_0)(t - t_0)$ und $\lim_{t \rightarrow t_0} H(t) = 0$. Der Zähler $h(t) - h(t_0)$ der linken Seite ist nach Annahme nicht negativ. Wählt man eine Folge von Punkten $t \in (0, 1)$ für die alle $t - t_0 > 0$ sind, folgt daher im Limes $h'(t_0) \geq 0$.



Wählt man eine Folge, deren Glieder $t - t_0 < 0$ sind, folgt $h'(t_0) \leq 0$. Wegen $t_0 \in (0, 1)$ sind beide Möglichkeiten von Folgen realisierbar. Aus $t_0 \in (a, b)$ folgt somit (*)

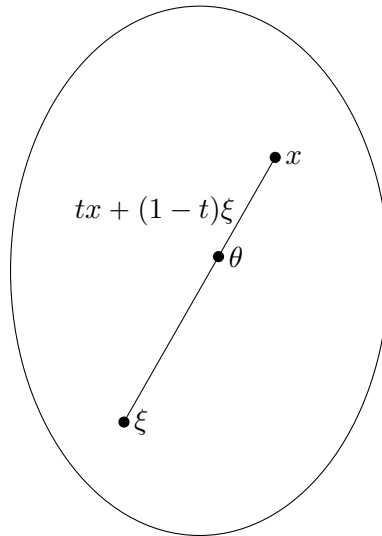
$$\boxed{h'(t_0) = 0} .$$

Satz 4.8 (Mittelwertsatz). Sei X zulässig im $\mathbb{R}^n$ und $f : X \rightarrow \mathbb{R}$ differenzierbar. Seien $x, \xi \in X$ feste Punkte und die Verbindungsgerade $\{tx + (1 - t)\xi \mid t \in [0, 1]\}$ liege ganz in X (z.B. wenn X ein Quader ist). Dann gibt einen Punkt $\theta \in X$

$$\theta = tx + (1 - t)\xi \quad , \quad 0 < t < 1$$

mit der Eigenschaft

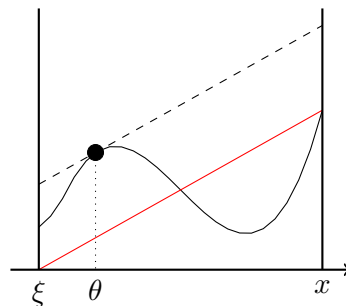
$$\boxed{f(x) - f(\xi) = Df(\theta)(x - \xi)} .$$



Beweis. Ist f affin linear, gilt diese Aussage für alle x und θ . Daher kann man eine geeignete affin lineare Abbildung von f subtrahieren und oBdA annehmen

$$f(x) = f(\xi) = 0.$$

Für $h(t) = f(tx + (1-t)\xi)$ liefert die Kettenregel die Formel $h'(t_0) = Df(\theta)(x - \xi)$ mit $\theta = t_0x + (1-t_0)\xi$. Dies reduziert auf folgende *Aufgabe*: Suche für die differenzierbare Funktion $h : [0, 1] \rightarrow \mathbb{R}$ mit $h(0) = h(1) = 0$ ein $0 < t_0 < 1$ mit $h'(t_0) = 0$. Die *Lösung*: Für ein Maximum oder Minimum von h in $t_0 \in (0, 1)$ gilt $h'(t_0) = 0$ nach (*). Beide Extremwerte existieren, da h stetig und $[0, 1]$ folgenkompakt ist. Wegen $h(0) = h(1) = 0$ existiert mindestens ein Extremwert in einem inneren Punkt $t_0 \in (0, 1)$. $\square$



Folgerung. Gilt $Df(\xi) = 0$ für alle $\xi \in U$ und ist U offen, dann ist f lokal konstant¹ auf U .

Lemma 4.9 (Extremwerte). Sei f eine differenzierbare reellwertige Funktion auf einer offenen Teilmenge U von $\mathbb{R}^n$. Nimmt f in $\xi \in U$ ein Maximum (analog Minimum) an, dann verschwindet die Jacobimatrix im Punkt ξ (man nennt solche Punkte **kritische Punkte**)

$$\boxed{f(\xi) = \max_{x \in U} f(x) \implies Df(\xi) = 0}.$$

¹d.h., für jeden Punkt $x \in U$ gibt es eine offene Kugel um x in U , auf der f konstant ist.

Beweis. $x_i = \xi_i$ ist ein Extremwert der eingeschränkten Funktion $f(\xi_1, \dots, x_i, \dots, \xi_n)$ für festes ξ . Daher gilt $\partial_i f(\xi_1, \dots, \xi_i, \dots, \xi_n) = \partial_i f(\xi) = 0$ nach (*) für alle i , also $Df(\xi) = 0$. $\square$

Die Umkehrung gilt bekanntlich nicht! Die Funktion $f(x) = -x^3$ hat einen kritischen Punkt bei $x = 0$, obwohl an dieser Stelle kein Extremwert vorliegt. Aber es gilt

Lemma 4.10. Sei $h: [0, r] \rightarrow \mathbb{R}$ eine zweimal stetig differenzierbare Funktion. Gilt $h'(0) = 0$ und $h''(\eta) < 0$ für alle $\eta \in (0, r)$, dann gilt $h(t) < h(0)$ für alle $0 < t \leq r$.

Beweis. Aus dem Mittelwertsatz folgert man die Existenz von Punkten $0 < \eta < \theta < t$ mit $h(t) - h(0) = t \cdot h'(\theta)$ und $h'(\theta) - h'(0) = \theta \cdot h''(\eta)$. Aus $h'(0) = 0$ folgt² daher $h'(\theta) < 0$, und damit $h(t) - h(0) < 0$. $\square$

4.5 Symmetrie der Hessematrix

Sei $U \subseteq \mathbb{R}^n$ zulässig und $f: U \rightarrow \mathbb{R}$ sei in $C^2(U)$, d.h. f sei eine zweimal stetig partiell differenzierbare reellwertige Funktion auf U . Damit sei gemeint, daß f zweimal partiell differenzierbar auf U ist für alle Koordinatenrichtungen, und daß diese partiellen Ableitungen stetige Funktionen auf U definieren. Unter dieser Annahme ist die **Hessematrix** $Hess_\xi(f)$

$$Hess_\xi(f) = \left(\frac{\partial^2 f(\xi)}{\partial x_i \partial x_j} \right)$$

als reelle $n \times n$ -Matrix für alle Punkte $\xi \in U$ definiert. Die Koeffizienten $Hess_x(f)_{ij}$ der Hessematrix sind nach Annahme stetige reellwertige Funktionen der Variable $x \in U$.

Satz 4.11. Unter den obigen Annahmen an f und U ist die Hessematrix $Hess_\xi(f)$ eine symmetrische reelle $n \times n$ -Matrix für alle $\xi \in U$.

Beweis. Es genügt der Fall einer Funktion in zwei Variablen $f(x, y)$. Wäre $f_{xy}(\xi) \neq f_{yx}(\xi)$, gäbe es aus Stetigkeitsgründen einen nicht degenerierten kleinen Quader Q um den Punkt ξ mit

$$\partial_x \partial_y f(\theta) \neq \partial_y \partial_x f(\eta) \quad \text{für alle } \theta, \eta \in Q.$$

Setzt man $g(x) = f(x, x_2) - f(x, \xi_2)$, ist $F := f(x_1, x_2) - f(x_1, \xi_2) - f(\xi_1, x_2) + f(\xi_1, \xi_2)$ gleich $F = g(x_1) - g(\xi_1)$. Für $x = (x_1, x_2) \in Q$ liefert zweimaliges Anwenden des Mittelwertsatzes

$$F = (x_1 - \xi_1) \cdot g'(\theta_1) = (x_1 - \xi_1)(x_2 - \xi_2) \cdot \partial_y \partial_x f(\theta_1, \theta_2).$$

Hierbei weiß man $\theta_1 \in (x_1, \xi_1)$ und $\theta_2 \in (x_2, \xi_2)$, also $\theta = (\theta_1, \theta_2) \in Q$. Ebenso gilt auch $F = h(x_2) - h(\xi_2)$ für $h(y) = f(x_1, y) - f(\xi_1, y)$. Dies liefert vollkommen analog

$$F = (x_1 - \xi_1)(x_2 - \xi_2) \cdot \partial_x \partial_y f(\eta_1, \eta_2)$$

für ein $\eta = (\eta_1, \eta_2) \in Q$. Offensichtlich ein Widerspruch, da $x_2 \neq \xi_2$ und $x_1 \neq \xi_1$ gewählt werden kann und dann die rechten Seiten, die beide F berechnen, verschieden wären! $\square$

²Im Fall $h(0) = h'(0) = 0$ folgt ausserdem $|h(t)| \leq t^2 \cdot |h''(\eta)|$.

4.6 Lokale Maxima

Eine symmetrische $n \times n$ -Matrix H mit reellen Koeffizienten H_{ij} heisst positiv definit und man schreibt $H > 0$, wenn für alle Vektoren $v = (v_1, \dots, v_n) \neq 0$ gilt

$${}^T_v H v = \sum_{1 \leq i, j \leq n} v_i v_j H_{ij} > 0.$$

Ist $-H$ positiv definit, nennt man H negativ definit und schreibt $H < 0$.

Satz 4.12. Sei U zulässig im $\mathbb{R}^n$ und $f \in C^2(U)$. Gilt

$$Df(\xi) = 0 \quad \text{und} \quad \text{Hess}_\xi(f) < 0 \quad (\text{resp. } \text{Hess}_\xi(f) > 0),$$

dann ist $f(\xi)$ ein lokales **striktes Maximum (Minimum)** von f auf jedem nichtdegenerierten Quader Q in U , der ξ enthält. Insbesondere gibt es eine offene Teilmenge V von U , welche ξ enthält, so dass gilt

$$f(\xi) = \max_{x \in V} f(x)$$

(resp. $f(\xi) = \min_{x \in V} f(x)$).

Setzt man $H(x) := \text{Hess}_x(f)$, gilt unter den Voraussetzungen von Satz 4.12

Lemma 4.13. Ist $H(\xi) < 0$, dann gibt es eine Konstante $C > 0$ und ein $r > 0$, so daß für alle $x \in Q$ in der offenen Kugel um ξ vom Radius r gilt

$${}^T_v H(x) v \leq -C \|v\|^2.$$

Beweis. ObdA genügt es dazu die Menge $S \subset \mathbb{R}^n$ aller Vektoren v von der Länge 1 zu betrachten, und die x aus einer abgeschlossenen beschränkten Kugel K von positivem Radius r um ξ . Dann ist $S \times K$ folgenkompakt und die Aussage folgt aus Satz 2.10, vorausgesetzt ${}^T_v H(x) v < 0$ gilt für alle $(v, x) \in S \times K$. Angenommen dies wäre nicht der Fall, egal wie klein man r wählt. Dann existiert eine Folge $x_m \rightarrow \xi$ und Vektoren $v_m \neq 0$ der Länge 1 mit

$${}^T_{v_m} H(x_m) v_m \geq 0.$$

Da die Einheitskugel S abgeschlossen und beschränkt und damit folgenkompakt ist, kann man durch Übergang zu einer Teilfolge annehmen $v_m \rightarrow v$ für einen Vektor v der Länge 1. Daraus folgt im Limes $m \rightarrow \infty$

$${}^T_v H(\xi) v \geq 0, \quad v \neq 0$$

im Widerspruch zur Annahme. □

Zum Beweis des Satzes 4.12 fixiere $Q \subseteq U$, $r > 0$ und $v \in S$ mit $\xi + tv \in Q$ für $0 \leq t \leq r$.

4 Differentiation

Beweis. Setze $h(t) = f(\xi + tv)$. Aus $h'(t) = \sum_{i=1}^n v_i \partial_i f(\xi + tv)$ und $Df(\xi) = 0$ folgt dann $h'(0) = 0$. Ein weiteres Anwenden der Kettenregel liefert

$$h''(t) = \sum_{i=1}^n \sum_{j=1}^n v_i v_j \partial_j \partial_i f(\xi + tv) = {}^T v \text{Hess}_{\xi+tv}(f) v .$$

Lemma 4.13 zeigt $h''(t) < 0$ für $v \in S$ und $t \in (0, r)$. Der Satz folgt damit aus Lemma 4.10. $\square$

4.7 Der Hauptsatz

Satz 4.14. Sei $a < b$ und sei $f: [a, b] \rightarrow \mathbb{R}$ eine stetige Funktion. Dann ist

$$F(x) := \int_{[a,x]} f(t) dt$$

eine differenzierbare Funktion auf dem Intervall $[a, b]$, und es gilt

$$F'(x) = f(x) .$$

Jede andere differenzierbare Funktion $G(x)$ auf $[a, b]$ mit der Eigenschaft $G'(x) = f(x)$ (solche Funktionen G nennt man **Stammfunktionen** von f) unterscheidet sich von $F(x)$ um eine reelle additive Konstante C , d.h. es gilt $F(x) = G(x) + C$.

Beweis. Sei $\xi \in [a, b]$. Für $h(x) = F(x) - F(\xi) - f(\xi) \cdot (x - \xi)$ ist $h(x) = o(x - \xi)$ zu zeigen. Dazu genügt, daß für jedes $\varepsilon > 0$ ein $\delta > 0$ existiert mit der Eigenschaft $|h(x)| < |x - \xi| \cdot \varepsilon$ für alle x mit der Eigenschaft $|x - \xi| < \delta$. Dabei können wir $h(x)$ durch $-h(x)$ ersetzen.

Je nachdem ob $x \geq \xi$ oder $x \leq \xi$ gilt nach Lemma 3.10

$$\pm h(x) = \int_{[a,x]} f(t) dt - \int_{[a,\xi]} f(t) dt - \int_{[\xi,x]} f(\xi) dt = \int_{[\xi,x]} (f(t) - f(\xi)) dt .$$

Aus der Boxungleichung (Abschnitt 3.1) folgt daher

$$|h(x)| \leq |x - \xi| \cdot \sup_{t \in [\xi, x]} |f(t) - f(\xi)| .$$

Die Funktion $f(x)$ ist stetig im Punkt ξ . Es folgt $\sup_{t \in [\xi, x]} |f(t) - f(\xi)| < \varepsilon$ für alle $t \in [\xi, x]$, falls $|x - \xi| < \delta$, $\delta = \delta(\varepsilon)$. Dies zeigt die erste Behauptung. Der verbleibende Zusatz folgt aus dem Mittelwertsatz: Die Ableitung von $F(x) - G(x)$ ist Null auf $[a, b]$. Nach Satz 4.8 ist daher $F(x) - G(x)$ konstant auf $[a, b]$. $\square$

Konvention. Man definiert ganz allgemein für eine stetige Funktion $f: [a, b] \rightarrow \mathbb{R}$ und beliebige $x, y \in [a, b]$ das **orientierte Integral**

$$\int_x^y f(t) dt$$

durch $\int_{[x,y]} f(t)dt$ im Fall $x \leq y$, bzw. durch $-\int_{[y,x]} f(t)dt$ im Fall $y \leq x$. Mit dieser Konvention gilt dann wegen Lemma 3.10 und dem Hauptsatz für jede Stammfunktion G von f auf $[a, b]$ die Formel

$$\boxed{\int_x^y f(t)dt = G(y) - G(x)}.$$

Folgerung. $\log(x) = \int_1^x \frac{dt}{t}$ ist differenzierbar³ auf $\mathbb{R}_{>0}$ mit Ableitung $\frac{1}{x}$.

4.8 Differentialgleichungen

Gegeben sei eine stetige Funktion $a = a(x, y)$ der reellen Variablen $x \in [c, d]$ und $y \in \mathbb{R}^N$

$$a : [c, d] \times \mathbb{R}^N \longrightarrow \mathbb{R}^N.$$

Gesucht ist eine differenzierbare Funktion $f : [c, d] \rightarrow \mathbb{R}^N$ mit der Eigenschaft

$$\boxed{f'(x) = a(x, f(x)) \quad \text{und} \quad f(x_0) = y_0}$$

für ein vorgegebenes $x_0 \in [c, d]$ und vorgegebenen **Anfangswert** $y_0 \in \mathbb{R}^N$. Hierbei bezeichne $f'(x)$ die komponentenweise Ableitung von f nach x .

Satz 4.15 (Picard). Sei $a(x, y)$ ausserdem Lipschitz-stetig auf $\mathbb{R}^N$ in der Variable y mit einer nicht von x abhängigen Lipschitzkonstante M . Dann existiert auf $[c, d]$ eine eindeutig bestimmte Lösung $f(x) \in C^1([c, d], \mathbb{R}^N)$ der Differentialgleichung $f'(x) = a(x, f(x))$ zu jedem gegebenem Anfangswert $f(x_0) = y_0$.

Nach Annahme gilt $\|a(x, y_1) - a(x, y_2)\|_{\mathbb{R}^N} \leq M \cdot \|y_1 - y_2\|_{\mathbb{R}^N}$ für eine Lipschitz-Konstante M , welche nicht (!) von der Variable $x \in [c, d]$ abhängt.

Beweis. Die Differentialgleichung mit Anfangsbedingung ist wegen unserem Hauptsatz 4.14 äquivalent zu einer Integralgleichung:

$$\left\{ \begin{array}{l} f(x) \text{ ist differenzierbar auf } [c, d] \\ f'(x) = a(x, f(x)), f(x_0) = y_0 \end{array} \right\} \iff \left\{ \begin{array}{l} f(x) \text{ ist stetig auf } [c, d] \\ f(x) = y_0 + \int_{x_0}^x a(t, f(t)) dt \end{array} \right.$$

Beweis der Äquivalenz von rechts nach links. Nach Annahme sind $a(t, y)$ und $f(t)$, und daher auch $a(t, f(t))$, stetig. Das vektorwertige Integral $F(x) = \int_{x_0}^x a(t, f(t)) dt$ ist komponentenweise definiert und alle Komponenten sind in der Variable x differenzierbare Funktionen (Hauptsatz) und für den Vektor der Ableitungen gilt $F'(x) = a(x, f(x))$. Aus $f(x) = y_0 + F(x)$ folgt daher durch Ableiten $f'(x) = a(x, f(x))$. Für $x = x_0$ gilt $f(x_0) = y_0$ wegen $F(x_0) = 0$.

³Wegen der Kettenregel gilt daher $\lim_{x \rightarrow 0} \log(1 + yx)/x = \lim_{x \rightarrow 0} \frac{\log(1+yx)-0}{x-0} = \log(1+yx)'|_{x=0} = y$. Anwenden der stetigen Funktion $\exp$ für $x_n = \frac{1}{n}$ liefert die nützliche Formel $\lim_{n \rightarrow \infty} (1 + \frac{y}{n})^n = \exp(y)$.

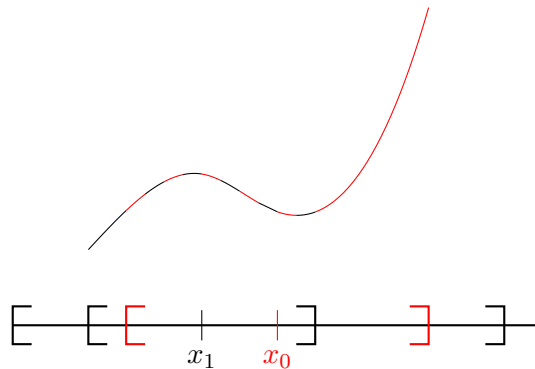
4 Differentiation

Beweis der Äquivalenz von links nach rechts. Ist $f(x)$ differenzierbar, dann ist $f(x)$ stetig nach Lemma 4.6. Aus dem Hauptsatz folgt andererseits

$$f(x) - f(x_0) = \int_{x_0}^x f'(t) dt = \int_{x_0}^x a(t, f(t)) dt.$$

Mit Hilfe der gezeigten Äquivalenz genügt es also die entsprechende Integralgleichung auf $[c, d]$ zu lösen. Wir lösen die Integralgleichung (und damit die Differentialgleichung) dazu zuerst *lokal* auf Teilintervallen genügend kleiner Länge $< \frac{1}{M\sqrt{N}}$.

Verheftung. Angenommen die Lösung der Differential-(oder Integral)gleichung existiert und ist lokal eindeutig auf jedem Teilintervall von $[c, d]$ der Länge $< \frac{1}{M\sqrt{N}}$. Man überdeckt dann das Intervall $[c, b]$ durch überlappende Teilintervalle der Länge $< \frac{1}{M\sqrt{N}}$, wählt Hilfspunkte x_i in den Überlappungen und wendet das lokale Resultat sukzessive für alle Hilfspunkte x_i an. Dies reduziert den allgemeinen Fall auf den Beweis der lokalen Version mit Intervalllänge $< \frac{1}{M\sqrt{N}}$.



Beweis der lokalen Version. Der Raum $X = C([c, d], \mathbb{R}^N)$ aller stetigen $\mathbb{R}^N$ -wertigen Funktionen, versehen mit der Supremums-Norm, ist ein vollständiger metrischer Raum (nach Satz 2.24). Die Selbstabbildung $F : X \rightarrow X$

$$F(f)(x) = y_0 + \int_{x_0}^x a(t, f(t)) dt$$

ist wohldefiniert, denn für stetiges $f : [c, d] \rightarrow \mathbb{R}^N$ ist auch $F(f) : [c, d] \rightarrow \mathbb{R}^N$ stetig und sogar komponentenweise differenzierbar. Die Lösung unserer (lokalen) Integralgleichung ist äquivalent zu der Fixpunktgleichung

$$F(f) = f, \quad f \in X.$$

Unsere Behauptung über die *lokale* Existenz und Eindeutigkeit ergibt sich jetzt sofort aus dem **Banachschen Fixpunktsatz**, denn im Fall $|d - c| < \frac{1}{M\sqrt{N}}$ ist $F : X \rightarrow X$ kontraktiv wegen

$$d_X(F(f), F(g)) = d_X \left(y_0 + \int_{x_0}^x a(t, f(t)) dt, y_0 + \int_{x_0}^x a(t, g(t)) dt \right)$$

$$\stackrel{\text{Def.}}{=} \sup_{x \in [c, d]} \left\| \int_{x_0}^x \left(a(t, f(t)) - a(t, g(t)) \right) dt \right\|_{\mathbb{R}^N}$$

Abschätzen des Integralvektors im $\mathbb{R}^N$ (siehe Seite 13) liefert für $\kappa = |d - c|M\sqrt{N}$

$$\begin{aligned} d_X(F(f), F(g)) &\leq \sqrt{N} \cdot |d - c| \cdot \sup_{t \in [c, d]} \|a(t, f(t)) - a(t, g(t))\|_{\mathbb{R}^N} \\ &\leq \sqrt{N} |d - c| M \cdot \sup_{t \in [c, d]} \|f(t) - g(t)\|_{\mathbb{R}^N} = \kappa \cdot d_X(f, g) \end{aligned}$$

mit Hilfe der Boxungleichung und der Lipschitzstetigkeit von h . Aus $|d - c| < \frac{1}{M \cdot \sqrt{N}}$ folgt die gewünschte Kontraktivität $\kappa < 1$. $\square$

Beispiel 4.16. Ist $a(x, y)$ linear in $y \in \mathbb{R}^N$

$$a(x, y) = A(x) \cdot y + b(x)$$

mit einer $N \times N$ -Matrixfunktion $A(x)$ und einem Vektor $b(x)$, welche stetig von x abhängen, dann sind die Voraussetzungen des Satzes von Picard erfüllt. [Benutze Satz 2.10 und den Beweis von Beispiel 2.2 (4).]

Beispiel 4.17. Um für $\eta_0, \dots, \eta_{n-1} \in \mathbb{R}^N$ allgemeinere Differentialgleichungen vom Typ

$$g^{(n)}(x) = H(x, g(x), \dots, g^{(n-1)}(x))$$

mit der Anfangswertbedingung $g(x_0) = \eta_0, \dots, g^{(n-1)}(x_0) = \eta_{n-1}$ zu behandeln, definiert man die Hilfsfunktion $f: [c, d] \rightarrow \mathbb{R}^{nN}$ durch

$$f(x) = \begin{pmatrix} g(x) \\ g'(x) \\ \vdots \\ g^{(n-1)}(x) \end{pmatrix} \in \mathbb{R}^{nN}$$

und erhält eine äquivalente Differentialgleichung

$$f'(x) = a(x, f(x)) \quad , \quad f(x_0) = y_0 \quad , \quad y_0 = (\eta_0, \dots, \eta_{n-1})$$

wobei $a: [c, d] \times \mathbb{R}^{nN} \rightarrow \mathbb{R}^{nN}$ definiert wird als

$$(x; y_0, \dots, y_{n-1}) \mapsto (y_1, \dots, y_{n-1}, H(x, y_1, \dots, y_{n-1})) .$$

Kombiniert man unsere letzten beiden Beispiele 4.16 und 4.17, erhält man folgende Aussage über lineare Differentialgleichungen auf einem Intervall $[c, d]$, $c < d$.

4 Differentiation

Satz 4.18. Seien $a_0(x), \dots, a_n(x)$ stetige reellwertige Funktionen auf $[c, d]$. Seien $x_0 \in [c, d]$ und $\eta_0, \dots, \eta_{n-1} \in \mathbb{R}$ gegeben. Dann besitzt die **lineare Differentialgleichung**

$$(*) \quad g^{(n)}(x) + a_1(x) \cdot g^{(n-1)}(x) + \dots + a_{n-1}(x) \cdot g'(x) + a_n(x) \cdot g(x) = a_0(x)$$

zu gegebenen Anfangsbedingungen

$$g(x_0) = \eta_0, \dots, g^{(n-1)}(x_0) = \eta_{n-1}$$

eine eindeutige Lösung $g : [c, d] \rightarrow \mathbb{R}$, welche n -mal stetig differenzierbar ist auf $[c, d]$.

Beweis. Die Methode von Beispiel 4.17 führt auf eine vektorwertige lineare Differentialgleichung $f'(x) = A(x) \cdot f(x) + b(x)$ wie in Beispiel 4.16, hier für die Matrixfunktion

$$A(x) = \begin{pmatrix} 0 & 1 & 0 & \dots & 0 & 0 \\ 0 & 0 & 1 & \dots & 0 & 0 \\ 0 & 0 & 0 & \dots & \dots & \dots \\ 0 & 0 & 0 & \dots & 0 & 1 \\ -a_n(x) & -a_{n-1}(x) & \dots & -a_2(x) & -a_1(x) & 0 \end{pmatrix}$$

Die Koordinaten des Vektors $b(x)$ sind hierbei Null bis auf die letzte $b_n(x) := a_0(x)$. $\square$

Ist $b(x) = 0$ oder äquivalent dazu $a_0(x) = 0$, nennt man die obige lineare Differentialgleichung **homogen**. Die Lösungen einer homogen linearen Differentialgleichung wie in Satz 4.18 bilden einen reellen Vektorraum V von Funktionen, wenn man die Forderung von Anfangswertbedingungen weglässt, denn für Lösungen $g(x)$ und $\tilde{g}(x)$ und beliebige reelle Konstanten α, β ist dann auch $\alpha \cdot g(x) + \beta \cdot \tilde{g}(x)$ eine Lösung, wie man sofort sieht.

Satz 4.19. Der Raum V aller Lösungen einer homogenen Differentialgleichung vom Typ $(*)$ ist ein endlich dimensionaler $\mathbb{R}$ -Vektorraum der Dimension n .

Beweis. Für $x_0 \in [c, d]$ ist die $\mathbb{R}$ -lineare Abbildung

$$ev_{x_0} : V \rightarrow \mathbb{R}^n, \quad g \mapsto (g(x_0), \dots, g^{(n-1)}(x_0))$$

injektiv (Eindeutigkeitsaussage von Satz 4.18) und surjektiv (Existenzaussage von Satz 4.18), also ein $\mathbb{R}$ -linearer Isomorphismus von $\mathbb{R}$ -Vektorräumen. $\square$

Beispiel 4.20. Die Differentialgleichung $f'(t) = f(t)$ auf $\mathbb{R}$ mit $f(0) = 1$ hat eine eindeutige Lösung. Die Kettenregel zeigt $\log(f(t))' = \frac{f'(t)}{f(t)} = 1$, also $\log(f(t)) = t + C$ mit $C = 0$ wegen $f(0) = 1$. Es folgt $f(t) = \exp(t)$ und damit⁴

$$\exp(t)' = \exp(t).$$

⁴Für $\alpha \in \mathbb{R}$ folgt $(x^\alpha)' = \alpha \cdot x^{\alpha-1}$ wegen $\exp(\alpha \cdot \log(x))' = \exp(\alpha \cdot \log(x)) \cdot \frac{\alpha}{x}$ sowie $\frac{1}{x} = \exp(-\log(x))$.

Beispiel 4.21 (Sinus und Cosinus). Wir definieren $\sin(x)$ resp. $\cos(x)$ als die eindeutig bestimmten (zweimal stetig differenzierbaren) Funktionen auf $\mathbb{R}$, die im zweidimensionalen $\mathbb{R}$ -Vektorraum V der Lösungen der homogenen Differentialgleichung vom Grad $n = 2$

$$g''(x) + g(x) = 0$$

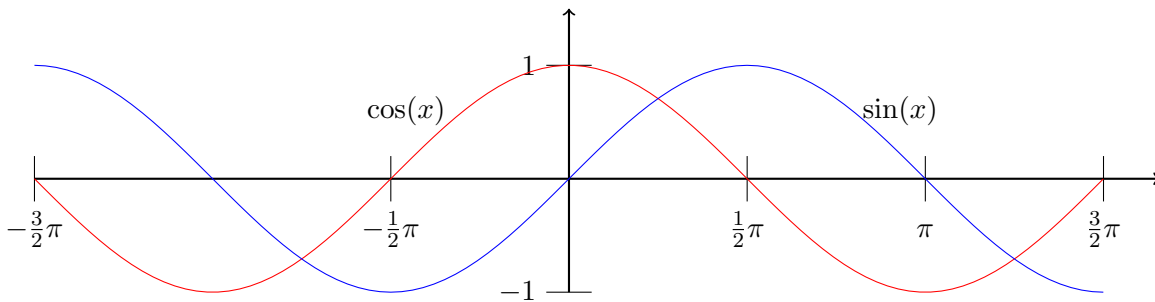
liegen und die Anfangswertbedingungen $g(0) = 0, g'(0) = 1$ resp. $g(0) = 1, g'(0) = 0$ erfüllen. Offensichtlich gilt hier $g \in V \implies g' \in V$. Es folgt $\sin(x)' = \cos(x)$ und $\cos(x)' = -\sin(x)$ und somit (nach Satz 4.8) durch Ableiten

$$\sin(x)^2 + \cos(x)^2 = 1.$$

Für $g \in V$ gilt $g(x) = g(0) \cdot \cos(x) + g'(0) \cdot \sin(x)$. Wegen $g(x) \in V \implies g(x + x_0) \in V$ (Kettenregel!) folgt $\cos(x + x_0) = \cos(x_0)\cos(x) - \sin(x_0)\sin(x)$ sowie $\sin(x + x_0) = \sin(x_0)\cos(x) + \cos(x_0)\sin(x)$. Dies beweist das sogenannte

Lemma 4.22 (Additionstheorem). Für $e(x) := \cos(x) + i \cdot \sin(x)$ in $\mathbb{C}$ und $x, x_0 \in \mathbb{R}$ gilt

$$e(x + x_0) = e(x) \cdot e(x_0).$$



Behauptung. Der Kern K des so definierten Homomorphismus $e : \mathbb{R} \rightarrow \{z \in \mathbb{C} \mid |z| = 1\}$ ist eine Untergruppe der additiven Gruppe von $\mathbb{R}$ und es gilt für eine reelle Zahl $2\pi > 0$

$$K = 2\pi \cdot \mathbb{Z}.$$

Die Funktionen $\sin(x)$ und $\cos(x)$ sind periodisch mit der genauen Periode 2π .

Beweis. Aus Stetigkeitsgründen existiert wegen $\cos(0) = 1$ ein $\delta > 0$ mit $\cos(x) > 0$ für $x \in I = (0, \delta)$. Für $0 \leq x_1 < x_2 \leq \delta$ gilt dann nach dem Mittelwertsatz $\sin(x_2) - \sin(x_1) = (x_2 - x_1) \cdot \cos(\theta) > 0$. Also $K \cap I = \emptyset$. Ist $K \neq \{0\}$, wird daher K von $2\pi := \inf(K \cap \mathbb{R}_{>0})$ als Gruppe erzeugt! Zum Nachweis von $K \neq \{0\}$ genügt ein $x_0 > 0$ mit $\cos(x_0) = 0$. [Dann ist $\sin(x_0) = \pm 1$. Es folgt $e(x_0) = \pm i$ und damit $e(4x_0) = 1$, also $4x_0 \in K$.]

Zur Existenz von x_0 . Wäre $\cos(x) > 0$ für alle $x > 0$, wäre nach dem Mittelwertsatz $\sin(x)$ strikt monoton steigend auf $(0, \infty)$. Also insbesondere wäre $\sin(x) > 0$, nach oben beschränkt durch 1 wegen $\cos(x)^2 + \sin(x)^2 = 1$, und $\cos(x)$ wäre monoton fallend nach unten beschränkt durch 0. Nach Satz 1.27 existiert deshalb der Limes $\zeta = \lim_{n \rightarrow \infty} e(n)$. Die Monotonie des $\sin(x)$ impliziert $\zeta \in \mathbb{C} \setminus \mathbb{R}$. Aus Lemma 4.22 folgt andererseits $\zeta = \zeta \cdot \zeta$, also $\zeta = 0, 1$. Ein Widerspruch! Daher nimmt $\cos(x)$ im Bereich $\mathbb{R}_{>0}$ nicht nur positive Werte an. Also existiert nach dem Zwischenwertsatz eine Nullstelle $x_0 > 0$ von $\cos(x)$. $\square$

4.9 Stetig partiell differenzierbare Funktionen

Sei $U \subseteq \mathbb{R}^n$ eine offene Teilmenge und sei

$$f : U \rightarrow \mathbb{R}^m$$

in $C^1(U, \mathbb{R}^m)$, d.h. eine auf U einmal stetig partiell differenzierbare Funktion.

Lemma 4.23. *Ist $f : U \rightarrow \mathbb{R}^m$ eine einmal stetig partiell differenzierbare Funktion, dann ist f differenzierbar auf U .*

Beweis. Wegen Lemma 4.5 können wir $m = 1$ annehmen. Für festes $\xi \in U$ und alle x nahe genug bei ξ ist die folgende Umformung erklärt

$$\begin{aligned} f(x) - f(\xi) &= [f(x_1, \dots, x_n) - f(\xi_1, x_2, \dots, x_n)] + [f(\xi_1, x_2, \dots, x_n) - f(\xi_1, \xi_2, x_3, \dots, x_n)] \\ &\quad + \dots + [f(\xi_1, \dots, \xi_{n-1}, x_n) - f(\xi_1, \xi_2, \dots, \xi_n)] . \end{aligned}$$

Betrachtet man die Funktion in der i -ten Klammer als Funktion der Variable x_i bei festgehaltenen anderen Variablen, ergibt der eindimensionale Mittelwertsatz 4.8.

$$f(x) - f(\xi) = \sum_{i=1}^n (x_i - \xi_i) \cdot \frac{\partial}{\partial x_i} f(\xi_1, \dots, \xi_{i-1}, \theta_i, x_{i+1}, \dots, x_n)$$

für gewisse θ_i zwischen x_i und ξ_i . Es folgt daher

$$f(x) - f(\xi) - \sum_{i=1}^n (x_i - \xi_i) \cdot \frac{\partial}{\partial x_i} f(\xi) = o(x - \xi) ,$$

denn rechts steht $\sum_i (x_i - \xi_i) \cdot [\frac{\partial}{\partial x_i} f(\xi_1, \dots, \xi_{i-1}, \theta_i, x_{i+1}, \dots, x_n) - \frac{\partial}{\partial x_i} f(\xi)]$ mit $|x_i - \xi_i| \leq \|x - \xi\|$ und der Term in eckigen Klammern ist stetig bei $x = \xi$ mit dem Limes

$$\lim_{x \rightarrow \xi} \left| \frac{\partial}{\partial x_i} f(\xi_1, \dots, \xi_{i-1}, \theta_i, x_{i+1}, \dots, x_n) - \frac{\partial}{\partial x_i} f(\xi) \right| = 0 ,$$

weil nach Annahme die partiellen Ableitungen $\frac{\partial}{\partial x_i} f$ stetig im Punkt ξ sind. Beachte $x \rightarrow \xi$ impliziert $\theta_i \rightarrow \xi_i$. Also ist f differenzierbar im Punkt ξ . $\square$

Lemma 4.24 (Kritische Punkte). *Sei $U \subset \mathbb{R}^n$ offen und $f : U \rightarrow \mathbb{R}^m$ einmal stetig partiell differenzierbar auf U . Gilt $Df(\xi) = 0$ für ein $\xi \in U$, dann existiert für jedes $\varepsilon > 0$ ein $\delta > 0$ so daß (für die Euklidische Norm oder die Quaternorm des $\mathbb{R}^n$) gilt*

$$\|x - \xi\| < \delta, \|y - \xi\| < \delta \implies \|f(x) - f(y)\| < \varepsilon \cdot \|x - y\| .$$

Beweis. ObdA ist $m = 1$. Wähle $\delta > 0$ so klein, daß die Kugel B vom Radius δ um ξ ganz in U enthalten ist. Aus dem Mittelwertsatz 4.8 folgt dann

$$|f(x) - f(y)| = \|Df(\theta)(x - y)\| \leq n \sup_i \left| \frac{\partial}{\partial x_i} f(\theta) \right| \cdot \|x - y\|$$

für alle $x, y \in B$, da die Verbindungsgerade zwischen x und y auch in B liegt. Aus der Annahme $\frac{\partial}{\partial x_1} f(\xi) = \dots = \frac{\partial}{\partial x_n} f(\xi) = 0$ und der Stetigkeit der partiellen Ableitungen im Punkt ξ folgt $n \left| \frac{\partial}{\partial x_i} f(\theta) \right| < \varepsilon$ für alle θ mit $\|\theta - \xi\| < \delta$, wenn $\delta > 0$ genügend klein gewählt wird. $\square$

4.10 Der Umkehrsatz

Sei U offen in $\mathbb{R}^n$ und sei

$$f: U \rightarrow \mathbb{R}^n$$

eine einmal stetig partiell differenzierbare, also insbesondere differenzierbare Funktion auf U . Wegen $n = m$ ist in jedem Punkt $\xi \in U$ die Jacobimatrix eine $n \times n$ -Matrix. Der folgende Satz zeigt, dass die Invertierbarkeit der Jacobimatrix $Df(\xi)$ im Punkt ξ eine hinreichende Bedingung für die Existenz einer lokalen Umkehrfunktion von f in der Nähe von ξ resp. $f(\xi)$ ist:

Satz 4.25. Für $\xi_0 \in U$ mit invertierbarem $Df(\xi_0)$ gibt es eine offene Teilmenge V von U , welche ξ_0 enthält, für die f eingeschränkt auf V eine bijektive Abbildung von V auf $W = f(V)$ definiert, so daß gilt

$$W = f(V) \subseteq \mathbb{R}^n \text{ ist } \underline{\text{offen}}.$$

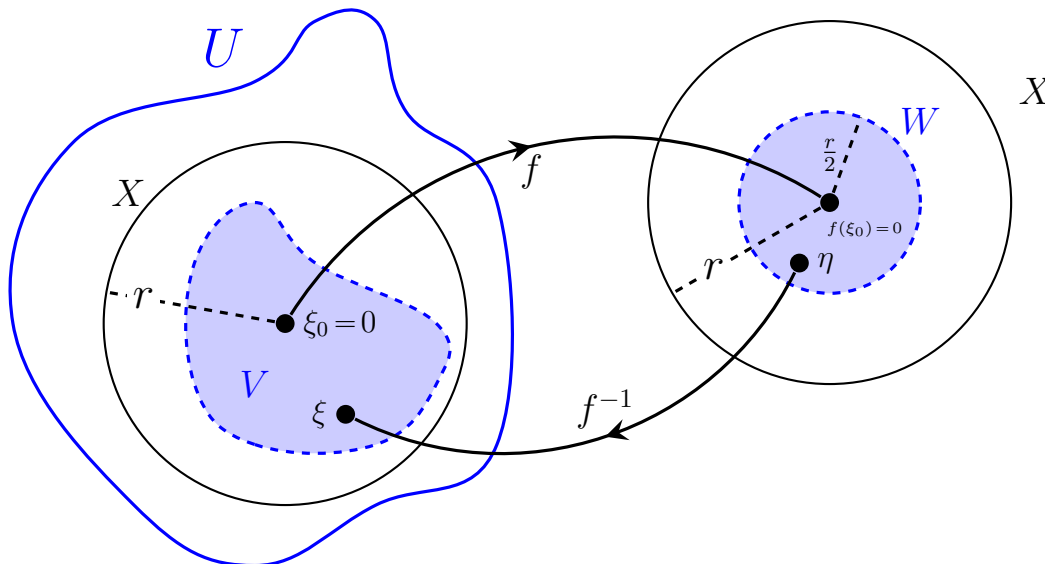
Weiterhin: Die lokale Umkehrfunktion

$$f^{-1}: W \rightarrow V$$

ist einmal stetig partiell differenzierbar, und damit auch differenzierbar auf $W = f(V)$.

Bemerkung. Die Invertierbarkeit von $Df(\xi)$ ist auch eine notwendige Bedingung für die Existenz einer differenzierbaren lokalen Umkehrfunktion $g = f^{-1}$ in der Nähe von ξ . Dies folgt aus der Kettenregel, denn $g \circ f = id$ impliziert $Dg(\eta) \circ Df(\xi) = D(id)(\xi) = id$ für $\eta = f(\xi)$. Somit ist $Dg(\eta)$ als Matrix zu $Df(\xi)$ invers.

Beweis. Wir geben zuerst den Beweis im Spezialfall $\xi_0 = f(\xi_0) = 0$ und $Df(\xi_0) = id$.



4 Differentiation

Die Hilfsfunktion $F = F_\eta$:

$$F(x) = x - f(x) + \eta$$

hat für gegebenes konstantes $\eta \in \mathbb{R}^n$ verschwindende Ableitung im Punkt ξ . Für $\varepsilon = 1/2$ gilt somit nach Lemma 4.24 für alle x, y vom Abstand zu $\xi_0 = 0$ kleiner als $\delta = \delta(1/2) > 0$

$$\|F(x) - F(y)\| < \frac{1}{2} \|x - y\|.$$

Der vollständige Raum X . Wir wählen eine abgeschlossene Kugel⁵ X um $\xi_0 = 0$ vom Radius r für ein $0 < r < \delta$. Dann ist $F: X \rightarrow \mathbb{R}^n$ Lipschitz-stetig mit Lipschitzkonstante $\frac{1}{2}$.

Bedingung an η . $F(0) = \eta$ und $\|F(x)\| \leq \|F(x) - F(0)\| + \|F(0)\| \leq \frac{1}{2} \|x\| + \|\eta\|$ implizieren für $x \in X$ (d.h. $\|x\| \leq r$) sowie gleichzeitig für

$$\|\eta\| < r/2$$

die Ungleichung $\|F(x)\| \leq r$. Es folgt

$$F: X \rightarrow X.$$

Fixpunktsatz. Der Fixpunktsatz von Banach liefert daher einen eindeutig bestimmten Fixpunkt $\xi \in X$ der kontraktiven Abbildung $F: X \rightarrow X$. Beachte $F(\xi) = \xi$ ist äquivalent zu $f(\xi) = \eta$. Mit anderen Worten $\xi = f^{-1}(\eta)$:

$$\exists! \xi \in X \text{ mit } f(\xi) = \eta, \text{ falls } \|\eta\| < r/2.$$

Konstruktion von V . Sei $W = B_{r/2}(0)$ die offene Kugel um Null aller $\eta \in \mathbb{R}^n$ mit $\|\eta\| < r/2$. Da $W \subseteq \mathbb{R}^n$ offen und f stetig ist, ist das Urbild $f^{-1}(W)$ offen in $\mathbb{R}^n$ (benutze Lemma 4.6 und Satz 2.12). Für $V = f^{-1}(W) \cap X$ gilt dann wie bereits gezeigt

$$f: V \xrightarrow{\cong} f(V) = W.$$

Kontroll-Abschätzungen. Die Kontraktivität $\|F(x) - F(y)\| < \frac{1}{2} \|x - y\|$ von F auf X liefert für $x, y \in X$ mit Hilfe der unteren und oberen Dreiecksungleichung⁶ (s. Seite 6)

$$\frac{1}{2} \|x - y\| < \|f(x) - f(y)\| < \frac{3}{2} \|x - y\|.$$

V ist offen. $\xi \in V \implies \eta = f(\xi) \in W$. Aus der unteren Kontroll-Abschätzung von f folgt für $x = \xi, y = 0$ dann $\frac{1}{2} \|\xi\| < \|\eta\|$. Andererseits $\|\eta\| < \frac{r}{2}$. Somit $\|\xi\| < r$. Also liegt ξ bereits in der

⁵Dies ist eine abgeschlossene Teilmenge des $\mathbb{R}^n$, also versehen mit der Euklidischen Metrik nach 2.8 ein vollständiger metrischer Raum.

⁶Benutze $\|u\| - \|v\| \leq \|u + v\| \leq \|u\| + \|v\|$ für $u = x - y$ und $v = F(y) - F(x)$ und $u + v = f(y) - f(x)$. Beachte $\|v\| < \frac{1}{2} \|u\|$ wegen der Kontraktivität von F .

offenen Kugel $X^0 \subset X$ vom Radius r um Null. Daher ist $V = f^{-1}(W) \cap X = f^{-1}(W) \cap X^0$ als Durchschnitt zweier offener Mengen selbst offen.

Stetigkeit von $f^{-1}: W \rightarrow V$. Eine unmittelbare Konsequenz von

$$\frac{1}{2} \|f^{-1}(\eta_1) - f^{-1}(\eta_2)\| < \|\eta_1 - \eta_2\|,$$

und dies gilt für alle $\eta_1, \eta_2 \in W$ wegen der linken Kontrollabschätzung.

Differenzierbarkeit von f^{-1} im Punkt 0. Für $\eta \neq 0 \iff f^{-1}(\eta) = \xi \neq 0$ gilt

$$\begin{aligned} \frac{\|f^{-1}(\eta) - f^{-1}(0) - id(\eta - 0)\|}{\|\eta - 0\|} &= \frac{\|f^{-1}(\eta) - \eta\|}{\|\eta\|} = \frac{\|\xi - f(\xi)\|}{\|f(\xi)\|} \\ &= \frac{\|\xi\|}{\|f(\xi)\|} \cdot \frac{\|f(\xi) - f(0) - id(\xi - 0)\|}{\|\xi - 0\|} \end{aligned}$$

Da f nach 4.6 auf V stetig ist, und f^{-1} stetig auf W ist, sind $\xi \rightarrow 0$ und $\eta = f(\xi) \rightarrow 0$ zueinander äquivalent. Da rechts der Limes $\xi \rightarrow 0$ existiert und Null ist (f ist differenzierbar im Punkt $\xi = 0$ mit der Ableitung id , und der Faktor $\|\xi\|/\|f(\xi)\|$ kann durch 2 abgeschätzt werden wegen der Kontrollabschätzungen), existiert der Limes $\eta \rightarrow 0$ links und ist auch Null. Somit ist f^{-1} differenzierbar im Punkt $\eta = 0$ mit der Ableitung $Df^{-1}(0) = id$.

Differenzierbarkeit von f^{-1} auf W . Hierzu nehmen wir an, dass der Radius r obdA so klein gewählt wurde, dass für alle $\xi \in X$ und damit für alle $\xi \in V$ die Ableitung von f im Punkt ξ invertierbar ist. Dann zeigt unser vorheriges Argument die Differenzierbarkeit von f^{-1} in allen Punkten $\eta = f(\xi) \in W$. Beachte die nachfolgende Reduktion auf den Spezialfall.

Stetig partielle Differenzierbarkeit von f^{-1} auf W . Wegen der Kettenregel ist $Df^{-1}(\eta)$ die zu $Df(\xi)$ inverse Matrix ist (wir unterscheiden jetzt nicht zwischen Df und Jf etc.). Die Cramersche Regel oder der Laplace Entwicklungssatz liefert daher die Formel

$$Df^{-1}(\eta) = (Df(\xi))^{-1} = \frac{Df(\xi)^{ad}}{\det(Df(\xi))}.$$

Beachte, ξ hängt stetig von η ab, da f^{-1} stetig ist. Die Einträge der adjungierten Matrix und die Determinante von $Df(\xi)$ sind Polynome in den Matrixkoeffizienten von $Df(\xi)$, also stetig in ξ , da f stetig partiell differenzierbar ist. Andererseits sind die partiellen Ableitungen von f^{-1} die Koeffizienten der Jacobimatrix $Df^{-1}(\eta)$, also wie oben erklärt stetige Funktionen von $\eta \in W$.

Reduktion auf den Spezialfall. Die zum Beweis des Umkehrsatzes gemachten Annahmen

$$\xi_0 = 0 \text{ und } \eta_0 = f(\xi_0) = 0 \text{ und } Df(\xi_0) = id$$

sind unbedenklich. Dazu modifiziert man ein allgemeines f mit affin linearen Abbildungen der Gestalt $\varphi(x) = L(x) + \xi_0$, $L \in Gl(n, \mathbb{R})$ resp. $\psi(x) = x - \eta_0$. Solche Abbildungen haben Jacobimatrix L resp. id und sind invertierbar auf ganz $\mathbb{R}^n$, und ihre Umkehrabbildungen sind wieder affin linear. Hat die Hilfsfunktion

$$\tilde{f} = \psi \circ f \circ \varphi$$

4 Differentiation

eine lokale Umkehrfunktion bei $x = 0$, dann besitzt unsere Funktion f wie behauptet eine lokale Umkehrfunktion⁷ bei $x = \xi_0$, nämlich

$$f^{-1} = \varphi \circ \tilde{f}^{-1} \circ \psi.$$

Andererseits⁸ gilt $\tilde{f}(0) = 0$ und $D\tilde{f}(0) = id$, falls L geeignet gewählt wird, nämlich

$$L = Df(\xi_0)^{-1}.$$

Genau an dieser Stelle geht die Invertierbarkeit der Jacobimatrix $Df(\xi_0)$ ein! □

4.11 Substitutionsregel

Sei $U \subset \mathbb{R}^n$ eine offene Menge und

$$f: U \rightarrow \mathbb{R}$$

eine stetige Funktion auf U . Wir nehmen an, f habe kompakten Träger in U . D.h. es gibt eine kompakte Teilmenge $K \subset U$ so daß $f(x)$ ausserhalb von K Null ist. Es bezeichne $C_c(U) \subseteq C(U)$ den Raum aller stetigen Funktionen mit kompaktem Träger in U . Funktionen in $C_c(U)$ können durch Null zu stetigen Funktionen auf $\mathbb{R}^n$ fortgesetzt werden. Dadurch ändert sich der Träger nicht. In diesem Sinn kann man $C_c(U)$ als Unterraum von $C_c(\mathbb{R}^n)$ auffassen. Einschränkung des Euklidischen Standardintegrals $\int_{\mathbb{R}^n}: C_c(\mathbb{R}^n) \rightarrow \mathbb{R}$ auf den Teilraum $C_c(U)$ definiert das abstrakte Integral

$$I(f) = \int_U f(x) dx \quad , \quad f \in C_c(U)$$

auf dem Verband $C_c(U)$.

Definition 4.26. *Ein Koordinatenwechsel ist eine bijektive und einmal stetig partiell differenzierbare Abbildung zwischen offenen Mengen U, V im $\mathbb{R}^n$*

$$\varphi: U \rightarrow V$$

für die gilt:

$$\det D\varphi(x) \neq 0 \text{ für alle } x \in U.$$

Man nennt φ einen **orientierten Koordinatenwechsel**, wenn $\det D\varphi(x) > 0$ gilt für alle $x \in U$.

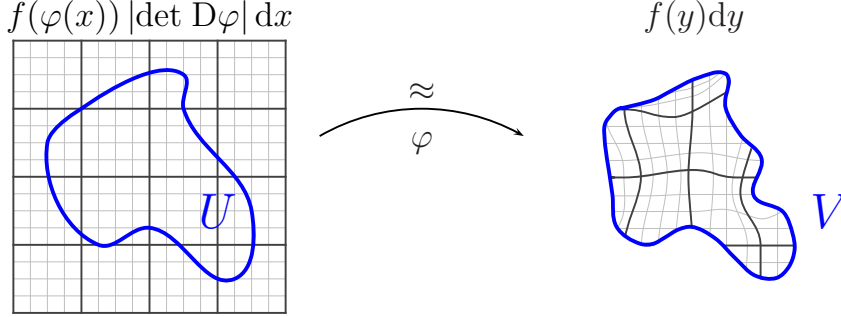
Aus dem Satz von der Umkehrfunktion folgt: Ist $\varphi: U \rightarrow V$ ein Koordinatenwechsel, dann ist auch $\psi = \varphi^{-1}: V \rightarrow U$ ein Koordinatenwechsel.

⁷ $f \circ (\varphi \circ \tilde{f}^{-1} \circ \psi) = \psi^{-1} \circ \tilde{f} \circ \tilde{f}^{-1} \circ \psi = id$ und $(\varphi \circ \tilde{f}^{-1} \circ \psi) \circ f = \varphi \circ \tilde{f}^{-1} \circ \tilde{f} \circ \varphi^{-1} = id$.

⁸Für die Funktion $\tilde{f}$ folgt die Existenz der lokalen Umkehrfunktion bei $x = 0$ aus dem vorherigen Abschnitt.

Satz 4.27 (Substitutionsregel). Sei φ ein Koordinatenwechsel. Dann liegt für jede Funktion $f: V \rightarrow \mathbb{R}$ aus $C_c(V)$ die Funktion $f(\varphi(x)) \cdot |\det D\varphi(x)|$ in $C_c(U)$, und es gilt

$$(*) \quad \boxed{\int_V f(y) dy = \int_U f(\varphi(x)) \cdot |\det D\varphi(x)| dx}.$$



Bemerkung. Analog wie im Beweis von Lemma 3.9 folgt daraus die analoge Substitutionsregel für alle Funktionen $f: V \rightarrow \mathbb{R}$, welche *Lebesgue integrierbar* sind im Sinne von Kapitel 6.

Beweis. 1.Schritt. Die Zerlegung $f = f^+ - f^-$ mit $f^+ = \max(f, 0)$ und $f^- = -\min(f, 0)$ und die $\mathbb{R}$ -Linearität des Integrals erlaubt es auf den Fall $f \geq 0$ zu reduzieren. Sei also oBdA $f \geq 0$, und damit auch $g(x) = f(\varphi(x)) |\det D\varphi(x)| \geq 0$.

2.Schritt. Es genügt für alle Koordinatenwechsel $\varphi: U \rightarrow V$ und alle $f \in C_c(V)$ mit $f \geq 0$ zu zeigen

$$(**) \quad \boxed{\int_V f(y) dy \leq \int_U f(\varphi(x)) \cdot |\det D\varphi(x)| dx},$$

denn angewandt auf $\psi: V \rightarrow U$ mit $\psi = \varphi^{-1}$ und $g(x) = f(\varphi(x)) |\det D\varphi(x)|$ gibt (**) die Ungleichung $\int_U g(x) dx \leq \int_V g(\psi(y)) |\det D\psi(y)| dy$. Nun ist $g(\psi(y)) |\det D\psi(y)|$ gleich $f(\varphi(\psi(y))) |\det D\varphi(\psi(y))| |\det D\psi(y)| = f(y) |\det D(\varphi \circ \psi)(y)| = f(y)$. Aus (**) folgt damit

$$\int_U f(\varphi(x)) \cdot |\det D\varphi(x)| dx \leq \int_V f(y) dy,$$

und damit, durch beide Abschätzungen zusammen, die Substitutionsregel (*).

3.Schritt. Gilt (**) für $\varphi: U \rightarrow V$ und $\psi: V \rightarrow W$, dann gilt (**) für $\psi \circ \varphi: U \rightarrow W$. [(**) für die Substitutionen $y = \varphi(x)$ und $z = \psi(y)$ liefert $\int_W g(z) dz \leq \int_V g(\psi(y)) \cdot |\det D\psi(y)| dy \leq \int_U g(\psi(\varphi(x))) \cdot |\det D(\psi(\varphi(x)))| \cdot |\det D\varphi(x)| dx = \int_U g((\psi \circ \varphi)(x)) |\det D(\psi \circ \varphi)(x)| dx$ vermöge der Kettenregel und der Produktformel für Determinanten. Ditto für (*).]

4.Schritt. Um die Aussage (*) - oder äquivalent (**) - für die Einschränkung $\varphi: U \rightarrow V$ von linearen Abbildungen $\varphi: \mathbb{R}^n \rightarrow \mathbb{R}^n$ zu zeigen, kann man sich wegen Schritt 3 auf den

4 Differentiation

Fall von Elementarmatrizen zurückziehen [Jede invertierbare Matrix ist ein Produkt von Diagonalmatrizen und oberen und unteren elementaren Dreiecksmatrizen (Scherungen).] Den Fall von Diagonalmatrizen behandelt man wie in Lemma 3.9. Der Fall einer elementaren Scherung ist ein einfacher Spezialfall des späteren Satzes von Fubini⁹, folgt aber auch aus einer simplen Modifikation des Arguments in Schritt 8. Wir können daher jetzt annehmen, im Fall von linearen Koordinatenwechseln sei (*) bereits gezeigt. Ditto für Translationen.

5.Schritt. Also angenommen es gäbe einen Fall, wo die Ungleichung (**) tatsächlich falsch wäre, also die linke Seite in (**) etwa um $\kappa > 0$ grösser wäre als die rechte. Wir legen dann den Träger K von $f(\varphi(x))$ in U in einen Quader $Q = Q_0 \subseteq \mathbb{R}^n$, sagen wir mit Seitenlänge c , und halbieren alle Seitenlängen sukzessive (Quaderschachtelung). Für jeden der iterierten Teilquader Q_m ist die Funktion $\chi_{\varphi(Q_m)}(y)f(y)$ integrierbar bezüglich eines erweiterten Integrals I^- (siehe Beispiel 2.25 und 3.19 zusammen mit Satz 3.20). Dann zeigt¹⁰¹¹ man leicht durch einen

Schubfachschluss(!): Es gibt eine absteigende Folge von Teilquadern $Q_m \subseteq Q_0$ mit $\text{vol}(Q_m) = 2^{-mn}\text{vol}(Q_0)$ sowie

$$\lim_{m \rightarrow \infty} \frac{\int \chi_{\varphi(Q_m)}(y)f(y)dy}{\text{vol}(Q_m)} \geq \frac{\kappa}{\text{vol}(Q_0)} + \lim_{m \rightarrow \infty} \frac{\int_{Q_m} f(\varphi(x)) \cdot |\det D\varphi(x)|dx}{\text{vol}(Q_m)}$$

und $\bigcap Q_m = \{x_0\}$ (Quaderschachtelung; siehe Übungsaufgaben).

Hinweis. Lasse den Limes weg und multipliziere mit $\text{vol}(Q_m)$. Wie findet man dann wohl den Quader Q_m in Q_{m-1} ? Natürlich wie folgt: Q_m ist einer der Teilquader mit maximaler Abweichung in der Ungleichung (**)!

6.Schritt. Wegen der Stetigkeit von f existiert der Limes¹²

$$\lim_{m \rightarrow \infty} \frac{\int_{Q_m} f(\varphi(x)) \cdot |\det D\varphi(x)|dx}{\text{vol}(Q_m)} = f(\varphi(x_0)) \cdot |\det D\varphi(x_0)|.$$

7.Schritt. Durch Komposition mit einer linearen Abbildung kann wegen Schritt 2), 3) und 4) weiterhin obdA $D\varphi(x_0) = id_{\mathbb{R}^n}$ angenommen werden. Dann gilt für $y_0 := \varphi(x_0)$

$$\lim_{m \rightarrow \infty} \frac{\int_{Q_m} f(\varphi(x)) \cdot |\det D\varphi(x)|dx}{\text{vol}(Q_m)} = f(y_0).$$

OBdA sei ausserdem $x_0 = 0$ und $y_0 = 0$.

⁹Im Scherungsfall ist obdA $n = 2$ und die Aussage folgt aus $\int_{\mathbb{R}} (\int_{\mathbb{R}} f(x, \lambda \cdot x + y)dy)dx = \int_{\mathbb{R}} (\int_{\mathbb{R}} f(x, y)dy)dx$ wegen der Translationsinvarianz $\int_{\mathbb{R}} h(y_0 + y)dy = \int_{\mathbb{R}} h(y)dy$ des Integrals.

¹⁰Schubfachschluss: Gilt $\kappa \geq \lambda_{Q_{m-1}}/\text{vol}(Q_{m-1}) = \sum_{\nu=1}^{2^n} \lambda_{Q_{m,\nu}}/\text{vol}(Q_{m-1})$, und ist Q_m einer der 2^n Teilquader $Q_{m,\nu}$ mit maximalem $\lambda_{Q_{m,\nu}}$, dann gilt $\kappa \geq 2^n \cdot \lambda_{Q_m}/\text{vol}(Q_{m-1}) = \lambda_{Q_m}/\text{vol}(Q_m)$.

¹¹5.Schritt: Eigentlich müsste als Integrationsbereich dastehen $Q_m \cap U$ respektive $\varphi(Q_m \cap U)$. Aber für $m \gg 0$ gilt $Q_m \subseteq U$ wegen $x_0 \in K \subseteq U$.

¹²6.Schritt: Es gilt $\text{vol}(Q_m) \cdot \min_{x \in Q_m} h(x) \leq \int_{Q_m} h(x)dx \leq \text{vol}(Q_m) \cdot \max_{x \in Q_m} h(x)$, und somit $\lim_{m \rightarrow \infty} \text{vol}(Q_m)^{-1} \int_{Q_m} h(x)dx = h(x_0)$ für jede stetige Funktion h .

8.Schritt. Sei $\varepsilon > 0$ beliebig. Wegen Schritt 7 ist $D\varphi(x_0) = id_{\mathbb{R}^n}$, also $D(\varphi - id_{\mathbb{R}^n})(x_0) = 0$. Es folgt $\|\varphi(x) - x\| < \varepsilon\|x\|$ für $\|x\| < \delta(\varepsilon)$ nach Lemma 4.24. Ist daher m groß genug, gilt wegen dieser Abschätzung¹³:

- $\varphi(Q_m)$ im Quader $Q'_m = (1+\varepsilon)Q_m$ der Seitenlänge $(1+\varepsilon)\frac{c}{2^m}$ enthalten. Zur Erinnerung: $\frac{c}{2^m}$ war die Seitenlänge von Q_m . Also $vol(Q'_m) = (1+\varepsilon)^n vol(Q_m)$.

Schritt 1 impliziert $\chi_{Q'_m}(y)f(y) \geq \chi_{\varphi(Q_m)}(y)f(y)$ und damit ist wegen der Monotonie des abstrakten Integrals I^- auf $C_c(\mathbb{R}^n)^-$

$$(1+\varepsilon)^n \cdot \frac{\int \chi_{Q'_m}(y)f(y)dy}{vol(Q'_m)} = \frac{\int \chi_{Q'_m}(y)f(y)dy}{vol(Q_m)} \geq \frac{\int \chi_{\varphi(Q_m)}(y)f(y)dy}{vol(Q_m)}.$$

Schritt 9. Die Stetigkeit von f und $y_0 \in \varphi(Q_m) \subseteq Q'_m$ liefert wie in Schritt 6

$$\lim_{m \rightarrow \infty} \frac{\int \chi_{Q'_m}(y)f(y)dy}{vol(Q'_m)} = f(y_0).$$

Zusammen mit Schritt 5), 6), 7) und 8) folgt daraus

$$(1+\varepsilon)^n \cdot f(y_0) \geq \frac{\kappa}{vol(Q_0)} + f(y_0),$$

oder $(1+\varepsilon)^n \cdot f(y_0) - f(y_0) \geq \frac{\kappa}{vol(Q_0)} > 0$ wegen der Schritt 5, 6, 7 und 8. Wählt man $\varepsilon > 0$ genügend klein, wird die linke Seite $f(y_0) \cdot [(1+\varepsilon)^n - 1]$ kleiner als jede feste positive Zahl im Widerspruch zu $\frac{\kappa}{vol(Q_0)} > 0$ von Schritt 5. Dies zeigt (***) und damit die Behauptung (*). $\square$

4.12 Differentialformen

Sei $U \subseteq \mathbb{R}^n$ zulässig und $C^\infty(U)$ der Raum aller unendlich oft partiell differenzierbaren reellwertigen Funktionen auf U . Wir betrachten Teilmengen $I \subseteq \{1, \dots, n\}$ der festen Kardinalität $|I| = i$. Einen formalen Ausdruck der Gestalt

$$\omega(x) = \sum_{I, |I|=i} \omega_I(x) \cdot dx_I,$$

dessen Koeffizienten Funktionen $\omega_I(x) \in C^\infty(U)$ sind, nennt man ω eine i -Form auf U . Man nennt i den **Grad** von ω . Den $\mathbb{R}$ -Vektorraum aller i -**Formen** auf U (mit komponentenweiser Addition und Skalarmultiplikation der Koeffizienten) bezeichnen wir mit

$$A^i(U).$$

Schreibweise. Sei $I = \{n_1, \dots, n_i\}$ mit $n_1 < \dots < n_i$, dann schreiben wir symbolisch $dx_I = dx_{n_1} \wedge \dots \wedge dx_{n_i}$ sowie $dx_\emptyset = 1$. Für die einelementigen Teilmengen $I = \{i\}$ schreiben wir meistens dx_i anstelle von $dx_{\{i\}}$. Wir erhalten damit für $A^\bullet(U) := \bigoplus_{i=0}^n A^i(U)$

¹³8.Schritt: Die hier benutzte Norm sei obdA die Norm $\|x\| = \max_{i=1, \dots, n} |x_i|$ und obdA $x_0 = y_0 = 0$. Dann gilt $(1-\varepsilon)x_i < \varphi_i(x) < (1+\varepsilon)x_i$, $\forall i = 1, \dots, n$. Daraus folgt $\varphi(Q_m) \subset Q'_m$ für einen Quader Q'_m wie behauptet.

4 Differentiation

- $A^0(U) = C^\infty(U)$
- $A^1(U) = C^\infty(U) \cdot dx_1 \oplus \dots \oplus C^\infty(U) \cdot dx_n$
- $A^2(U) = C^\infty(U) \cdot dx_1 \wedge dx_2 \oplus \dots \oplus C^\infty(U) \cdot dx_{n-1} \wedge dx_n$
- $\dots$
- $A^n(U) = C^\infty(U) \cdot \omega_n$ für $\omega_n := dx_1 \wedge \dots \wedge dx_n$.

Das $\wedge$ -Produkt. Wir definieren $dx_I \wedge dx_J := 0$, falls $I \cap J \neq \emptyset$. Anderenfalls setzen wir $dx_I \wedge dx_J = \text{sign}(\sigma) dx_{I \cup J}$, wobei σ die Permutation ist, welche $n_1, \dots, n_i, m_1, \dots, m_j$ in eine aufsteigende Reihenfolge bringt. Hierbei seien $n_1 < \dots < n_i$ und $m_1 < \dots < m_j$ so gewählt, dass $dx_I = dx_{n_1} \wedge \dots \wedge dx_{n_i}$ und $dx_J = dx_{m_1} \wedge \dots \wedge dx_{m_j}$ gilt. Wir erhalten eine wohldefinierte $\mathbb{R}$ -bilineare Abbildung¹⁴

$$A^i(U) \times A^j(U) \xrightarrow{\wedge} A^{i+j}(U)$$

$$\left(\sum_I \omega_I(x) \cdot dx_I, \sum_J \omega_J(x) \cdot dx_J \right) \mapsto \sum_I \sum_J \omega_I(x) \omega_J(x) \cdot dx_I \wedge dx_J.$$

Das $\wedge$ -Produkt ist per Definition distributiv.

Beispiel. Aus der Definition des $\wedge$ -Produkts folgt unmittelbar

- $dx_i \wedge dx_i = 0$
- $dx_i \wedge dx_j = -dx_j \wedge dx_i$

Allgemeiner folgt aus der Definition: $dx_I \wedge dx_J = (-1)^{|I||J|} dx_J \wedge dx_I$. Also für beliebige $\eta \in A^i(U)$ und $\omega \in A^j(U)$

$$\boxed{\eta \wedge \omega = (-1)^{ij} \cdot \omega \wedge \eta}.$$

Die Cartanableitungen. Wir definieren die Cartan Ableitungen d und damit den sogenannten **Differentialformenkomplex**

$$A^0(U) \xrightarrow{d} A^1(U) \xrightarrow{d} \dots \xrightarrow{d} A^{n-1}(U) \xrightarrow{d} A^n(U) \xrightarrow{d} A^{n+1}(U) = 0$$

wobei die Cartan Ableitungen definiert sind durch

$$d\left(\sum_I \omega_I(x) \cdot dx_I\right) = \sum_I \sum_{i=1}^n \frac{\partial \omega_I}{\partial x_i}(x) \cdot dx_i \wedge dx_I.$$

Beispiel. Für eine Funktion $f(x) \in A^0(U) = C^\infty(U)$ bedeutet dies

$$df(x) = \sum_{i=1}^n \frac{\partial f}{\partial x_i}(x) \cdot dx_i.$$

¹⁴Wir zeigen später, dass das $\wedge$ -Produkt assoziativ ist im Sinne von $(dx_I \wedge dx_J) \wedge dx_K = dx_I \wedge (dx_J \wedge dx_K)$. Da es auch offensichtlich distributiv ist, wird dadurch $A^\bullet(U)$ zu einem (nichtkommutativen) Ring.

Man nennt dann $df \in A^1(U)$ das **totale Differential** von f (im Prinzip ist es dasselbe wie die Jacobimatrix von f , nur etwas anders geschrieben). Die Abbildung

$$d = d_i : A^{i-1}(U) \rightarrow A^i(U)$$

nennt man im Fall $i = 1$ den **Gradient grad** und im Fall $i = n$ die **Divergenz div**. Im Fall $n = 3, i = 2$ (der klassischen Vektoranalysis auf $\mathbb{R}^3$) benutzt man die Bezeichnung **Rotation**.

Spezialfall. Sei $f(x) = x_i$ die i -te Koordinatenfunktion, mit anderen Worten die Zusammensetzung $U \hookrightarrow \mathbb{R}^n \xrightarrow{pr_i} \mathbb{R}$. Dann gilt

$$df = 1 \cdot dx_i = dx_i,$$

oder kurz $d(x_i) = dx_i$. Dies (!) rechtfertigt erneut die Schreibweise dx_i .

Die Produktformel (Superderivationseigenschaft). Für $\eta \in A^i(U)$ und $\omega \in A^j(U)$ gilt

$$\boxed{d(\eta \wedge \omega) = d\eta \wedge \omega + (-1)^i \eta \wedge d\omega}.$$

Beweis. Wegen der Bilinearität des $\wedge$ -Produkts können wir obdA annehmen $\eta = f(x) \cdot dx_I$ und $\omega = g(x) \cdot dx_J$ für $f, g \in C^\infty(U)$. Dann gilt

$$d(\eta \wedge \omega) = d(fg \cdot dx_I \wedge dx_J) = d(fg) \wedge dx_I \wedge dx_J$$

nach der Definition der Cartanableitung. Die übliche Produktformel für die partiellen Ableitungen einer Funktion liefert

$$d(fg) = gdf + fdg.$$

Also $d(\eta \wedge \omega) = (gd(f) + fd(g)) \wedge (dx_I \wedge dx_J) = d(f)dx_I \wedge gdx_J + (-1)^i f dx_I \wedge (d(g) \wedge dx_J) = d\eta \wedge \omega + (-1)^i \eta \wedge d\omega$. Hierbei wurde benutzt $dx_i \wedge (dx_I \wedge dx_J) = (dx_i \wedge dx_I) \wedge dx_J = (-1)^i (dx_I \wedge dx_i) \wedge dx_J = (-1)^i dx_I \wedge (dx_i \wedge dx_J)$ vermöge des Assoziativgesetzes, welches hier als Übungsaufgabe verbleibt¹⁵. $\square$

Integration. Sei $\omega = f(x) \cdot dx_1 \wedge \cdots \wedge dx_n$ eine Form **höchsten Grades** mit kompaktem Träger, d.h. $f \in C_c^\infty(U)$ oder kurz $\omega \in A_c^n(U) := C_c^\infty(U) \cdot \omega_n$. Dann wird per Definition das Integral $\int_U \omega$ erklärt durch das n -dimensionale Standardintegral der Funktion $f(x)$

$$\int_U \omega := \int_U f(x) dx_1 dx_2 \cdots dx_n.$$

Der Pullback φ^* . Sei $\varphi : U \rightarrow V$ eine unendlich oft differenzierbare Abbildung für zulässige Mengen $U \subset \mathbb{R}^n$ und $V \subset \mathbb{R}^m$. Dann gibt es eine *graderhaltende* Abbildung

$$\varphi^* : A^\bullet(V) \rightarrow A^\bullet(U)$$

eindeutig bestimmt durch die folgenden vier Eigenschaften

¹⁵Für $(dx_I \wedge dx_J) \wedge dx_K = dx_I \wedge (dx_J \wedge dx_K)$ benutze Induktion nach $i + j + k$ und bei festem $i + j + k$ Induktion nach $\max(i, j, k)$. Der Induktionsanfang $i = j = k = 1$ ist trivial. Sei $j > 1$, also $dx_J = dx_U \wedge dx_V$. Dann gilt $(dx_I \wedge dx_J) \wedge dx_K = (dx_I \wedge (dx_U \wedge dx_V)) \wedge dx_K = ((dx_I \wedge dx_U) \wedge dx_V) \wedge dx_K = (dx_I \wedge dx_U) \wedge (dx_V \wedge dx_K) = dx_I \wedge (dx_U \wedge (dx_V \wedge dx_K)) = dx_I \wedge (dx_J \wedge dx_K)$. Ist $k > 1$ und $dx_K = dx_U \wedge dx_V$, dann $dx_I \wedge (dx_J \wedge dx_K) = dx_I \wedge (dx_J \wedge (dx_U \wedge dx_V)) = dx_I \wedge ((dx_J \wedge dx_U) \wedge dx_V) = (dx_I \wedge (dx_J \wedge dx_U)) \wedge dx_V = ((dx_I \wedge dx_J) \wedge dx_U) \wedge dx_V = (dx_I \wedge dx_J) \wedge (dx_U \wedge dx_V) = (dx_I \wedge dx_J) \wedge dx_K$. Analog für $i > 1$.

4 Differentiation

1. φ^* ist $\mathbb{R}$ -linear
2. φ^* ist multiplikativ $\varphi^*(\omega \wedge \eta) = \varphi^*(\omega) \wedge \varphi^*(\eta)$ für alle $\omega, \eta \in A^\bullet(V)$
3. φ^* vertauscht mit der Cartan Ableitung: $\varphi^*(d\omega) = d\varphi^*(\omega)$ für alle $\omega \in A^\bullet(V)$
4. Für Nullformen $\omega = f(y)$ aus $A^0(V)$, d.h. Funktionen, gilt

$$\boxed{\varphi^*(f)(x) = f(\varphi(x))}.$$

Beachte $\varphi^*(\sum_I \omega_I(y) dy_I) = \sum_I \omega_I(\varphi(x)) \varphi^*(dy_I)$ und $\varphi^*(dy_I) = \varphi^*(dy_{m_1} \wedge \cdots \wedge dy_{m_i}) = \varphi^*(dy_{m_1}) \wedge \cdots \wedge \varphi^*(dy_{m_i})$ sowie $\varphi^*(dy_k) = d\varphi^*(y_k) = d\varphi_k$ für alle $k = 1, \dots, m$. Also

$$\boxed{\varphi^*(dy_k) = \sum_{l=1}^n \frac{\partial \varphi_k}{\partial x_l}(x) dx_l} \quad , \quad \frac{\partial \varphi_k}{\partial x_l}(x) = (D\varphi(x))_{kl}.$$

Aus der **Leibniz Formel** für die Determinante der Matrix $D\varphi(x)$ folgt daher im Spezialfall $n = m$ für Formen $\omega = f(y) \cdot dy_1 \wedge \cdots \wedge dy_n \in A^n(V)$

$$\boxed{\varphi^*(f(y) \cdot dy_1 \wedge \cdots \wedge dy_n) = f(\varphi(x)) \cdot \det D\varphi(x) \cdot dx_1 \wedge \cdots \wedge dx_n}.$$

Aus der Substitutionsformel (Satz 4.27) folgt daher sofort

Korollar 4.28. Für orientierte Koordinatenwechsel¹⁶ $\varphi : U \rightarrow V$ und Formen $\omega = f(y) \cdot \omega_n$ in $A_c^n(V)$ gilt

$$\boxed{\int_U \varphi^*(\omega) = \int_V \omega}.$$

Lemma 4.29. Zweimaliges Anwenden der Cartanableitung $d^2 : A^i(U) \rightarrow A^{i+2}(U)$ gibt die Nullabbildung

$$\boxed{d^2 = 0}.$$

Beweis. Im Spezialfall $i = 0$ ist wegen $d^2(f) = d(\sum_{\mu=1}^n \frac{\partial f}{\partial x_\mu}(x) \cdot dx_\mu)$

$$d^2(f) = \sum_{\nu=1}^n \sum_{\mu=1}^n \frac{\partial^2 f}{\partial x_\nu \partial x_\mu}(x) \cdot dx_\nu \wedge dx_\mu.$$

Da $dx_\nu \wedge dx_\mu$ alternierend in ν und μ ist, und andererseits wegen Satz 4.11 die zweiten partiellen Ableitungen symmetrisch in ν und μ sind, verschwindet dieser Ausdruck. Damit ist der Fall $i = 0$ gezeigt. Für $i > 0$ ist oBdA $\omega = f(x) \cdot dx_I$ für $f \in C^\infty(X)$. Nach der Definition der Cartanableitung ist $d^2(f(x) \cdot dx_I) = d^2(f(x)) \wedge dx_I$ und $d^2 f = d(df) = 0$ wurde bereits gezeigt (der Fall $i = 0$). $\square$

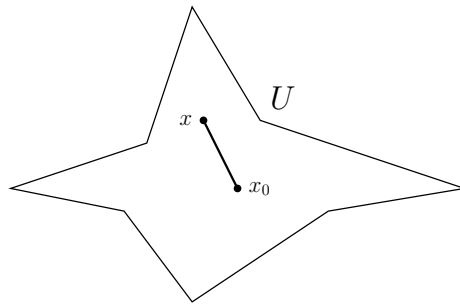
Dies zeigt, dass alle **exakten** Formen $\omega = d\eta$ **geschlossene** Formen sind, d.h. $d\omega (= d^2\eta) = 0$. Für Differentialformen vom Grad > 0 gilt folgende Umkehrung:

¹⁶zwischen offenen Teilmengen $U, V \subseteq \mathbb{R}^n$ wie in Abschnitt 4.11. Orientiert bedeutet dabei $\text{sign}(\det D\varphi(x)) > 0$.

Satz 4.30 (Poincare Lemma). Sei $U \subseteq \mathbb{R}^n$ offen und **sternförmig** und sei $\omega \in A^j(U)$. Dann gilt

- $d\omega = 0$ für $j > 0 \implies \exists \eta \in A^{j-1}(U)$ mit $\omega = d\eta$.
- $d\omega = 0$ im Fall $j = 0 \implies \omega \in A^0(U)$ ist konstant.

Sternförmigkeit bedeutet: Es gibt einen Punkt $x_0 \in U$, einen sogenannten *Sternmittelpunkt*, so dass für alle $x \in U$ der *Verbindungsweg* $x_0 + \{t(x - x_0) \mid 0 \leq t \leq 1\}$ auch in U liegt. Ein Sternmittelpunkt ist nicht eindeutig bestimmt, wie man an folgendem Beispiel für x_0 sieht:



Zum Beweis des Poincare Lemmas genügt die Existenz von Operatoren $I: A^j(U) \rightarrow A^{j-1}(U)$ mit der Eigenschaft $\omega - \delta_{j0} \cdot \omega(x_0) = I(d\omega) + dI(\omega)$.

Der eindimensionale Fall. Für eine zulässige Teilmenge $U \subset \mathbb{R}$ reduziert sich der allgemeine Differentialformenkomplex auf die Grade 0 und 1. Es bleibt also nur

$$A^0(U) = C^\infty(U) \xrightarrow{d} A^1(U) = C^\infty(U) \cdot dx ,$$

und für $f(x) \in C^\infty(U)$ ist die Cartan Ableitung d gegeben durch

$$f(x) \mapsto df(x) = f'(x) \cdot dx .$$

Das Poincare Lemma im eindimensionalen Fall ist damit fast der Hauptsatz der Differential- und Integralrechnung. Im Grad $j = 1$ besagt es nämlich, daß jede Funktion $g(x) \in C^\infty(U)$ eine Stammfunktion besitzt, da im letzten Grad $j = n$ (hier ist $n = 1$) automatisch $d\omega = 0$ gilt für die Differentialform $\omega = g(x)dx$. Im Grad 0 besagt das Poincare Lemma, daß Stammfunktionen eindeutig sind bis auf eine lokalkonstante Funktion. Allerdings gibt es zwei Einschränkungen. Erstens, wir beschränken uns hier auf C^∞ -Funktionen anstelle von stetigen Funktionen. Zweitens, erst der Satz von Stokes wird die noch fehlende Verbindung zur Integrationstheorie herstellen. Die höherdimensionale Integrationstheorie werden wir dazu noch verfeinern müssen.

4.13 Beweis des Poincare Lemmas

Lemma 4.31. $U \subseteq \mathbb{R}^n$ sei offen und sternförmig¹⁷. Dann ist für $f(x) \in C^\infty(U)$ das Integral

$$I^{(j)}(f) = \int_0^1 t^j f(tx) dt \quad , \quad (\text{für } j \in \mathbb{N})$$

als Funktion von $x \in U$ definiert, und als solche wieder eine Funktion in $C^\infty(U)$.

Wir formulieren nun, unter Vorgriff auf das Kapitel 5, einen allgemeinen Satz aus der Theorie Lebesgue integrierbarer Funktionen, welcher in unserem Fall $Y = [0, 1]$ für stetig differenzierbare Funktionen wegen $C(Y) \subset L(Y)$ unmittelbar anwendbar ist. Die Behauptung von Lemma 4.31 folgt unmittelbar aus folgendem (später bewiesenen **Vertauschungssatz**; Beweis siehe Abschnitt 6.5)

Satz 4.32. Sei $Y = \mathbb{R}^n$, $Y = \mathbb{Z}$ oder $Y = \mathbb{N}$. Sei $f(x, y) : [a, b] \times Y \rightarrow \mathbb{R}$ partiell differenzierbar nach x . Ist $f(x, y)$ für feste x in $L(Y)$, und gilt unabhängig von x die Abschätzung $|\partial_x f(x, y)| \leq F(y)$ für ein $F \in L(Y)$, dann ist

$$g(x) = \int_Y f(x, y) dy$$

differenzierbar auf $[a, b]$ als Funktion von x und es gilt $g'(x) = \int_Y \partial_x f(x, y) dy$.

Für $f(x) \in C^\infty(U)$ sei $f_\alpha := \frac{\partial f(x)}{\partial x_\alpha}$. Aus dem Hauptsatz folgt $f(x) = t^j f(tx) \big|_0^1 = \int_0^1 \frac{d}{dt} (t^j f(tx)) dt$ (für $j > 0$). Aus $\frac{d}{dt} f(tx) = \sum_\alpha x_\alpha f_\alpha(tx)$ folgt dann

Lemma 4.33. $f(x) = j \cdot I^{(j-1)}(f)(x) + \sum_{\alpha=1}^n x_\alpha I^{(j)}(f_\alpha)(x)$ für alle $j > 0$.

Wir kommen nun zum eigentlichen *Beweis des Poincare Lemmas*.

Beweis. Für $\beta \in \{1, \dots, n\}$ und $J \subset \{1, \dots, n\}$ definieren wir die **Kontraktionen**

$$\boxed{\partial_\beta \vee dx_J = 0} \quad (\text{im Fall } \beta \notin J) \quad \boxed{\partial_\beta \vee dx_J = \varepsilon \cdot dx_{J \setminus \{\beta\}}} \quad (\text{im Fall } \beta \in J)$$

mit dem eindeutig bestimmten Vorzeichen $\varepsilon \in \{\pm 1\}$ so daß $dx_\beta \wedge \varepsilon \cdot dx_{J \setminus \{\beta\}} = dx_J$.

Der Operator I . Wir definieren nun für alle $j > 0$ den $\mathbb{R}$ -linearen Operator

$$I : A^j(U) \longrightarrow A^{j-1}(U) ,$$

durch $I(\sum_J \omega_J(x) \cdot dx_J) = \sum_J I^{(j-1)}(\omega_J)(x) \cdot i_E(dx_J)$ mit $i_E(dx_J) := \sum_{\beta=1}^n x_\beta \cdot \partial_\beta \vee dx_J$.

Für $j = 0$ folgt das Poincare Lemma aus dem Mittelwertsatz 4.8. Genauer gilt

Bemerkung: Im Grad $j = 0$ gilt $\boxed{(I \circ d)f(x) = f(x) - f(0)}$ (Beweis wie bei Lemma 4.33).

¹⁷Wir nehmen zur Vereinfachung der Notation $x_0 = 0$ für den Sternmittelpunkt an.

Im Fall $j > 0$ setze $\eta = I(\omega) \in A^{j-1}(U)$. Aus $d\omega = 0$ folgt dann das Poincare Lemma $\omega = (d \circ I + I \circ d)\omega = d(I(\omega)) + 0 = d\eta$ auf Grund der folgenden

Homotopieformel: Für $\omega \in A^j(U)$ und Grad $j > 0$ gilt

$$\boxed{\omega = (d \circ I + I \circ d) \omega}.$$

Da diese Formel linear in ω ist, kann man für ihren Beweis $\omega = f(x) \cdot dx_J$ annehmen:

1.Schritt. Es gilt $I(\omega) = I^{(j-1)}(f) \cdot \sum_{\beta} x_{\beta} \cdot (\partial_{\beta} \vee dx_J)$. Wegen der Produktformel für die Cartanableitung ist

$$(d \circ I)(\omega) = d\left(I^{(j-1)}(f)\right) \wedge \left(\sum_{\beta} x_{\beta} \cdot (\partial_{\beta} \vee dx_J)\right) + I^{(j-1)}(f) \cdot d\left(\sum_{\beta} x_{\beta} \cdot (\partial_{\beta} \vee dx_J)\right).$$

Es gilt $d(I^{(j-1)}(f)) = \sum_{\alpha} (\partial_{\alpha} \int_0^1 t^{j-1} f(tx) dt) dx_{\alpha} = \sum_{\alpha} (\int_0^1 t^{j-1} \partial_{\alpha} f(tx) dt) dx_{\alpha}$ wegen dem Vertauschungssatz 4.32. Also $d(I^{(j-1)}(f)) = \sum_{\alpha} (\int_0^1 t^j f_{\alpha}(tx) dt) dx_{\alpha} = \sum_{\alpha} I^{(j)}(f_{\alpha}) dx_{\alpha}$. Ausserdem ist $d(\sum_{\beta} x_{\beta} (\partial_{\beta} \vee dx_J)) = \sum_{\beta \in J} dx_{\beta} \wedge (\partial_{\beta} \vee dx_J) = |J| \cdot dx_J = j \cdot dx_J$. Es folgt

$$(d \circ I)(\omega) = j \cdot I^{(j-1)}(f) \cdot dx_J + \left(\sum_{\alpha} I^{(j)}(f_{\alpha}) \cdot dx_{\alpha}\right) \wedge \left(\sum_{\beta} x_{\beta} \cdot (\partial_{\beta} \vee dx_J)\right).$$

2.Schritt. Andererseits ist $(I \circ d)(\omega) = I(\sum_{\alpha} f_{\alpha} \cdot dx_{\alpha} \wedge dx_J)$ gleich

$$(I \circ d)(\omega) = \sum_{\alpha} I^{(j)}(f_{\alpha}) \sum_{\beta} x_{\beta} \cdot (\partial_{\beta} \vee (dx_{\alpha} \wedge dx_J)).$$

3.Schritt. Wegen¹⁸ $dx_{\alpha} \wedge (\partial_{\beta} \vee dx_J) + \partial_{\beta} \vee (dx_{\alpha} \wedge dx_J) = \delta_{\alpha\beta} \cdot dx_J$ vereinfacht sich die Formel bei der Addition von $(d \circ I)(\omega)$ und $(I \circ d)(\omega)$ und liefert für $j > 0$ dank Lemma 4.33 die Homotopieformel

$$(dI + Id)(\omega) = \left(j \cdot I^{(j-1)}(f) + \sum_{\alpha} x_{\alpha} I^{(j)}(f_{\alpha})\right) \cdot dx_J = f(x) \cdot dx_J = \omega.$$

□

Bemerkung. Die Sternförmigkeit von U ist wesentlich für das Poincare Lemma. Die gelochte Ebene $U = \mathbb{R}^2 \setminus \{0\}$ ist nicht sternförmig! Die 1-Form $\omega = \text{Im}(\frac{dz}{z})$, in reellen Koordinaten

$$\omega = \frac{xdy - ydx}{x^2 + y^2},$$

¹⁸oder $\frac{\partial}{\partial \theta_{\beta}} \theta_{\alpha} + \theta_{\alpha} \frac{\partial}{\partial \theta_{\beta}} = \delta_{\alpha\beta}$ im Sinne von Abschnitt 5.15. Es genügt dies auf den Erzeugern $\theta^J \in \mathcal{S}_n$ zu zeigen.

Zum Beweis unterscheidet man die Fälle $|\{\alpha, \beta\} \cap J| = 0, 1, 2$. Diese Formel steht in Analogie zur Heisenberg Kommutatorformel $\frac{\partial}{\partial x_{\beta}} x_{\alpha} - x_{\alpha} \frac{\partial}{\partial x_{\beta}} = \delta_{\alpha\beta}$.

4 Differentiation

hat eine Singularität im Ursprung. Aber es gilt $\omega \in A^1(U)$ und man zeigt leicht $d\omega = 0$ auf U . Für die Abbildung $\varphi : \mathbb{R} \rightarrow U$ definiert durch $\varphi(t) = (\cos(t), \sin(t))$ gilt

$$\varphi^*(\omega) = \frac{\cos(t) \sin(t)' - \sin(t) \cos(t)'}{\cos(t)^2 + \sin(t)^2} dt = dt.$$

Hätte ω eine Stammfunktion η auf U , d.h. würde $d\eta = \omega$ gelten für ein $\eta \in A^0(U)$, wäre

$$\varphi^*(\omega) = \varphi^*(d\eta)(t) = d\varphi^*(\eta)(t) = d(\eta(\varphi(t))) = \eta'(\varphi(t)) \cdot dt.$$

Ein Vergleich ergäbe daher $\eta'(\varphi(t)) dt = dt$, und nach dem Hauptsatz wäre folglich bis auf eine Integrationskonstante C

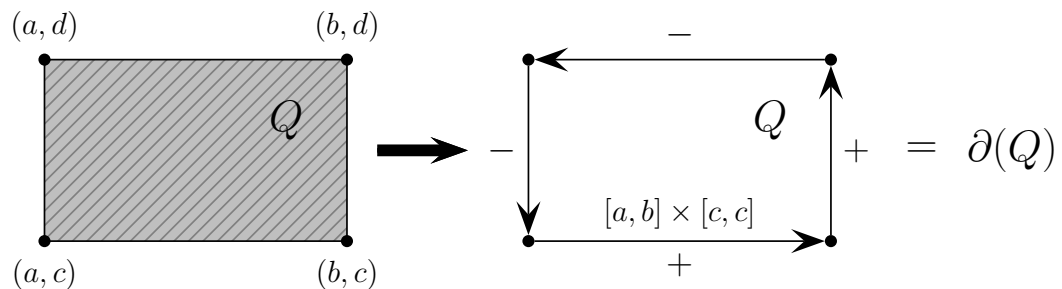
$$\eta(\varphi(t)) = t + C.$$

Ein Widerspruch, denn $\eta(\varphi(t))$ ist periodisch in t (mit Periode 2π), während die Funktion $t + C$ für kein C in t periodisch ist. *Das Poincare Lemma gilt also nicht für $U = \mathbb{R}^2 \setminus \{0\}$.*

4.14 Satz von Stokes für Quader

Sei $Q = \prod_{i=1}^n [a_i, b_i]$ ein nichtdegenerierter Quader im $\mathbb{R}^n$. Wir betrachten orientierte Quader $\varepsilon \cdot Q$ für eine **Orientierung** $\varepsilon = \pm 1$. Der Rand ∂Q eines Quaders Q ist in natürlicher Weise die Vereinigung von $2n$ nicht degenerierten orientierten Quadern im $\mathbb{R}^{n-1}$ (aber degeneriert im $\mathbb{R}^n$). Wir erläutern dies obdA im Fall $n = 2$ des Quaders $Q = +[a_1, b_1] \times [a_2, b_2]$ (positiv orientiert). Hier ist

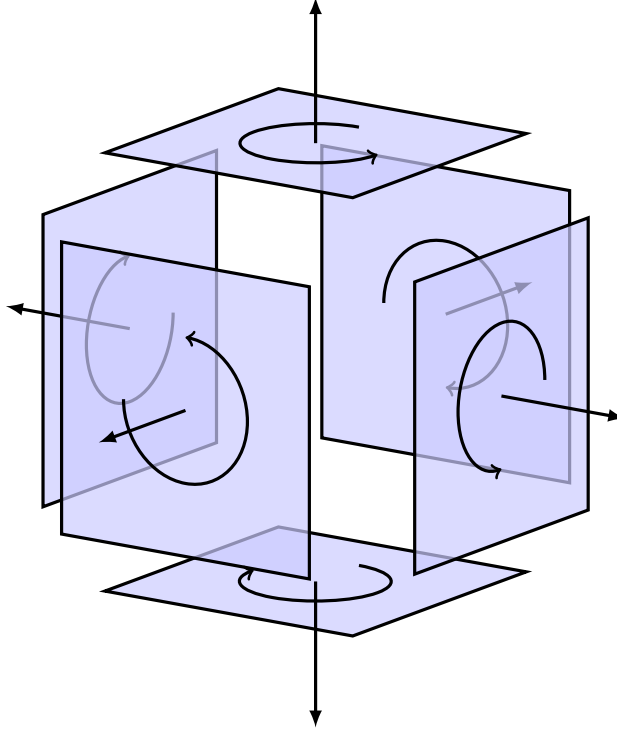
$$\partial(Q) := +[a_1, b_1] \times [a_2, a_2] + [b_1, b_1] \times [a_2, b_2] - [a_1, b_1] \times [b_2, b_2] - [a_1, a_1] \times [a_2, b_2].$$



Die mathematisch positive Orientierung des $\mathbb{R}^2$ oder von Q entspricht einer Orientierung im Gegenuhrzeigersinn. Diese Orientierung auf Q induziert eine Orientierung des Randes ∂Q . Fügt man mehrere (positiv orientierte) Quader nahtlos zusammen, dann haben Ränder zweier jeweils aneinander stossenden Quader immer umgekehrte Orientierung. Bildet man den Rand als formale Summe von Kanten heben sich diese inneren Kanten dann automatisch weg und es verbleibt

nur der umrandende Kantenzug übrig. Dies führt sehr leicht zu einer Verallgemeinerung des nächsten Satzes. Die Situation ist in höheren Dimension vollkommen analog. Im Fall $n = 3$ ist

$$\begin{aligned} \partial(Q) := & \{b_1\} \times [a_2, b_2] \times [a_2, b_2] - [a_1, b_2] \times \{b_2\} \times [a_3, b_3] + [a_1, b_1] \times [a_2, b_2] \times \{b_3\} \\ & - \{a_1\} \times [a_2, b_2] \times [a_2, b_2] + [a_1, b_2] \times \{a_2\} \times [a_3, b_3] - [a_1, b_1] \times [a_2, b_2] \times \{a_3\} . \end{aligned}$$



Für nichtdegenerierte kompakte Quader $Q \subset \mathbb{R}^n$ gilt

Satz 4.34. (*Baby Stokes*) Für jede Differentialform $\omega \in A^{n-1}(Q)$ gilt

$$\boxed{\int_Q d\omega = \int_{\partial Q} \omega|_{\partial Q}} .$$

Beweis. Per Definition ist $\omega|_{\partial Q} := i^*(\omega)$ für $i : \partial(Q) \hookrightarrow Q$ der *Pullback* von ω auf die $2n$ Randflächen von Q für die ‘offensichtlichen’ linearen Inklusionen i . Man definiert dann für

$$\omega = \sum_{i=1}^n f_i(x) \cdot dx_{\{1, \dots, n\} \setminus \{i\}} .$$

das Integral als die Summe über die $2n$ Beiträge des Randes ∂Q verstehen mit dem zugehörigen Orientierungsfaktor. Dieser ist wie folgt¹⁹ definiert: $\{b_i\} \times \prod_{\nu \neq i} [a_\nu, b_\nu]$ bekommt die Orientierung $(-1)^{i-1}$ und $\{a_i\} \times \prod_{\nu \neq i} [a_\nu, b_\nu]$ die Orientierung $-(-1)^{i-1}$. Da es sich bei den

¹⁹Anschaulich orientiert man die einzelnen Wände des Quaders Q mit Hilfe der verallgemeinerten Dreifinger-Regel, indem man jeweils den Daumen *von innen nach aussen* senkrecht durch die jeweilige Wand von Q richtet.

4 Differentiation

Summanden um $(n - 1)$ -dimensionale Euklidische Standardintegrale handelt, ist immer jeweils die i -te Integration bei diesen Summanden wegzulassen. Für $Q = \prod_{i=1}^n [a_i, b_i]$ mit $a_i < b_i, \forall i$ ist dann

$$\begin{aligned} \int_{\partial Q} \omega|_{\partial Q} &:= \sum_{i=1}^n (-1)^{i-1} \int_{a_1}^{b_1} \dots \int_{a_n}^{b_n} f_i(x_1, \dots, x_{i-1}, b_i, x_{i+1}, \dots, x_n) \cdot dx_{\{1, \dots, n\} \setminus \{i\}} \\ &\quad - (-1)^{i-1} \sum_{i=1}^n \int_{a_1}^{b_1} \dots \int_{a_n}^{b_n} f_i(x_1, \dots, x_{i-1}, a_i, x_{i+1}, \dots, x_n) \cdot dx_{\{1, \dots, n\} \setminus \{i\}}. \end{aligned}$$

Die Aussage des Satzes ist additiv in ω , daher ist oBdA $\omega = f_i(x) \cdot dx_{\{1, \dots, n\} \setminus \{i\}}$ für ein festes i . Damit vereinfacht sich die obige Formel zu

$$\int_{\partial Q} \omega|_{\partial Q} := (-1)^{i-1} \int_{a_1}^{b_1} \dots \int_{a_n}^{b_n} (f_i(x_1, \dots, x_{i-1}, x, x_{i+1}, \dots, x_n)|_{x=a_i}^{x=b_i}) \cdot dx_{\{1, \dots, n\} \setminus \{i\}}.$$

Andererseits gilt für $\omega = f_i(x) \cdot dx_{\{1, \dots, n\} \setminus \{i\}}$ dann $d\omega = \partial_i f_i(x) (-1)^{i-1} \cdot dx_{\{1, \dots, n\}}$ wegen $dx_i \wedge dx_{\{1, \dots, n\} \setminus \{i\}} = (-1)^{i-1} \cdot dx_{\{1, \dots, n\}}$. Das Integral $\int_Q d\omega$ ist somit gegeben durch

$$\int_Q d\omega = \int_{a_1}^{b_1} \dots \int_{a_n}^{b_n} \partial_i f_i(x) (-1)^{i-1} \cdot dx_1 \dots dx_n,$$

stimmt also (nach Fubini, Korollar 6.14) überein mit $\int_{\partial Q} \omega|_{\partial Q}$, wenn man den Hauptsatz (Satz 4.14) für die Integration über die i -te Koordinate benutzt. $\square$

Satz 4.35. Sei U offen im $\mathbb{R}^n$ und $\omega \in A^{n-1}(U)$. Verschwindet $\int_{\partial Q} \omega|_{\partial Q}$ für jeden Quader Q in U , dann gilt $d\omega = 0$.

Beweis. Sei $d\omega = f(x) dx_1 \wedge \dots \wedge dx_n$. Da f gleichmäßig stetig ist auf Q , gibt es für $\varepsilon > 0$ ein $\delta > 0$ mit $|f(x) - f(\xi)| < \varepsilon$ für $\|x - \xi\| < \delta$. Für $Q \subset K_\delta(\xi)$ folgt daher aus der Boxungleichung

$$-\varepsilon < f(\xi) - \frac{1}{\text{vol}(Q)} \int_Q d\omega < \varepsilon.$$

Unsere Annahme und Satz 4.34 zeigen $\frac{1}{\text{vol}(Q)} \int_Q d\omega = \frac{1}{\text{vol}(Q)} \int_{\partial Q} \omega|_{\partial Q} = 0$. Es folgt $|f(\xi)| < \varepsilon$ für alle $\varepsilon > 0$, also $f(\xi) = 0$ für alle $\xi \in U$. $\square$

4.15 Analytische Funktionen

Beispiel 4.36. Im metrischen Raum $Y = \mathbb{Z}$ (mit der Euklidischen Metrik) sind die folgenkompakten Teilmengen die endlichen Teilmengen von Y und jede (!) Funktion $f : \mathbb{Z} \rightarrow \mathbb{R}$ ist stetig. Die stetigen Funktionen mit kompaktem Träger auf $Y = \mathbb{Z}$ sind daher die Funktionen mit

endlichem Träger. Der Raum $C_c(\mathbb{Z})$ kann also identifiziert werden mit dem Raum der Folgen $\mathbb{Z} \ni n \mapsto f(n) \in \mathbb{R}$, für die fast alle Folgenglieder $f(n)$ Null sind. Daher definiert die endliche (!) Summe $I(f) = \sum_{n \in \mathbb{Z}} f(n)$ ein abstraktes Integral $I : B(\mathbb{Z}) := C_c(\mathbb{Z}) \rightarrow \mathbb{R}$ im Sinne von Abschnitt 3.18. Analoges gilt für die Teilmenge $\mathbb{N}$ anstelle von $\mathbb{Z}$. Man überlegt sich nun leicht: Eine Funktion $f : \mathbb{Z} \rightarrow \mathbb{R}^+$ ist in der monotonen Hülle $B^+(\mathbb{Z})$ genau dann, wenn $f(n) < \infty$ nur für endlich viele $n \in \mathbb{Z}$ gilt. In diesem Abschnitt wenden wir den späteren Satz 4.32 im Fall $Y = \mathbb{Z}$ oder genauer $Y = \mathbb{N}$ an um zu zeigen, daß die im folgenden betrachteten Potenzreihen gliedweise abgeleitet werden dürfen.

Für eine Folge a_l reeller (oder komplexer) Zahlen betrachten wir die formale **Potenzreihe** $\sum_{l=0}^{\infty} a_l x^l$ und fragen, für welche Werte $x = x_0 \in \mathbb{R}$ (oder $x = x_0 \in \mathbb{C}$) der Limes

$$\lim_{N \rightarrow \infty} \sum_{l=0}^N a_l \cdot x^l =: \sum_{l=0}^{\infty} a_l \cdot x^l$$

existiert. Wenn dieser Limes existiert, nennt man die Potenzreihe konvergent im Punkt $x = x_0$. Allgemeiner kann man anstatt der Monome $P_l(x) = a_l x^l$ auch homogene Polynome $P_l(x)$ vom Grad l in mehreren Variablen betrachten.

Seien $P_l(x)$ für $l = 0, 1, 2, \dots$ homogene Polynome $P_l : \mathbb{R}^n \rightarrow \mathbb{R}$ vom Grad l , dann gilt $P_l(t \cdot x) = t^l \cdot P_l(x)$ für alle $t \in \mathbb{R}$. Wie man leicht aus der Homogenität folgert, nimmt ein homogenes Polynom $P_l(x)$ auf einer Kugel $\|x\| \leq \rho$ ihr Maximum auf dem Rand $\|x\| = \rho$ an. Zur Untersuchung der Konvergenz von $\sum_{l=0}^{\infty} P_l(x)$ definiert man den **Konvergenzradius**

$$R := \sup \left\{ \rho \mid \exists C_\rho > 0 \text{ mit } \max_{\|y\|=\rho} (|P_l(y)|) \leq C_\rho \text{ für alle } l \right\}.$$

Der Konvergenzradius R ist eine Zahl in $\mathbb{R}_{\geq 0}^+ = \mathbb{R}_{\geq 0} \cup \{+\infty\}$. Nach Definition gilt für alle $x \in \mathbb{R}^n$ mit der Eigenschaft $\|x\| \leq \rho < R$ wegen der Homogenität von $P_l(x)$

$$|P_l(x)| \leq C_\rho \cdot \frac{\|x\|^l}{\rho^l}.$$

Im eindimensionalen Fall $n = 1$ ist dies gleichbedeutend mit der Bedingung

$$|a_l| \leq C_\rho \cdot \rho^l.$$

Satz 4.37. *Ist der Konvergenzradius $R > 0$, so ist die offene Kugel $X = \{x \in \mathbb{R}^n \mid \|x\| < R\}$ nichtleer und*

$$f(x) = \sum_{l=0}^{\infty} P_l(x)$$

konvergiert absolut für $x \in X$. Für $\rho' \in [0, R)$ ist die Konvergenz absolut und gleichmäßig auf der folgenkompakten Teilmenge $K = \{x \in \mathbb{R}^n \mid \|x\| \leq \rho'\}$ von X . Die Grenzfunktion $f(x)$ ist daher eine stetige Funktion der Variable x auf K , und definiert durch Variation von ρ' eine stetige Funktion auf X .

4 Differentiation

Beweis. Für gegebenes $x \in X$ sei $\rho < R$ und $0 \leq q < 1$, so daß x in der kompakten Kugel K vom Radius $\rho' = q\rho$ enthalten ist. Wir wollen $|P_l(x)|$ auf K abschätzen. Aus $\|x\| \leq \rho' = q\rho$ folgt $\|x\|^l / \rho^l = q^l$ und damit

$$|P_l(x)| \leq C_\rho \cdot q^l, \quad x \in K.$$

Für die Partialsummen $f_n(x) := \sum_{l=0}^n P_l(x)$ gilt dann wegen Lemma 1.16 für alle $m \geq n$

$$\|f_m - f_n\|_K = \max_{x \in K} |f_m(x) - f_n(x)| = \max_{x \in K} \left| \sum_{l=n+1}^m P_l(x) \right| \leq \sum_{l=n+1}^m C_\rho q^l \leq C_\rho \frac{q^{n+1}}{1-q}.$$

Wegen $|q| < 1$ geht $C_\rho \frac{q^{n+1}}{1-q}$ gegen Null für $n \rightarrow \infty$ (mit Schranken unabhängig von x). Die stetigen Polynome $f_n(x)$ bilden daher eine Cauchyfolge im Raum $C(K)$ der stetigen Funktionen auf K . Wegen Satz 2.24 ist $C(K)$ vollständig. Daher konvergieren die $f_n(x)$ gleichmäßig auf K gegen eine auf K stetige Grenzfunktion $f(x)$. Da $\{x \mid \|x\| < R\}$ die Vereinigung solcher K ist für geeignete $\rho < R$ und $q < 1$, folgt unsere Behauptung. $\square$

Wir betrachten jetzt den eindimensionalen Fall. Seien q, ρ wie im obigen Beweis. Wähle $\varepsilon > 0$ genügend klein, so daß $\tilde{q} := q(1 + \varepsilon) < 1$. Für $\tilde{C} := 1/\varepsilon$ gilt dann $l < \tilde{C}(1 + \varepsilon)^l$. Wir wollen nun den **Vertauschungssatz** 4.32 (für $Y = \mathbb{N}$) auf die Funktionen $f_l(x) = a_l x^l$ anwenden:

Sei $f_l(x) : [a, b] \rightarrow \mathbb{R}$ eine Folge auf $[a, b]$ differenzierbarer Funktionen. Ist die reelle Folge $f_l(x)$ in $L(\mathbb{Z})$ für jedes feste $x \in [a, b]$, und gilt unabhängig von x für die Ableitungen eine Abschätzung $|f'_l(x)| \leq F_l$ für eine reelle Folge $F_l \in L(\mathbb{Z})$. Dann ist

$$f(x) = \lim_{n \rightarrow \infty} \sum_{l=0}^n f_l(x) =: \sum_{l=0}^{\infty} f_l(x)$$

wohldefiniert und differenzierbar auf $[a, b]$ mit der Ableitung $f'(x) = \sum_{l=0}^{\infty} f'_l(x)$.

$|f_l(x)| = |a_l x^l|$ lässt sich durch $C_\rho q^l$, $|q| < 1$ abschätzen, und $|f'_l(x)| = |l a_l x^{l-1}|$ lässt sich durch $F_l = C \tilde{q}^{l-1}$, $|\tilde{q}| < 1$ abschätzen für alle $|x| \leq \rho q$; letzteres obdA für $|x| = q\rho$ wegen

$$|f'_l(x)| = |a_l l x^{l-1}| = |a_l x^l| \frac{l}{|x|} \leq C_\rho q^l \frac{l}{|x|} \leq C_\rho \tilde{C} (1 + \varepsilon)^l q^l / \rho q = C \cdot \tilde{q}^{l-1}.$$

Wegen Lemma 1.16 gilt damit $f_l, F_l \in L(\mathbb{Z})$. Daher sind die Voraussetzungen für den Vertauschungssatz 4.32 erfüllt. Wendet man ihn an, erhält man: $f(x)$ differenzierbar im Intervall $[-q\rho, q\rho]$ und die Potenzreihe kann dort gliedweise abgeleitet werden, und die abgeleitete Potenzreihe besitzt (mindestens) denselben Konvergenzradius R . Iteriert man schließlich diesen Schluß, und betrachtet anschliessend den Limes $q \rightarrow 1$ und $\rho \rightarrow R$, so folgt

Satz 4.38. Ist der Konvergenzradius R der Potenzreihe $f(x) = \sum_{l=0}^{\infty} a_l x^l$ echt grösser als 0, dann ist im Bereich $\|x\| < R$ die Funktion $f(x)$ unendlich oft differenzierbar und die Potenzreihe ist gliedweise ableitbar, d.h. es gilt

$$f'(x) = \sum_{l=0}^{\infty} l a_l x^{l-1}$$

und der Konvergenzradius der abgeleiteten Potenzreihe ist wieder R .

Ist der Konvergenzradius $R > 0$, folgt aus Satz 4.38 durch n -faches Ableiten der Potenzreihe $f(x)$ die Formel $f^{(n)}(x) = n! \cdot a_n + \sum_{l>n} a_l \cdot (x^l)^{(n)}$. Setzt man $x = 0$, wird die Restsumme über alle Summanden $l > n$ gleich Null. Es folgt

Lemma 4.39. Ist der Konvergenzradius R der Potenzreihe $f(x) = \sum_{n=0}^{\infty} a_n x^n$ echt größer als 0, dann gilt die **Taylorformel**

$$a_n = \frac{f^{(n)}(0)}{n!}.$$

Insbesondere sind alle Koeffizienten a_n durch die Funktion $f(x)$ eindeutig bestimmt.

Nun einige einfache, aber fundamentale **Beispiele für Potenzreihen**:

Die Funktion $f(t) = (1+t)^\alpha$ ist auf dem Intervall $(-1, \infty)$ differenzierbar und die eindeutig bestimmte Lösung der linearen Differentialgleichung $(1+t)f'(t) = \alpha \cdot f(t)$ zum Anfangswert $f(0) = 1$. Es folgt für $\binom{\alpha}{n} := \frac{1}{n!} \prod_{i=1}^n (\alpha + 1 - i)$

Lemma 4.40. Für alle $t \in (-1, 1)$ und $\alpha \in \mathbb{R}$ gilt²⁰

$$(1+t)^\alpha = \sum_{n=0}^{\infty} \binom{\alpha}{n} \cdot t^n.$$

Beweis. Beide Seiten erfüllen auf $(-1, 1)$ dieselbe Differentialgleichung, auf der rechten Seite wegen $n \cdot \binom{\alpha}{n} + (n+1) \cdot \binom{\alpha}{n+1} = \alpha \cdot \binom{\alpha}{n}$ und $\binom{\alpha}{0} = 1$ bei gliedweisem Ableiten. $\square$

Im ganzzahligen Fall $\alpha = m \in \mathbb{N}$ reduziert sich Lemma 4.40 auf das klassische **Binomial Theorem** und $\binom{m}{n} = \frac{1}{n!} \prod_{i=1}^n (m+1-i)$ wird dann Null für alle $n \geq m+1$. Man erhält dadurch folgende Formel für die **Binomialkoeffizienten** $\binom{m}{n}$

$$\binom{m}{n} = \frac{m!}{n!(m-n)!}, \quad 0 \leq n \leq m.$$

Analog zeigt man

Lemma 4.41. Für alle $t \in (-1, 1)$ gilt $\log(1+t) = \sum_{n=1}^{\infty} (-1)^{n-1} \frac{t^n}{n}$.

Satz 4.42. Für alle $t \in \mathbb{R}$ gilt²¹

$$\exp(t) = \sum_{n=0}^{\infty} \frac{t^n}{n!}.$$

Allgemeiner ist für eine $m \times m$ -Matrix X die **Matrix Exponentialfunktion**

$$f(t) = \exp(tX) = \sum_{n=0}^{\infty} \frac{X^n}{n!} t^n$$

²⁰Für $0 < \rho := |x| < 1$ gilt $|\binom{\alpha}{l} \rho^l| \leq C \rho^l \prod_{i=l_0}^l |1 - \frac{\alpha+1}{i}| \leq C$ für alle $l \geq l_0$, wenn l_0 so groß ist daß $|\frac{\alpha+1}{l_0}| < 1$. Hierbei ist C eine geeignet gewählte von l unabhängige Konstante. Für $\sum_{n=0}^{\infty} \binom{\alpha}{n} \cdot t^n$ ist damit der Konvergenzradius $R \geq 1$.

²¹Für beliebiges ρ gilt $|\frac{\rho^l}{l!}| \leq C \cdot |\frac{\rho}{l_0}|^l \leq C$ für alle $l \geq l_0$, falls l_0 so groß ist daß $|\frac{\rho}{l_0}| < 1$. Hierbei ist oBdA $C = \prod_{1 \leq i < l_0} \frac{l_0}{i}$. Für $\sum_{n=0}^{\infty} \frac{t^n}{n!}$ ist daher der Konvergenzradius $R = +\infty$.

4 Differentiation

erklärt als matrixwertige Funktion $f : \mathbb{R} \rightarrow M_{m,m}(\mathbb{R})$, und ist dabei eindeutig bestimmt durch die lineare Differentialgleichung ($\circ$ ist hier die Matrixmultiplikation)

$$\frac{d}{dt}f(t) = X \circ f(t) \quad , \quad f(0) = id .$$

Für $t, s \in \mathbb{R}$ gilt die **Matrix-Funktionalgleichung**

$$\boxed{\exp(sX) \circ \exp(tX) = \exp((s+t)X)} .$$

Insbesondere ist $\exp(tX)$ eine invertierbare Matrix

$$\exp(tX) \in Gl(m, \mathbb{R})$$

mit Umkehrmatrix $\exp(-tX)$. Zum Beweis der Funktionalgleichung genügt, daß beide Seiten dieselbe Differentialgleichung $\frac{d}{ds}F(s) = X \circ F(s)$ erfüllen mit $F(0) = \exp(tX)$.

Ähnlich gilt $\exp(tX) \in Gl(m, \mathbb{C})$ für komplexe Matrizen $X \in M_{m,m}(\mathbb{C})$. Eine komplexe Matrix $M \in M_{m,m}(\mathbb{C})$ nennt man **unitär**, wenn $M^\dagger \circ M = id$ gilt; hierbei sei $M^\dagger = {}^T\overline{M}$ die hermitesch transponierte Matrix. Gilt

$$X^\dagger = -X ,$$

nennt man eine Matrix $X \in M_{m,m}$ **antihermitesch**. Die antihermiteschen Matrizen bilden eine **Lie Algebra**, denn für antihermitesche Matrizen X, Y gilt

$$[X, Y]^\dagger = (XY - YX)^\dagger = Y^\dagger X^\dagger - X^\dagger Y^\dagger = (-Y)(-X) - (-X)(-Y) = -(XY - YX) = -[X, Y] .$$

Die hermiteschen Matrizen dagegen bilden keine Lie Algebra!

Lemma 4.43. Für $M \in M_{m,m}(\mathbb{C})$ gilt: $\exp(tX)$ ist unitär für alle $t \in \mathbb{R}$ genau dann wenn X antihermitesch ist

$$\boxed{\exp(tX) \text{ ist unitär } \forall t \in \mathbb{R} \iff X \text{ ist antihermitesch} \iff \frac{X}{2\pi i} \text{ ist hermitesch}} .$$

Beweis. Ist $\exp(tX)$ unitär für alle $t \in \mathbb{R}$, gilt $\exp(tX)^\dagger \circ \exp(tX) = id$. Durch Ableiten folgt aus der Produktregel $(X \circ \exp(tX))^\dagger \circ \exp(tX) + \exp(tX)^\dagger \circ X \circ \exp(tX) = 0$. Setzt man $t = 0$, folgt $X^\dagger + X = 0$.

Ist umgekehrt X antihermitesch, dann gilt $\exp(tX)^\dagger \circ \exp(tX) = \exp(tX^\dagger) \circ \exp(tX) = \exp(-tX) \circ \exp(tX) = \exp(0) = id$. Im ersten Schritt wurde die Stetigkeit der Abbildung $X \mapsto X^\dagger$ benutzt, im zweiten Schritt die Annahme $X^\dagger = -X$ und im letzten Schritt die Matrix-Funktionalgleichung der Exponentialfunktion. $\square$

Übungsaufgabe. Zeige für reelles t die Aussage

$$\boxed{\exp(it) = \cos(t) + i \cdot \sin(t)} .$$

5 Ausgewählte Themen I

5.1 Wegintegrale

Sei U eine offene Teilmenge im $\mathbb{R}^n$ und $\omega \in A^1(U)$ eine 1-Form auf U

$$\omega = \sum_{\nu=1}^n F_{\nu}(x) dx_{\nu} .$$

Ein **Weg** $\gamma: [a, b] \rightarrow U$ ist eine stückweise stetig differenzierbare Funktion auf $[a, b]$ mit Stützstellen $a = t_0 \leq t_1 \leq \dots \leq t_{r-1} \leq t_r = b$; d.h. $\gamma|_{[t_{i-1}, t_i]}$ ist stetig differenzierbar auf jedem der Teilintervalle $[t_{i-1}, t_i]$. Wir definieren dann das **Wegintegral**

$$\int_{\gamma} \omega := \int_a^b \gamma^*(\omega) = \sum_{i=1}^r \int_{t_{i-1}}^{t_i} \gamma^*(\omega) .$$

Beachte $\gamma^*(dx_{\nu}) = d\gamma^*(x_{\nu}) = d\gamma_{\nu}(t) = \dot{\gamma}_{\nu}(t)dt$ (Ableitung nach t in jedem der Intervalle $[t_{i-1}, t_i]$) für $x = \gamma(t)$. Also ist für $F = (F_1, \dots, F_n)$, bis auf den Term dt ,

$$\gamma^*(\omega)(t) = \sum_{\nu=1}^n F_{\nu}(\gamma(t)) \dot{\gamma}_{\nu}(t) dt = \langle F(\gamma(t)), \dot{\gamma}(t) \rangle \cdot dt$$

das *Standardskalarprodukt* $\langle F(\gamma(t)), \dot{\gamma}(t) \rangle$ von dem *Vektor* $F(\gamma(t))$ mit dem *Tangentenvektor* $\dot{\gamma}(t)$. Besitzt ω eine Stammfunktion ϕ im Sinne von $\omega = d\phi$ für ein $\phi \in A^0(U)$, dann gilt

$$\int_{\gamma} d\phi = \phi(B) - \phi(A)$$

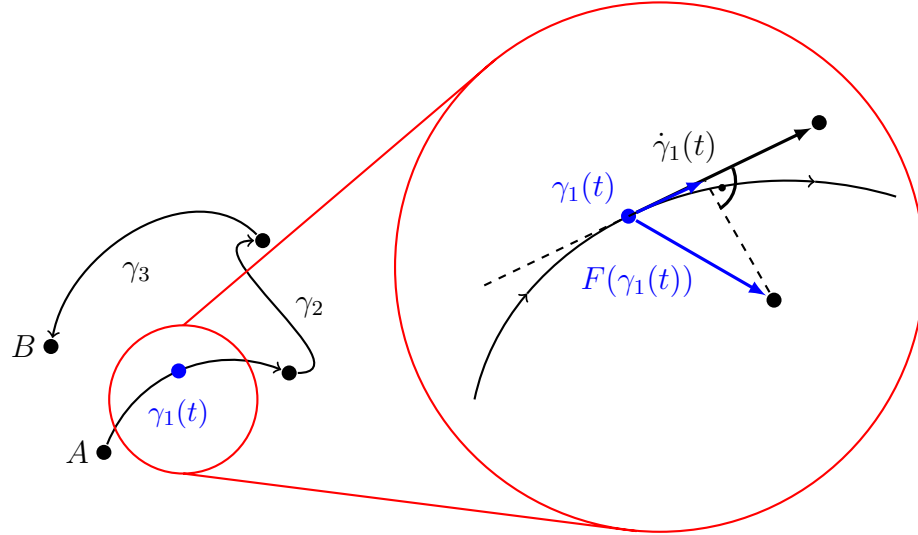
für $P = \gamma(a)$ (Anfangspunkt) und $B = \gamma(b)$ (Endpunkt des Weges); insbesondere hängt dann das Wegintegral nur vom Anfangspunkt A und vom Endpunkt B des Weges ab. Zum Beweis:

$$\int_{t_{i-1}}^{t_i} \gamma^*(\omega) = \int_{t_{i-1}}^{t_i} \gamma^*(d\phi) = \int_{t_{i-1}}^{t_i} d(\phi(\gamma(t))) = \int_{t_{i-1}}^{t_i} \frac{d}{dt} \phi(\gamma(t)) dt = \phi(\gamma(t)) \Big|_{t=t_{i-1}}^{t=t_i} .$$

Also ist $\int_{\gamma} d\phi$ als Teleskopsumme über $i = 1, \dots, r$ gleich $\phi(\gamma(t_r)) - \phi(\gamma(t_0)) = \phi(B) - \phi(A)$.

In der klassischen Mechanik (speziell im Fall $n = 3$) kann eine 1-Form ω im $\mathbb{R}^n$ als **Kraft** aufgefaßt werden repräsentiert durch den Vektor $F = (F_1, \dots, F_n)$, und $\int_{\gamma} \omega$ als **Arbeit** entlang des Weges γ . Ist die Kraft von der Form $\omega = d\phi$, nennt man ϕ ein **Potential**. Eine Kraft heißt

konservativ, wenn die Arbeit nicht vom Weg γ , sondern nur von dem Anfangspunkt $P = \gamma(a)$ und vom Endpunkt $B = \gamma(b)$ des Weges abhängt. Offensichtlich äquivalent dazu ist, daß das Wegintegral $\int_{\gamma} \omega$ für jeden geschlossenen Weg γ in U , d.h mit $\gamma(a) = \gamma(b)$, verschwindet. Existiert ein Potential ϕ , ist offensichtlich die Kraft konservativ.



Satz 5.1. Eine Form $\omega \in A^1(U)$ ist konservativ genau dann, wenn ein Potential $\phi \in A^0(U)$ existiert mit $d\phi = \omega$. Insbesondere ist dann $d\omega = 0$.

Beweis. Sei ω konservativ. ObdA sei U wegzusammenhängend, d.h. man kann jeden Punkt in U mit einem fixierten Punkt $x_0 \in U$ durch einen Weg verbinden. Setze $\phi(x) := \int_{\gamma} \omega$ für einen beliebigen Weg γ mit Anfangspunkt x_0 und Endpunkt x . Wir behaupten $d\phi = \omega$. Um dies in einem Punkt $\xi \in U$ zu zeigen, kann man U durch eine offene nichtleere Kreisscheibe $V = K_{\varepsilon}(\xi) \subset U$ ersetzen, denn $\phi(x + \xi) = \phi(\xi) + \int_{\tilde{\gamma}} \omega$ für irgendeinen Weg $\tilde{\gamma}$ in V von ξ nach x in V . Zum Beispiel $\tilde{\gamma}(t) = \xi + t \cdot x$ für $t \in [0, 1]$. Es gilt dann $\dot{\tilde{\gamma}}(t) = x$ und somit $\phi(x + \xi) - \phi(\xi) = \int_0^1 \langle x, F(\xi + tx) \rangle dt = \sum_{\nu} x_{\nu} \cdot \int_0^1 F_{\nu}(\xi + tx) dt$. Partielles Ableiten bei $x = 0$ gibt wegen Satz 4.32 und der Produktregel $\partial_{\nu} \phi(\xi) = \int_0^1 F_{\nu}(\xi) dt = F_{\nu}(\xi)$. Also gilt $d\phi = \omega$ im Punkt ξ . $\square$

Bemerkung 5.2. Insbesondere gilt daher in einer sternförmigen Teilmenge $V \subseteq U$ mit Sternmittelpunkt $x_0 = 0$ für $d\phi = \omega = \sum_{\nu} F_{\nu} dx_{\nu}$ die folgende Formel

$$\phi(x) - \phi(0) = \sum_{\nu=1}^n x_{\nu} \cdot \int_0^1 F_{\nu}(tx) dt .$$

Das Poincare Lemma liefert schliesslich

Satz 5.3. Sei U sternförmig. Dann ist $\omega \in A^1(U)$ konservativ genau dann, wenn $d\omega = 0$ gilt. Das Potential ϕ ist bis auf eine Konstante eindeutig durch ω bestimmt.

Beispiel. Wir betrachten folgende komplexwertige C^∞ -Form auf $U = \mathbb{R}^2 \setminus \{0\}$

$$\omega = \frac{dz}{z} = \frac{(x - iy) \cdot (dx + idy)}{x^2 + y^2} = \frac{xdx + ydy}{r^2} + i \frac{xdy - ydx}{x^2 + y^2} = d \log(r) + i \cdot \operatorname{Im}(\omega).$$

Ihr Realteil $\frac{xdx+ydy}{x^2+y^2}$ ist konservativ mit Potential $\log(r)$ in U . Ihr Imaginärteil $\operatorname{Im}(\omega) = \frac{xdy-ydx}{x^2+y^2}$ ist geschlossen aber nicht konservativ, denn nach Sektion 4.13 gilt $\int_\gamma \operatorname{Im}(\omega) = 2\pi$ für den geschlossenen Kreisweg $\gamma : [0, 2\pi] \rightarrow U$ definiert durch $\gamma(t) = (\cos(t), \sin(t))$. Dies liefert für den Kreisweg γ das Wegintegral

$$\boxed{\int_\gamma \frac{dz}{z} = 2\pi i}.$$

5.2 Holomorphe Funktionen

Eine Differentialform mit komplexen Koeffizienten $\omega = \alpha + i\beta$ wird definiert durch zwei reelle Differentialformen α und β , die man den Real- bzw. Imaginärteil von ω nennt. Für komplexwertige 1-Formen erklärt man das Wegintegral durch

$$\int_\gamma \omega := \int_\gamma \alpha + i \cdot \int_\gamma \beta.$$

Sei nun U eine offene Teilmenge der komplexen Ebene $\mathbb{C} = \mathbb{R}^2$. Für Funktionen f auf U schreiben wir dann $f(z) = f(x, y)$, falls $z = x + iy \in U$. Sei $\beta \in A^1(U)$ eine reelle 1-Form auf U . Diese lässt sich schreiben in der Gestalt $\beta = v(x, y)dx + u(x, y)dy$ für Funktionen $u, v \in C^\infty(U)$. Diese reelle 1-Form lässt sich interpretieren als Imaginärteil der komplexen 1-Form auf U

$$\omega = \left(u(x, y)dx - v(x, y)dy \right) + i \cdot \left(v(x, y)dx + u(x, y)dy \right).$$

Für die komplexwertige C^∞ -Funktion $f(z) = u(x, y) + iv(x, y)$ auf U und die komplexwertige 1-Form $dz := dx + idy$ gilt dann (durch Ausmultiplizieren)

$$\boxed{\omega = f(z) \cdot dz = (u + iv)(dx + idy)}.$$

Lemma 5.4. *Mit den obigen Annahmen und Bezeichnungen sind folgende Eigenschaften äquivalent:*

1. Auf U gilt $d\omega = 0$.
2. Auf U gelten die **Cauchy-Riemann Differentialgleichungen** $\partial_y v = \partial_x u$ und $\partial_x v = -\partial_y u$, oder kurz: $(\partial_x + i\partial_y)(u + iv) = 0$.
3. Die Jacobimatrix $Df(z)$ der Abbildung $f : U \rightarrow \mathbb{C} = \mathbb{R}^2$ hat für alle $z \in U$ die Gestalt

$$Df(z) = \begin{pmatrix} a(z) & b(z) \\ -b(z) & a(z) \end{pmatrix}.$$

Beweis. Dies folgt unmittelbar aus $d\omega = -(\partial_y u + \partial_x v)dx \wedge dy - i(\partial_y v - \partial_x u)dx \wedge dy$. $\square$

Definition 5.5. Eine komplexwertige Funktion $f : U \rightarrow \mathbb{C}$ der Gestalt $f(z) = u(z) + iv(z)$ mit $u, v \in C^\infty(U)$ heißt **holomorph** auf U , wenn die drei äquivalenten Bedingungen des letzten Lemmas für f erfüllt sind.

Beispiel. Die Funktion $f(z) = z$ ist offensichtlich holomorph auf ganz $\mathbb{C}$. Wegen $u(z) = x$ und $v(z) = y$ verifiziert man sofort die Cauchy-Riemann Differentialgleichungen.

Lemma 5.6. Die auf U holomorphen Funktionen bilden einen Unterring $\mathcal{O}(U)$ des Rings $C^\infty(U, \mathbb{C})$. Für $f, g \in \mathcal{O}(U)$ mit $g(z) \neq 0$ für alle $z \in U$ ist auch $f(z)/g(z)$ holomorph auf U .

Lemma 5.7. Sind $f : U \rightarrow V$ und $g : V \rightarrow W$ holomorphe Funktionen und U, V, W offene Teilmengen von $\mathbb{C}$, dann ist auch $g \circ f : U \rightarrow W$ holomorph.

Lemma 5.8. Ist $f : U \rightarrow \mathbb{C}$ holomorph, dann sind Realteil $u(x, y)$ und Imaginärteil $v(x, y)$ reelle **harmonische**¹ Funktionen auf U , ebenso $\log(|f(z)|)$ im Komplement der Nullstellen.

Das erste der drei letzten Lemmata folgt sofort aus der Produktregel mit Hilfe der Cauchy-Riemann Differentialgleichungen $D(fg) = D(f)g + fD(g) = 0$ für $D = \partial_x + i\partial_y$ und $f, g \in \mathcal{O}(U)$. Das zweite folgt aus der Kettenregel mit Hilfe von Lemma 5.4(3). Das dritte folgt aus $\Delta = \partial_x^2 + \partial_y^2 = (\partial_x - i\partial_y)(\partial_x + i\partial_y)$, d.h. $(\partial_x + i\partial_y)f = 0 \Rightarrow \Delta f = 0$.

5.3 Vektorfelder I

Sei U eine offene Teilmenge $U \subseteq \mathbb{R}^n$. Ein Vektorfeld X auf U ist ein Differentialoperator

$$X = \sum_{\nu=1}^n a_\nu(x) \cdot \partial_\nu$$

mit Koeffizienten $a_\nu(x) \in C^\infty(U)$. Hierbei faßt man die ∂_ν als Differentialoperatoren auf. Das heißt, man kann ein Vektorfeld X auf eine Funktion $f \in C^\infty(U)$ anwenden in der Form $X(f)(x) = \sum_{\nu=1}^n a_\nu(x) \cdot \partial_\nu(f)(x)$, und erhält wieder eine Funktion $X(f) \in C^\infty(U)$. Auf diese Weise definiert ein Vektorfeld X eine $\mathbb{R}$ -lineare Abbildung

$$X : C^\infty(U) \rightarrow C^\infty(U)$$

und **Derivation**, d.h. es gilt $X(\lambda \cdot f + \mu \cdot g) = \lambda \cdot X(f) + \mu \cdot X(g)$ für $\lambda, \mu \in \mathbb{R}$ sowie

$$X(f \cdot g) = X(f) \cdot g + f \cdot X(g) \quad , \quad f, g \in C^\infty(U)$$

was sich ja unmittelbar aus der Produktregel für Ableitungen ergibt.

Sind $X = \sum_{\nu=1}^n a_\nu(x) \cdot \partial_\nu$ und $Y = \sum_{\mu=1}^n b_\mu(x) \cdot \partial_\mu$ Vektorfelder, dann ist der **Kommutator**

$$[X, Y] = X \circ Y - Y \circ X$$

¹d.h. wird von $\partial_x^2 + \partial_y^2$ annulliert; siehe Abschnitt 5.9.

wieder ein Vektorfeld. In der Tat ist a priori $[X, Y]$ ein Differentialoperator zweiter Ordnung, aber die zweiten Ableitungen $\sum_{\nu, \mu} a_\nu(x) b_\mu(x) (\partial_\nu \partial_\mu - \partial_\mu \partial_\nu)$ kürzen sich wegen Satz 4.11 weg. Die genaue Rechnung zeigt $X \circ Y - Y \circ X = \sum_{\nu=1}^n c_\nu(x) \cdot \partial_\nu$ mit den Koeffizienten $c_\nu(x) = \sum_\mu a_\mu(x) \partial_\mu(b_\nu)(x) - b_\mu(x) \partial_\mu(a_\nu)(x)$ in $C^\infty(U)$.

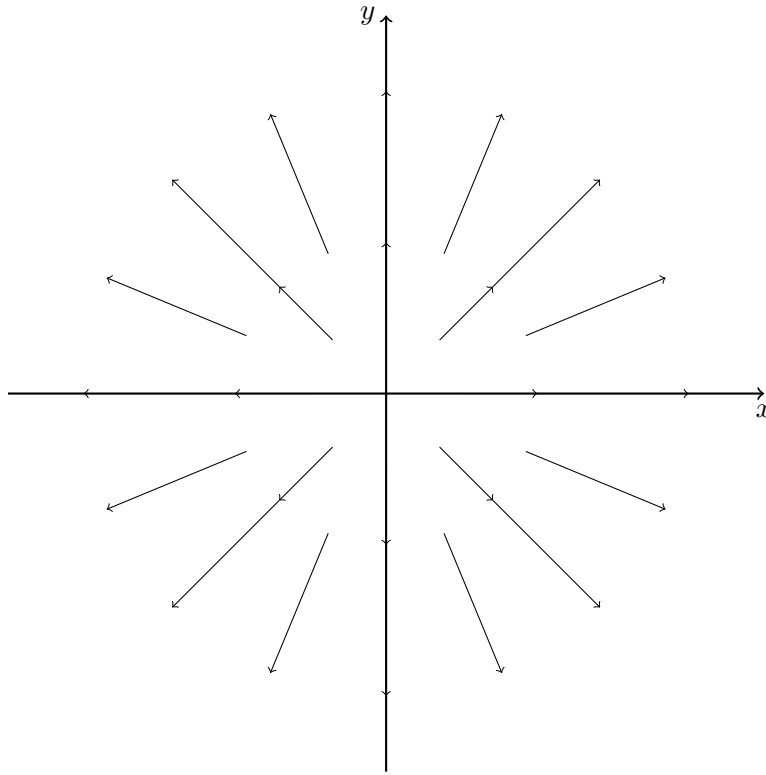
Ein Vektorfeld $X = \sum_{\nu=1}^n a_\nu(x) \partial_\nu$ auf U definiert einen **Kontraktionsoperator**

$$i_X: A^j(U) \rightarrow A^{j-1}(U)$$

vermöge² der Definition $i_X(\omega) = \sum_{\nu=1}^n a_\nu(x) \cdot (\partial_\nu \lrcorner \omega)$. Es folgt damit $i_X(df) = X(f)$.

Ein Vektorfeld $X = \sum_{\nu=1}^n a_\nu(x) \cdot \partial_\nu$ lässt sich visualisieren, indem man an jedem Punkt $\xi \in U$ den Vektor $a(\xi) = (a_1(\xi), \dots, a_n(\xi)) \in \mathbb{R}^n$ anfügt. Die Funktion $a: U \rightarrow \mathbb{R}^n$ definiert ein ‘Feld’ von Vektoren, wenn man $a(x)$ als ‘Tangentialvektor’ im Punkt x visualisiert.

Beispiel. Das **Eulerfeld**³ $E = \sum_{\nu} x_\nu \partial_\nu$ mit $a(x_1, \dots, x_n) = (x_1, \dots, x_n)$. Im Fall $n = 2$



Eine nicht identisch verschwindende Funktion $f(x)$ auf $\mathbb{R}^n \setminus \{0\}$ heisst **homogen** vom Grad α , wenn für alle reellen $t > 0$ gilt

$$f(t \cdot x) = t^\alpha \cdot f(x).$$

² Auf dem antisymmetrischen Polynomring $\mathcal{S}_n$ (siehe Abschnitt 5.14) definiert man die $\mathbb{R}$ -linearen ‘Ableitungen’ $\frac{\partial}{\partial \theta_\nu}: \mathcal{S}_n \rightarrow \mathcal{S}_n$ durch $\frac{\partial}{\partial \theta_\nu}(\theta_\nu \cdot g(\theta)) = g(\theta)$, sowie $\frac{\partial}{\partial \theta_\nu} \theta^I = 0$ für $\nu \notin I$. Ersetzt man formal die dx_J durch $\theta^J \in \mathcal{S}_n$, entspricht $\frac{\partial}{\partial \theta_\nu} \theta^J$ der Form $\partial_\nu \lrcorner dx_J$ (siehe Abschnitt 4.13). In diesem Sinne ist $i_X = \sum_{\nu=1}^n a_\nu(x) \cdot \frac{\partial}{\partial \theta_\nu}$.

³ Der zugehörige Fluß (Abschnitt 5.4) ist $\varphi_t(x) = \exp(t) \cdot x$ wegen $a(x) := \frac{d}{dt} \varphi_t(x)|_{t=0} = \frac{d}{dt} (\exp(t) \cdot x)|_{t=0} = x$. Beachte $\frac{d}{dt} f(\exp(t) \cdot x)|_{t=0} = \sum_{\nu=1}^n x_\nu \frac{\partial}{\partial x_\nu} f(x) = E(f)$ für $f \in C^\infty(U)$.

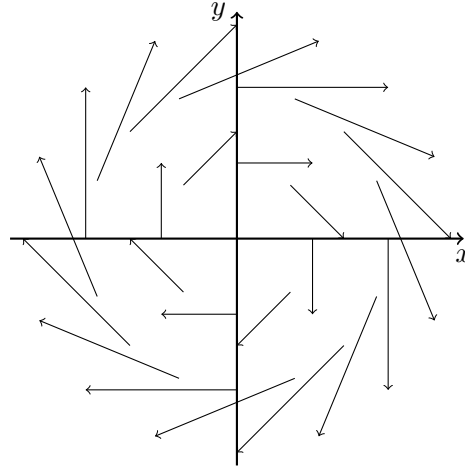
5 Ausgewählte Themen I

Der Grad $\alpha \in \mathbb{R}$ ist dann eindeutig bestimmt. Eine homogene Funktionen ist durch ihre Werte auf der Sphäre X vom Radius $\|x\| = 1$ eindeutig bestimmt. Ist $f(x)$ ein Polynom, dann ist $\alpha = l$ notwendiger Weise eine natürliche Zahl $l = 0, 1, 2, 3, \dots$.

Lemma 5.9. Sei $f(x)$ differenzierbar auf $\mathbb{R}^n \setminus \{0\}$. Dann ist $f(x)$ homogen vom Grad α genau dann, wenn für den **Euler Operator** $E = \sum_{\nu} x_{\nu} \partial_{\nu}$ gilt $E(f) = \alpha \cdot f$.

Beweis. Sei $x \in \mathbb{R}^n$ fixiert. Im Fall $E(f) = \alpha f$ ist $g(t) = f(tx)$ wegen $\alpha f(tx) = E(f)(tx) = \sum_{\nu} x_{\nu} \partial_{\nu} f(tx) = t \sum_{\nu} x_{\nu} (\partial_{\nu} f)(tx) = t \frac{d}{dt} f(tx)$ eine Lösung der Differentialgleichung $\frac{d}{dt} g(t) = \frac{\alpha}{t} \cdot g(t)$ mit $g(1) = f(x)$. Also $g(t) = t^{\alpha} \cdot f(x)$ wegen Satz 4.18. Die Umkehrung ist trivial. $\square$

Ein anderes Beispiel liefert das **Drehfeld**⁴ $L_{21} = -L_{12} = y\partial_x - x\partial_y$ mit $a(x, y) = (y, -x)$



Ist $f(x)$ homogen vom Grad λ , dann sind die partiellen Ableitungen $\partial_{\nu} f$ homogen vom Grad $\lambda - 1$. Die Differentialoperatoren $E = \sum_{\nu=1}^n x_{\nu} \partial_{\nu}$ und $L_{\nu\mu} = x_{\nu} \partial_{\mu} - x_{\mu} \partial_{\nu}$ erhalten daher den Homogenitätsgrad. Beachte $L_{\nu\mu} = -L_{\mu\nu}$. Wir nehmen daher immer $\nu \neq \mu$ an. Dann ist⁵

$$[L_{\alpha\beta}, L_{\beta\gamma}] = L_{\alpha\gamma}$$

für $\alpha \neq \gamma$ und ist Null sonst. Ebenso $[L_{\alpha\beta}, L_{\gamma\delta}] = 0$, falls $\{\alpha, \beta\} \cap \{\gamma, \delta\} = \emptyset$. Also ist $\sum_{\nu < \mu} \mathbb{R} \cdot L_{\nu\mu}$ eine **Lie Algebra**, d.h. abgeschlossen unter Kommutatorbildung.

Lemma 5.10. Der Operator $L^2 := \sum_{\alpha < \beta} (L_{\alpha\beta})^2$ vertauscht mit allen $L_{\nu\mu}$

$$[L^2, L_{\nu\mu}] = 0.$$

⁴Der zugehörige Fluß ist $\varphi_t(x, y) = (\cos(t)x + \sin(t)y, -\sin(t)x + \cos(t)y)$ im Sinne von Abschnitt 5.4 und Lemma 4.22, d.h. es gilt $a(x, y) := \frac{d}{dt} \varphi_t(x, y)|_{t=0} = (y, -x)$.

⁵ $L_{\alpha\beta} L_{\beta\gamma} - L_{\beta\gamma} L_{\alpha\beta} = (x_{\alpha} \partial_{\beta} - x_{\beta} \partial_{\alpha})(x_{\beta} \partial_{\gamma} - x_{\gamma} \partial_{\beta}) - (x_{\beta} \partial_{\gamma} - x_{\gamma} \partial_{\beta})(x_{\alpha} \partial_{\beta} - x_{\beta} \partial_{\alpha}) = x_{\alpha} \partial_{\gamma} - x_{\gamma} \partial_{\alpha} = L_{\alpha\gamma}$.

Beweis. Sei obdA $L_{\nu\mu} = L_{12}$. Dann vertauscht L_{12} mit $L_{\alpha\beta}^2$ ausser wenn genau einer der beiden Indizes α, β in $\{1, 2\}$ liegt. Es gilt $[L_{12}, L_{1\alpha}^2] = (L_{12}L_{1\alpha} - L_{1\alpha}L_{12})L_{1\alpha} + L_{1\alpha}(L_{12}L_{1\alpha} - L_{1\alpha}L_{12}) = [L_{12}, L_{1\alpha}]L_{1\alpha} + L_{1\alpha}[L_{12}, L_{1\alpha}] = -L_{2\alpha}L_{1\alpha} - L_{1\alpha}L_{2\alpha}$ für $\alpha \neq 1, 2$. Analog $[L_{12}, L_{2\alpha}^2] = [L_{12}, L_{2\alpha}]L_{2\alpha} + L_{2\alpha}[L_{12}, L_{2\alpha}] = L_{1\alpha}L_{2\alpha} + L_{2\alpha}L_{1\alpha}$. Aufsummation über alle α gibt Null wegen $[L_{12}, L_{12}^2] = 0$. $\square$

Es gilt $L_{\nu\mu}(r^2) = 0$ für $r^2 = \sum_{j=1}^n x_j^2$; umgekehrt gilt für $n \geq 2$

Jedes Polynom $P = P(x_1, \dots, x_n)$ mit $L_{\nu\mu}(P) = 0$ für alle ν, μ ist ein Polynom in r^2 .

[Für $n > 2$ kann man per Induktion annehmen $P = Q(x_1, \rho^2)$ für $\rho^2 = x_2^2 + \dots + x_n^2$, und damit $P = \sum_{i \geq m} x_1^i P_i(r^2)$ wegen $\rho^2 = r^2 - x_1^2$; oBdA $P_m(r^2) \neq 0$. Durch Subtraktion von $P_0(r^2)$ ist oBdA $m > 0$; man muss dann zeigen $P = 0$. Wäre $P \neq 0$, folgt aus $L_{1i}(P) = 0$ für $i \geq 2$ sofort $x_i P_m(r^2) = 0$ im Widerspruch zu $P_m(r^2) \neq 0$. Der Fall $n = 2$ geht ähnlich (hierbei nimmt man am besten oBdA an, P sei homogen vom Grad ℓ).]

5.4 Vektorfelder II

Sei $U \subset \mathbb{R}^n$ offen. Die C^∞ -Koordinatenwechsel $\varphi : U \rightarrow U$ (siehe 4.26) bilden eine Gruppe unter Komposition. Einen Homomorphismus $t \mapsto \varphi_t$ der additiven Gruppe $\mathbb{R}$ in diese Gruppe nennt man einen **Fluß** auf U ; d.h. es soll gelten $\varphi_0 = id_U$ und die **Funktionalgleichung**

$$\boxed{\varphi_s \circ \varphi_t = \varphi_{s+t}} \quad \text{für alle } s, t \in \mathbb{R}.$$

Wir fordern hierbei stillschweigend $\varphi_t(x) \in C^\infty(\mathbb{R} \times U)$. Damit definiert φ_t ein **zugeordnetes Vektorfeld** $X = X_\varphi$ auf U

$$\boxed{X = \sum_{\nu=1}^n a_\nu(x) \cdot \partial_\nu} \quad \text{für} \quad \boxed{a_\nu(x) := \frac{d\varphi_t(x)_\nu}{dt} \Big|_{t=0}}.$$

Das Vektorfeld X bzw. die Funktion $a : \mathbb{R}^n \rightarrow \mathbb{R}^n$, $a(x) = (a_1(x), \dots, a_n(x))$ legen den Fluß eindeutig fest durch die Anfangsbedingung $\varphi_0(x) = x$ und die Differentialgleichung

$$\boxed{\frac{d}{dt} \varphi_t(x) = a(\varphi_t(x))}.$$

[Nach Definition ist $a(\xi) = \frac{d}{ds} \varphi_s(\xi)|_{s=0}$. Für $\xi = \varphi_t(x)$ ist $\frac{d}{ds} \varphi_s(\xi)|_{s=0} = \frac{d}{ds} \varphi_s(\varphi_t(x))|_{s=0} = \frac{d}{ds} \varphi_{s+t}(x)|_{s=0} = \frac{d}{dt} \varphi_t(x)$ wegen $\varphi_s \circ \varphi_t = \varphi_{s+t}$. Es folgt $\frac{d}{dt} \varphi_t(x) = a(\varphi_t(x))$.]

Lieableitung. Sei $\omega \in A^j(U)$. Für einen Fluß φ_t auf U gilt $\varphi_t^*(\omega) \in A^j(U)$ für alle $t \in \mathbb{R}$. Wir definieren die Lieableitung $L_X : A^j(U) \rightarrow A^j(U)$ durch

$$\boxed{L_X(\omega) := \frac{d}{dt} \varphi_t^*(\omega) \Big|_{t=0}}.$$

Man zeigt durch Ableiten der Funktionalgleichung $\varphi_{s+t} = \varphi_s \circ \varphi_t = \varphi_t \circ \varphi_s$ leicht⁶

$$\boxed{\frac{d}{dt} \varphi_t^*(\omega) = \varphi_t^*(L_X(\omega)) = L_X(\varphi_t^*(\omega))} \quad \forall t \in \mathbb{R}.$$

⁶beachte $\frac{d}{dt} \varphi_t^*(\omega) = \frac{d}{ds} \varphi_{s+t}^*(\omega)|_{s=0} = \frac{d}{ds} (\varphi_t^*(\varphi_s^*(\omega)))|_{s=0} = \varphi_t^*(\frac{d}{ds} \varphi_s^*(\omega)|_{s=0}) = \varphi_t^*(L_X(\omega))$ oder gleich $\frac{d}{dt} \varphi_t^*(\omega) = \frac{d}{ds} \varphi_{s+t}^*(\omega)|_{s=0} = \frac{d}{ds} (\varphi_s^*(\varphi_t^*(\omega)))|_{s=0} = \frac{d}{ds} \varphi_s^*(\varphi_t^*(\omega))|_{s=0} = L_X(\varphi_t^*(\omega))$

5 Ausgewählte Themen I

Die Kettenregel liefert $L_X(f)(x) := \frac{d}{dt}f(\varphi_t(x))|_{t=0} = \sum_{\nu=1}^n (\partial_\nu f)(x) \cdot a_\nu(x) = X(f)(x)$, also die erste der Eigenschaften

1. L_X ist $\mathbb{R}$ -linear, und für $f \in A^0(U)$ gilt $L_X(f) = X(f)$.
2. L_X vertauscht mit der Cartanableitung d .
3. L_X eine **Derivation**: $L_X(\eta \wedge \omega) = L_X(\eta) \wedge \omega + \eta \wedge L_X(\omega)$ für $\eta, \omega \in A^\bullet(U)$.

[Eigenschaft 2) folgt aus der entsprechenden Eigenschaft von φ_t^* . Eigenschaft 3) folgt aus der $\frac{d}{dt}$ -Produktregel wegen $L_X(\eta \wedge \omega) = \frac{d}{dt}\varphi_t^*(\eta \wedge \omega)|_{t=0} = \frac{d}{dt}(\varphi_t^*(\eta) \wedge \varphi_t^*(\omega))|_{t=0} = (\frac{d}{dt}\varphi_t^*(\eta) \wedge \varphi_t^*(\omega))|_{t=0} + (\varphi_t^*(\omega) \wedge \frac{d}{dt}\varphi_t^*(\eta))|_{t=0}$.]

Lemma 5.11 (Cartan Formel). Auf $A^\bullet(U)$ stimmen L_X und $i_X \circ d + d \circ i_X$ überein.

Beweis. Da $A^\bullet(U)$ als Algebra von $C^\infty(U)$ und $d(C^\infty(U))$ erzeugt wird, legen die Eigenschaften 1)-3) die Lieableitung eindeutig fest. Es genügt daher, daß $D := i_X \circ d + d \circ i_X$ die obigen Eigenschaften 1)-3) erfüllt. 1) folgt aus $D(f) = i_X(df) = X(f)$ für $f \in C^\infty(U)$. Wegen $d \circ D = d \circ i_X \circ d = D \circ d$ vertauscht D mit d , und schliesslich ist D eine Derivation⁷. $\square$

Bemerkung 5.12. Die für den Beweis des *Poincaré Lemmas* wichtige *Homotopieformel* $\omega = I(d\omega) + dI(\omega)$ erhält man recht elegant aus Lemma 5.11:⁸

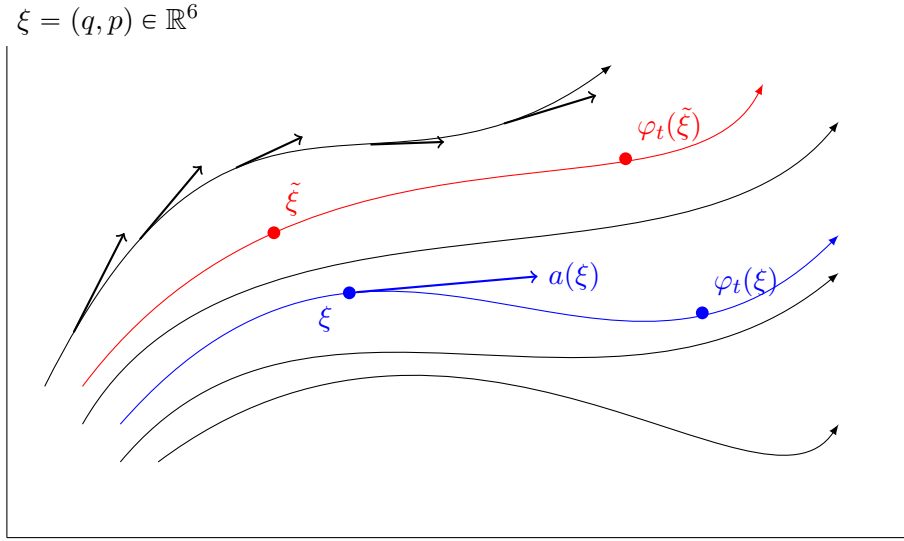
Konstruktion von Flüssen. Sei $K \subset U$ eine kompakte Teilmenge. Ein Fluß φ_t auf U heisst K -lokal, wenn $\varphi(x) = x$ gilt für $x \notin K$. Aus $\varphi_t(x) = x$ für $x \notin K$ folgt $a_\nu(x) = 0$ für $x \notin K$ und alle ν . Damit hat das assoziierte Vektorfeld X Träger in K .

Lemma 5.13. Auf U besteht eine 1-1 Korrespondenz zwischen K -lokalen Flüssen $\varphi_t = \varphi_t^X$ und Vektorfeldern $X = X_\varphi$ mit Träger in K .

Beweis. Für $X = \sum_\nu a_\nu(x)\partial_\nu$ mit Träger in K kann $a(x) = (a_1(x), \dots, a_n(x))$ durch Null zu einer C^∞ -Funktion $a : \mathbb{R}^n \rightarrow \mathbb{R}^n$ fortgesetzt werden. Für festes $\xi \in \mathbb{R}^n$ besitzt $\dot{f}(t) = a(f(t))$ eine eindeutige Lösung $f : \mathbb{R} \rightarrow \mathbb{R}^n$ mit $f(0) = \xi$ nach Satz 4.15 [Die Lipschitzbedingung $\|a(x) - a(y)\|_{\mathbb{R}^n} \leq M\sqrt{n}\|x - y\|_{\mathbb{R}^n}$ für $M = \max_{\theta \in K}(\|Da(\theta)\|) < \infty$ (Satz 2.10) und alle $x, y \in \mathbb{R}^n$ folgt aus dem Mittelwertsatz.] Für $\varphi_t(\xi) := f(t)$ gilt $\varphi_0(\xi) = \xi$ sowie $\varphi_s(\varphi_t(\xi)) = \varphi_{s+t}(\xi)$, denn beide Seiten erfüllen dieselbe Differentialgleichung in s und die Anfangswerte bei $s = 0$ sind gleich wegen $\varphi_0 = id$. Wegen $\varphi_s \circ \varphi_t = \varphi_{s+t}$ ist $\varphi_t = \varphi_{-t}^{-1}$ bijektiv auf $\mathbb{R}^n$. Das φ_t zugeordnete Vektorfeld ist X [wegen $\frac{d}{dt}\varphi_t(\xi)|_{t=0} = f'(0) = a(f(0)) = a(\varphi_0(\xi)) = a(\xi)$]. Also $\varphi_t(\xi) = \xi$ für $\xi \notin K$, d.h. $\varphi_t(K) = K$ für alle $t \in \mathbb{R}$. $\square$

⁷Paritätsändernde $\mathbb{R}$ -lineare Abbildungen C mit $C(\eta \wedge \omega) = C(\eta) \wedge \omega + (-1)^{grad(\eta)} \eta \wedge C(\omega)$ nennt man Superderivationen (auf einem graduerten Ring). Der Superkommutator $D = AB + BA$ von zwei Superderivationen A, B ist eine Derivation, somit gilt $D(\eta \wedge \omega) = D(\eta) \wedge \omega + \eta \wedge D(\omega)$. Beachte, i_X und d sind Superderivationen.

⁸Sei $x_0 = 0$ der Sternmittelpunkt von $U \subseteq \mathbb{R}^n$. Der Fluß $\varphi_t(x) = \exp(-t) \cdot x$ auf $\mathbb{R}^n$ bildet für $t \leq 0$ das sternförmige Gebiet $U \subseteq \mathbb{R}^n$ in sich ab mit $\lim_{t \rightarrow -\infty} \varphi_t(x) = x_0$ für alle $x \in U$. Setze $J(\omega) = \int_{-\infty}^0 \varphi_t^*(\omega)$ sowie $I(\omega) = J(i_X(\omega))$ für das assoziierte Vektorfeld X (das Eulerfeld E). Das Integral ist definiert (!), und $\varphi_0^*(\omega)(x) - \varphi_{-\infty}^*(\omega)(x) = \int_{-\infty}^0 \frac{d}{dt}\varphi_t^*(\omega)(x)dt = \int_{-\infty}^0 \varphi_t^*(L_X(\omega))dt = J \circ (d \circ i_X + i_X \circ d)(\omega)$ nach der Cartan Formel. Wegen $J \circ d = d \circ J$ (Satz 4.32) ist dies $d \circ J \circ i_X + J \circ i_X \circ d = d \circ I + I \circ d$. Beachte $\varphi_0^*(\omega) = \omega$ sowie $\varphi_{-\infty}^*(\omega)(x) = \delta_{j0} \cdot \omega(x_0)$ (!) für $\omega \in A^j(U)$.



Bemerkung 5.14. Das Vektorfeld $Y = (1 + x^2) \frac{\partial}{\partial x}$ auf $\mathbb{R}$ definiert keinen Fluß φ_t^Y [denn die Lösung $f(t) = \frac{\sin(t)}{\cos(t)}$ der Differentialgleichung $f'(t) = a(f(t)) = 1 + f(t)^2$ mit $f(0) = 0$ ‘explodiert’ bei $t = \frac{\pi}{2}$]. Daher konstruieren wir in Abschnitt 5.17 C^∞ -Funktionen $\psi(x) \geq 0$, welche auf einer gegebenen kompakten Teilmenge $K' \subset \mathbb{R}^n$ konstant 1 sind und ausserhalb zu Null abklingen mit kompaktem Träger K in $\mathbb{R}^n$. Ersetzt man ein beliebiges Vektorfeld Y auf $\mathbb{R}^n$ durch $X = \psi(x) \cdot Y$, so hat das Vektorfeld X Träger in der kompakten Menge K ; nach Lemma 5.13 gibt es den zugehörigen Fluß φ_t^X auf $\mathbb{R}^n$. Innerhalb von K' hängt dieser Fluß nur von dem ursprünglichen Vektorfeld Y ab und nicht von der Wahl der Abschneidefunktion ψ . Dies erlaubt es Lieableitungen für beliebige Vektorfelder zu definieren.

5.5 Poissonklammer

Sei $U \subset \mathbb{R}^n$ offen. Eine 2-Form $\omega \in A^2(U)$ heisst **nicht entartet**, wenn für alle $x \in U$ die aus den Koeffizienten $\omega_{\nu\mu}(x)$ von ω gebildete, antisymmetrisch ergänzte reelle $n \times n$ -Matrizen $(\omega_{\nu\mu}(x))$ invertierbar ist für alle $x \in U$

$$\omega = \sum_{\nu < \mu} \omega_{\nu\mu}(x) \cdot dx_\nu \wedge dx_\mu \quad , \quad \boxed{\det(\omega_{\nu\mu}(x)) \neq 0} .$$

Mit Hilfe einer nicht entarteten 2-Form ω kann man Vektorfelder X auf U mit 1-Formen ξ auf U identifizieren: $X = \sum_\nu a_\nu(x) \cdot \frac{\partial}{\partial x_\nu}$ wird $\xi = \sum_\mu a_\nu(x) \omega_{\nu\mu}(x) \cdot dx_\mu$ zugeordnet

$$\xi = i_X(\omega) .$$

Umgekehrt der Form $\xi = \sum_i f_i(x) \cdot dx_i$ in $A^1(U)$ das Vektorfeld $X = \sum_{i,j=1}^n f_i(x) \omega^{ij}(x) \cdot \frac{\partial}{\partial x_j}$; hierbei bezeichne $\omega^{\nu\mu}(x)$ die zu $\omega_{\nu\mu}(x)$ inverse Matrix.

Definition 5.15. Man nennt (U, ω) eine **symplektische Mannigfaltigkeit**, wenn ω eine nicht entartete geschlossene 2-Form ist: $\boxed{d\omega = 0}$.

5 Ausgewählte Themen I

Sei (U, ω) eine symplektische Mannigfaltigkeit. Die **Poisson Klammer** $\{f, g\} = \{f, g\}_\omega$ zweier Funktionen $f, g \in C^\infty(U)$ ist definiert als die C^∞ -Funktion

$$\{f, g\} := X_f(g).$$

Hierbei bezeichne X_f das Vektorfeld, welches bezüglich ω der 1-Form $\xi = df$ zugeordnet ist: Also $\xi = \sum_i \frac{\partial f(x)}{\partial x_i} \cdot dx_i \leftrightarrow X_f$. Es gilt $\{f, g\} = -\{g, f\}$, denn nach Definition von X_f ist

$$\{f, g\}(x) = \sum_{i,j=1}^n \frac{\partial f(x)}{\partial x_i} \omega^{ij}(x) \frac{\partial g(x)}{\partial x_j}.$$

Jacobi Identität. Für $f, g, h \in C^\infty(U)$ gilt

$$\{h, \{f, g\}\} + \{g, \{h, f\}\} + \{f, \{g, h\}\} = 0.$$

Beweis. Für $f_i = \partial_i f$, $g_j = \partial_j g$ gilt $\{h, \{f, g\}\} = \sum_{i,j,k,\nu} h_k \omega^{k\nu} \partial_\nu (f_i \omega^{ij} g_j)$. Summiert über die zyklischen Vertauschungen gibt dies $\sum_{i,j,k,\nu} f_i g_j h_k (\omega^{i\nu} \partial_\nu \omega^{jk} + \omega^{j\nu} \partial_\nu \omega^{ki} + \omega^{k\nu} \partial_\nu \omega^{ij})$ sowie drei weitere Terme vom Typ $\sum_{i,j,k,\nu} f_\nu i g_j h_k (\omega^{k\nu} \omega^{ij} + \omega^{j\nu} \omega^{ki})$ etc. für $f_{\nu i} := \partial_\nu f_i$ etc. Die drei weiteren Terme sind Null [$f_{\nu i}$ ist symmetrisch in den Indizes ν und i]. Wegen $d\omega = 0$ entfällt auch der erste Term [dieser ist invariant unter zyklischen Vertauschungen der Indizes i, j, k . Benutze nun $\partial_\nu \omega^{ij} = -\sum_{\alpha,\beta} \omega^{i\alpha} \partial_\nu (\omega_{\alpha\beta}) \omega^{\beta j}$. Summiert man über die drei möglichen zyklischen Vertauschungen von i, j, k und verwendet $d\omega = 0$, d.h. $\partial_k \omega_{ij} + \partial_j \omega_{ki} + \partial_i \omega_{jk} = 0$, folgt die Behauptung.] $\square$

Die Poisson Klammer macht $C^\infty(U)$ zu einer (**abstrakten**) **Lie Algebra**, und die bereits oben definierte lineare Zuordnung $f \mapsto X_f$ wird zu einem Lie Algebrenhomomorphismus, d.h. es gilt

$$X_{\{f,g\}} = [X_f, X_g]$$

wegen $X_{\{f,g\}} h = \{\{f, g\}, h\} = \{f, \{g, h\}\} + \{g, \{h, f\}\} = X_f(X_g h) - X_g(X_f h) = [X_f, X_g](h)$.

Hamiltonscher Fluß. Sei $h \in C_c^\infty(U)$, $X_h = \sum_{\nu,\mu} \frac{\partial h}{\partial x_\nu}(x) \omega^{\nu\mu}(x) \partial_\mu$ das Vektorfeld $X = X_h$ und $\varphi_t = \varphi_t^X$ ein zugehöriger Fluß. Für $f \in C^\infty(U)$ gilt wegen $\{h, f\} = X_h(f)$ die Gleichung (*)

$$\frac{d}{dt} f(\varphi_t(x)) = \{h, f\}(\varphi_t(x)),$$

[links steht $\frac{d}{dt} \varphi_t^*(f) = \varphi_t^*(L_{X_h}(f))$ and der Stelle x ; beachte $\varphi_t^*(f)(x) = f(\varphi_t(x))$ und $L_{X_h}(f) = X_h(f) = \{h, f\}$ für $f \in A^0(U)$]. Funktionen $f \in C^\infty(U)$, die konstant sind auf allen Flußlinien, nennt man **Invarianten** oder **Integrale** des Hamiltonschen Flusses. Notwendig und hinreichend dafür ist $X_h(f) = 0$ oder $\{f, h\} = 0$. Daher ist h selbst eine Invariante. Summen, Produkte etc. von Invarianten sind wieder Invarianten. Ebenso ist die Poisson Klammer zweier Invarianten wieder eine Invariante (wegen der Jacobi Identität). Schließlich definiert ein Hamiltonscher Fluß **symplektische Koordinatenwechsel** φ_t von (U, ω) (**kanonische Transformationen**), d.h. es gilt

$$\varphi_t^*(\omega) = \omega \quad \forall t \in \mathbb{R},$$

denn $\frac{d}{dt}\varphi_t^*(\omega) = \varphi_t^*(L_{X_h}(\omega)) = 0$ wegen $L_{X_h}(\omega) = i_{X_h}(d\omega) + d(i_{X_h}(\omega)) = d(i_{X_h}(\omega)) = d(dh) = 0$ (Lemma 5.11) unter Benutzung von $d\omega = 0$ sowie $i_{X_h}(\omega) = dh$.

Beispiel. Auf $U = \mathbb{R}^{2n}$ mit den Koordinaten $x = (q_1, \dots, q_n, p_1, \dots, p_n)$ ist die **Liouville 2-Form** $\omega = \sum_{\nu=1}^n dq_\nu \wedge dp_\nu$ nicht entartet mit konstanten Koeffizienten

$$(\omega_{\nu\mu}) = -(\omega^{\nu\mu}) = \begin{pmatrix} 0 & E \\ -E & 0 \end{pmatrix}.$$

Es gilt $\omega = -d\alpha$ für die 1-Form $\alpha = \sum_{\nu=1}^n p_\nu dq_\nu$ (**Wirkungsform**), also $d\omega = 0$. Man nennt die sympl. Mannigfaltigkeit $(\mathbb{R}^{2n}, \omega)$ den **Phasenraum** (das **Kotangentialbündel**) von $\mathbb{R}^n$.

Sei φ_t ein Fluß auf dem Phasenraum zu dem Hamiltonschen Vektorfeld X einer Funktion h . Wegen $X = \sum_{\nu=1}^n (\frac{dh}{dp_\nu}) \cdot \partial_{q_\nu} + \sum_{\mu=1}^n (-\frac{dh}{dq_\mu}) \cdot \partial_{p_\mu}$ gilt für die Koordinatenfunktionen $f(x) = p_i$ resp. $g(x) = q_j$ dann $\{p_i, p_j\} = \{q_i, q_j\} = 0$ sowie $\{p_i, q_j\} = \delta_{ij}$; und die Gleichung (*) verkürzt sich symbolisch zu $\dot{q}_i = \{h, q_i\} = \frac{\partial h}{\partial p_i}$ und $\dot{p}_i = \{h, p_i\} = -\frac{\partial h}{\partial q_i}$.

Die Lieableitung der Wirkungsform α ist $L_X(\alpha) = d\ell \in A^1(\mathbb{R}^{2n})$ für die wie folgt definierte **Legendre Funktion** $\ell = \ell(q, p)$ auf $\mathbb{R}^{2n}$

$$\ell(q, p) := \sum_{j=1}^n p_j \cdot \frac{dh}{dp_j}(q, p) - h(q, p).$$

[Benutze Lemma 5.11 und $i_X(\alpha) = i_X(\sum_j p_j dq_j) = \sum_j p_j \frac{dh}{dp_j}$ sowie $i_X(d\alpha) = -i_X(\omega) = -dh$].

Auf den Flußkurven $f(t) = \varphi_t^X(\xi)$ ist h constant und es gilt $f^*(\alpha)(t) = (\dot{f}(t), \alpha(f(t))) \cdot dt$ (Abschnitt 5.1) und $\dot{f}(t) = a(f(t))$. Also ist $f^*(\alpha)(t) = i_X(\alpha)(f(t)) \cdot dt$, und wegen $i_X(\alpha) = \sum_j p_j \frac{dh}{dp_j}$ ist damit

$$\int_{t_0}^{t_1} \ell(f(t)) dt = (t_0 - t_1) \cdot h(\xi) + \int_{t_0}^{t_1} f^*(\alpha).$$

5.6 Variationsrechnung

Für eine reellwertige C^∞ -Funktion $L = L(q, x, t)$ der Variablen $(q, x, t) \in \mathbb{R}^{2n+1}$

$$L: \mathbb{R}^{2n+1} \rightarrow \mathbb{R},$$

für $t \in \mathbb{R}$ und Vektorvariable $q, x \in \mathbb{R}^n$, definiert man die reellwertige **Hamilton Funktion**

$$H(q, x, t) := \sum_{j=1}^n p_j \cdot x_j - L(q, x, t)$$

mit Hilfe der folgenden reellwertigen Hilfsfunktionen (**den kanonischen Impulsen**)

$$p_j := \frac{dL}{dx_j}(q, x, t) \quad \text{und} \quad p := p(q, x, t) = (p_1, \dots, p_n).$$

Annahme. $\det(\text{Hess}_x(\mathcal{L})) \neq 0$ für $\text{Hess}_x(L) = (\frac{d^2 L(q, x, t)}{dx_i dx_j})_{1 \leq i, j \leq n}$.

5 Ausgewählte Themen I

Dann definiert $\psi: (q, x, t) \mapsto (q, p, t) = (q, p(q, x, t), t)$ einen (lokalen) Koordinatenwechsel φ im $\mathbb{R}^{2n+1}$ nach Satz 4.25. Man kann damit $L = L(q, x, t)$ lokal als Funktion der Variablen (q, p, t) schreiben: $L(q, x, t) = \ell(q, p(q, x, t), t)$ für eine geeignete Funktion $\ell = \ell(q, p, t)$; man erhält für die Vektorvariablen $q, p \in \mathbb{R}^n, t \in \mathbb{R}$

$$h(q, p, t) := \sum_j p_j \cdot x_j(q, p, t) - \ell(q, p, t) = H(q, x, t).$$

Behauptung. Es gilt $\ell(q, p, t) = \sum_{j=1}^n p_j \cdot \frac{dh}{dp_j}(q, p, t) - h(q, p, t)$ wegen⁹

$$x_i = \frac{dh}{dp_i}(q, p, t) \quad ; \text{ ausserdem } \frac{dL}{dq_i}(q, x, t)|_{x=x(q,p,t)} = -\frac{dh}{dq_i}(q, p, t).$$

Variationsrechnung. Hier sucht man C^∞ -Funktionen $f: \mathbb{R} \rightarrow \mathbb{R}^n$ mit der Eigenschaft

$$\left[\frac{d}{dt} \left(\frac{dL}{dx_i}(f(t), \dot{f}(t), t) \right) = \frac{dL}{dq_i}(f(t), \dot{f}(t), t) \right] \quad , \quad i = 1, \dots, n;$$

(siehe Satz 5.16). Für eine gegebene Lösung $f(t)$ definieren wir $\varphi: \mathbb{R} \rightarrow \mathbb{R}^{2n+1}$,

$$t \mapsto \varphi_t = (q(t), p(t), t) \in \mathbb{R}^{2n+1},$$

durch $q(t) = q(\varphi_t) := f(t)$, $x(t) := \dot{f}(t)$ und $p(t) = p(\varphi_t) := p(q(t), x(t), t) = p(f(t), \dot{f}(t), t)$. Beachte $q(t), x(t), p(t)$ definieren Funktionen vom Typ $\mathbb{R} \rightarrow \mathbb{R}^n$. Es gilt $\dot{q}_i(t) = \dot{f}_i(t) = x_i(t)$ und $\dot{p}_i(t) = \frac{d}{dt}(p_i(t)) := \frac{d}{dt} \left(\frac{dL}{dx_i}(f(t), \dot{f}(t), t) \right) = \frac{dL}{dq_i}(f(t), \dot{f}(t), t)$, also wegen obiger Behauptung

$$\left[\frac{d}{dt}(q_i(\varphi_t)) = \frac{dh}{dp_i}(\varphi_t) \quad , \quad \frac{d}{dt}(p_i(\varphi_t)) = -\frac{dh}{dq_i}(\varphi_t) \right]$$

oder kurz $\dot{q}_i = \frac{\partial h}{\partial p_i}$ und $\dot{p}_i = -\frac{\partial h}{\partial q_i}$ im Punkt φ_t . Somit definiert φ_t einen **Hamiltonschen Fluß** im **Phasenraum** zu der **Hamilton Funktion** h , zumindestens wenn h nicht von t abhängt.

Sei $L(q, x, t): \mathbb{R}^{2n} \times [a, b] \rightarrow \mathbb{R}$ eine C^∞ -Funktion mit $t \in [a, b]$ sowie $q, x \in \mathbb{R}^n$. Sei $f: [a, b] \rightarrow \mathbb{R}^n$ eine C^2 -differenzierbare Abbildung, die das folgende Integral $\mathbb{L}(f)$ minimiert

$$\mathbb{L}(f) := \int_a^b L(f(t), \dot{f}(t), t) dt.$$

Satz 5.16 (Euler-Lagrange). Unter den obigen Voraussetzungen gilt für $i = 1, \dots, n$

$$\left[\frac{d}{dt} \left(\frac{dL}{dx_i}(f(t), \dot{f}(t), t) \right) = \frac{dL}{dq_i}(f(t), \dot{f}(t), t) \right].$$

⁹Beachte $\frac{dh}{dp_i}(q, p, t) = x_i + \sum_j \frac{dx_j}{dp_i} p_j - \sum_j \frac{dL}{dq_j}(q, x, t) \frac{dq_j}{dp_i} - \sum_j \frac{dL}{dx_j}(q, x, t) \frac{dx_j}{dp_i} = x_i$ wegen $\frac{dt}{dq_j} = \frac{dq_j}{dp_i} = 0$ und $\sum_j \frac{dL}{dx_j}(q, p, t) \frac{dx_j}{dp_i} = \sum_j p_j \frac{dx_j}{dp_i}$. Weiterhin ist $\frac{dh}{dq_i}(q, p, t) = \sum_j \frac{dx_j}{dq_i} p_j - \frac{dL}{dq_i}(q, x, t) - \sum_j \frac{dL}{dx_j}(q, x, t) \frac{dx_j}{dq_i}$ wegen $\frac{dp_j}{dq_i} = 0$ und $\frac{dt}{dq_j} = 0$. Benutze nun $\frac{dL}{dx_j}(q, x, t) \frac{dx_j}{dq_i} = p_j \frac{dx_j}{dq_i}$.

Beweis. Für jede feste differenzierbare Kurve $g: [a, b] \rightarrow \mathbb{R}^n$ mit $g(a) = g(b) = 0$ besitzt die in der Variable $s \in \mathbb{R}$ nach Satz 4.32 differenzierbare reellwertige Funktion $\phi(s) := \mathbb{L}(f + s \cdot g)$ im Punkt $s = 0$ ein Minimum. Aus Lemma 4.9 folgt $\frac{d\phi}{ds}(0) = 0$. Die Differentiation nach s kann nach Satz 4.32 mit dem Integral vertauscht werden. Dies gibt

$$\frac{d\phi}{ds}(0) = \int_a^b \frac{d}{ds} L(f(t) + s \cdot g(t), \dot{f}(t) + s\dot{g}(t), t)|_{s=0} dt = 0$$

oder wegen der Kettenregel $\int_a^b \left(\frac{d}{dq} L(f(t), \dot{f}(t), t) \cdot g(t) + \frac{d}{dx} L(f(t), \dot{f}(t), t) \cdot \dot{g}(t) \right) dt = 0$; der einfacheren Schreibweise zuliebe ist hier oBdA $n = 1$. Durch partielle Integration folgt wegen $g(a) = g(b) = 0$ dann (*)

$$\int_a^b h(t) \cdot g(t) dt = 0$$

für $h(t) = \frac{d}{dq} L(f(t), \dot{f}(t), t) - \frac{d}{dt} \left(\frac{d}{dx} L(f(t), \dot{f}(t), t) \right)$. Dies gilt für alle Funktionen $g \in C_c^\infty(U)$ auf $U = (a, b)$. Somit verschwindet die stetige Funktion $h(t)$ in jedem Punkt $t = t_0 \in U$, und damit auch auf $[a, b]$, da man anderenfalls wegen Abschnitt 5.17 eine Funktion $g \geq 0$ mit $g(t_0) = 1$ und Träger in einem beliebig kleinen Intervall $(-\varepsilon + t_0, t_0 + \varepsilon)$ finden könnte, für die das Integral (*) nicht verschwindet. $\square$

5.7 Satz von Darboux

Sei (M, ω) eine symplektische Mannigfaltigkeit, d.h. M ist eine offene Teilmenge von $\mathbb{R}^n$ und ω eine nicht degenerierte geschlossene 2-Form auf M . Für jedes $\xi \in M$ definiert ω eine alternierende $m \times m$ -Matrix $J = (\omega_{\nu\mu}(\xi))$ mit $\det(J) \neq 0$. Nach einem Satz von Frobenius (Lineare Algebra) gilt dann $m = 2n$, und für fixiertes ξ existiert eine lineare Abbildung $A \in Gl(2n, \mathbb{R})$ mit der Eigenschaft

$${}^T A J A = \begin{pmatrix} 0 & E \\ -E & 0 \end{pmatrix}.$$

Durch geeignete affin lineare Koordinatenwechsel kann man daher (für lokale Betrachtungen in einem festen Punkt $\xi \in M$) ohne Einschränkung annehmen: $\xi = 0$ und $\omega(\xi) = \omega_L(\xi)$. Für die gegebene symplektische Mannigfaltigkeit (M, ω) gilt überraschenderweise

Satz 5.17. Für jedes $\xi \in M$ existiert eine offene Teilmenge $V \subset M$, die ξ enthält, sowie eine offene Teilmenge $U \subseteq \mathbb{R}^{2n}$ und ein lokaler Koordinatenwechsel $\varphi: U \rightarrow V \subseteq M$ mit

$$\varphi^*(\omega) = \omega_L$$

für die **Liouville Form** ω_L auf U . D.h., in einer Umgebung V von ξ ist (M, ω) als symplektische Mannigfaltigkeit isomorph zu einer offenen Teilmenge U des Phasenraums.

Beweis. OBdA $\omega(\xi) = \omega_L(\xi)$. Für $\omega_t = t \cdot \omega + (1 - t)\omega_L$ folgt $\omega_t(\xi) = \omega_L(\xi)$ für alle t . Es gibt daher ein $r > 0$ mit $\det(\omega_t)(x) \neq 0$ für alle $t \in [0, 1]$ und alle $x \in M$ mit $\|x - \xi\| \leq r$. [Anderenfalls gäbe es eine Folge (x_n, t_n) mit $t_n \in [0, 1]$, $\det(\omega_{t_n}(x_n)) = 0$ und $\|x_n - \xi\| \rightarrow 0$.

5 Ausgewählte Themen I

Wegen Kompaktheit würde dann $t_n \rightarrow t$ gelten für eine Teilfolge im Widerspruch zu $1 = \det(\omega_t)(\xi) = \lim_n \det(\omega_{t_n})(x_n) = 0$. Wir ersetzen dann M durch eine offene (sternförmige) Kugel in M vom Radius $< r$ um ξ , und nennen diese wieder M .

Für festes $t \in [0, 1]$ gilt $d(\omega_L - \omega) = 0$. Nach Satz 4.30 existiert eine 1-Form $\alpha \in A^1(M)$ mit $d\alpha = \omega_L - \omega$. Für festes t bestimmt dies eindeutig Vektorfelder X_t mit

$$i_{X_t}(\omega_t) = \alpha.$$

Man sieht durch direkte Inspektion, dass die Koeffizienten $a_\nu(t, x)$ von $X_t = \sum_\nu a_\nu(t, x) \cdot \partial_\nu$ als Funktion von $(x, t) \in M \times [0, 1]$ dann C^∞ -Funktionen definieren. Schneidet man diese Vektorfelder mit einer geeigneten Funktion $\psi \in C_c^\infty(M)$ räumlich ab (siehe 5.14), welche identisch 1 ist auf einer Umgebung von ξ , ist die Differentialgleichung

$$\frac{d}{dt}\varphi_t(\eta) = a(t, \varphi_t(\eta))$$

mit Anfangsbedingung $\varphi_0(\eta) = \eta$ ähnlich wie im zeitunabhängigen Fall 5.13 lösbar auf $[0, 1]$ für jeden fest fixierten Anfangspunkt $\eta \in M$. [Hierbei bezeichne $a(t, x) = (a_1(t, x), \dots, a_m(t, x))$ bereits die Koeffizienten der abgeschnittenen Vektorfelder.] Die Lösung φ_t definiert zwar keinen Fluß mehr mit Funktionalgleichung, dennoch gilt immer noch¹⁰

$$\frac{d}{dt}(\varphi_t^*(\omega_t)) = \left(\frac{d}{dt}\varphi_t^*\right)(\omega_t) + \varphi_t^*\left(\frac{d}{dt}\omega_t\right) = \varphi_t^*(L_{X_t}(\omega_t)) + \varphi_t^*(\omega - \omega_L).$$

Aus der *Cartan Formel* $L_{X_t}(\omega_t) = i_{X_t}(d\omega_t) + d(i_{X_t}(\omega_t)) = d(i_{X_t}(\omega_t)) = d\alpha = \omega_L - \omega$ folgt daher $\frac{d}{dt}(\varphi_t^*(\omega_t)) = 0$ und damit $\varphi_t^*(\omega_t) = \varphi_0(\omega_0)$. Wegen $\omega_0 = \omega_L$ und $\omega_1 = \omega$ zeigt dies $\varphi^*(\omega) = \omega_L$ für $\varphi := \varphi_1$. $\square$

5.8 Kanonische Transformationen

Definiere Abbildungen $p = p(q, P)$ und $Q = Q(q, P)$ für $p, q, P \in \mathbb{R}^n$

$$(\mathbb{R}^{2n}, \omega) \longleftarrow \mathbb{R}^{2n} \longrightarrow (\mathbb{R}^{2n}, \omega)$$

$$(q, p) \leftarrow (q, P) \rightarrow (Q, P)$$

mit Hilfe einer *Hilfsfunktion* $W = W(q, P) \in C^\infty(\mathbb{R}^{2n})$ vermöge der Formeln

$$p_i := \partial_{q_i} W(q, P) \quad \text{und} \quad Q_j := \partial_{P_j} W(q, P).$$

Ist $\det(\partial W(q, P)/\partial q_j \partial P_i)(x_0) \neq 0$ für $x_0 \in \mathbb{R}^{2n}$, folgt aus Satz 4.25 die Existenz einer offenen Teilmenge $U \subseteq \mathbb{R}^{2n}$ mit $x_0 \in U$, so daß $f(q, P) := (q, p(q, P))$ und $g(q, P) := (Q(p, P), P)$ lokale Koordinatenwechsel definieren, $f : U \cong f(U) \subset \mathbb{R}^{2n}$ und $g : U \cong g(U) \subset \mathbb{R}^{2n}$.

¹⁰Es gilt $\frac{d}{dt}\varphi_t^*(\omega) = \varphi_t^*(L_{X_t}(\omega))$ für alle $\omega \in A^\bullet(M)$. Benutze: Beide Seiten definieren $\mathbb{R}$ -lineare Derivationen auf $A^\bullet(M)$ und vertauschen mit d . Reduziere damit auf den Fall $\omega = f \in A^0(M)$. Benutze dann die Kettenregel $\frac{d}{dt}f(\varphi_t(x)) = \sum_\nu (\partial_\nu f)(\varphi_t(x))\dot{\varphi}_t(x)_\nu = \sum_\nu \dot{a}_\nu(t, x)(\partial_\nu f)(\varphi_t(x)) = X_t(f)(\varphi_t(x)) = \varphi_t^*(X_t(f))(x)$.

Lemma 5.18. $\varphi := g \circ f^{-1} : (q, p) \mapsto (Q, P)$ definiert eine kanonische Transformation zwischen den offenen Teilmengen $f(U)$ und $g(U)$ des Phasenraums $(\mathbb{R}^{2n}, \omega)$. D.h. $\varphi^*(\omega) = \omega$

$$\sum_i dq_i \wedge dp_i = \sum_{i,j} \frac{W(q,P)}{\partial q_j \partial P_i} \cdot dq_i \wedge dP_j = \sum_i dQ_i \wedge dP_i .$$

Beweis. Es gilt $\sum_i dq_i \wedge dp_i = \sum_i dq_i \wedge d(\partial_{q_i} W(q, P)) = \sum_{i,j} \partial_{P_j} \partial_{q_i} W(q, P) dq_i \wedge dP_j + \sum_{i,j} \partial_{q_j} \partial_{q_i} W(q, P) dq_i \wedge dq_j = \sum_{i,j} \partial_{P_j} \partial_{q_i} W(q, P) dq_i \wedge dP_j$. Analog gilt $\sum_i dQ_i \wedge dP_i = \sum_i d(\partial_{P_i} W(q, P)) \wedge dP_i = \sum_{i,j} \partial_{P_j} \partial_{P_i} W(q, P) dP_j \wedge dP_i + \sum_{i,j} \partial_{q_j} \partial_{P_i} W(q, P) dq_j \wedge dP_i = \sum_{i,j} \partial_{q_j} \partial_{P_i} W(q, P) dq_j \wedge dP_i$. Durch Umbenennung von i und j folgt die Behauptung. $\square$

Sei $\varphi : U \rightarrow V$ ein lokaler Koordinatenwechsel und $\omega \in A^2(V)$, $d\omega = 0$ eine nicht entartete 2-Form. Für $f, g \in C^\infty(V)$ gilt dann $\varphi^*(\{f, g\}_\omega) = \{\varphi^*(f), \varphi^*(g)\}_{\varphi^*(\omega)}$.

Findet man für eine gegebene Funktion $h = h(q, p)$ eine Funktion $W = W(q, P)$ mit einer einfacheren Funktion $h(q, P) := h(q, \frac{\partial}{\partial q} W)$, dann wird die Bestimmung des Hamiltonschen Flußes von h sich in den Koordinaten (q, P) , und gegebenenfalls dann auch Q, P , vereinfachen.

5.9 Harmonische Funktionen

Der Laplace Operator

$$\Delta = \sum_{i=1}^n \partial_i^2$$

bildet homogene Funktionen vom Grad α auf homogene Funktionen vom Grad $\alpha - 2$ ab und es gilt (Übungsaufgabe und obdA $n = 2$)

$$[\Delta, L_{\nu\mu}] = 0 .$$

Für $r^2 = \sum_{i=1}^n x_i^2$ erhält der Operator $r^2 \Delta$ den Grad α und es gilt $[r^2 \Delta, E] = 0$ für $E = \sum_i x_i \partial_i$.

Beispiel. Für $r \neq 0$ und $\alpha \in \mathbb{R}$ gilt auf $\mathbb{R}^n \setminus \{0\}$ die Formel

$$\Delta r^\alpha = \alpha(\alpha + n - 2) \cdot r^{\alpha-2} .$$

[Aus $2r \partial_i(r) = 2x_i$ folgt $\partial_i(r) = \frac{x_i}{r}$. Deshalb gilt $\Delta(r^\alpha) = \alpha \sum_i \partial_i(x_i r^{\alpha-2}) = n\alpha r^{\alpha-2} + \alpha(\alpha - 2)r^2 r^{\alpha-4}$. Analog zeigt man $\Delta \log(r) = (n - 2)r^{-2}$ und $\Delta(r^{2-n} \log(r)) = (2 - n)r^{-n}$.

Definition. Sei $U \subset \mathbb{R}^n$ offen. Eine Funktion $f \in C^\infty(U)$ heisst **harmonisch**, wenn für den Laplace Operator Δ gilt

$$\Delta f = 0 .$$

Konstante Funktionen und lineare Polynome sind harmonische Funktionen. Wie wir oben gezeigt haben ist die homogene Funktion

$$f(x) = \frac{1}{r^{n-2}} , \quad \kappa := n - 2$$

5 Ausgewählte Themen I

auf $U = \mathbb{R}^n \setminus \{0\}$ harmonisch. Diese Funktion besitzt eine Singularität¹¹ im Ursprung in den Dimensionen $n \neq 2$. Der Fall $n = 2$ ist exzeptionell. Man hat hier nur die singuläre harmonische Funktion $f(x) = \log(r)$ auf $\mathbb{R}^2 \setminus \{0\}$, welche aber nicht homogen vom Grad Null ist: Es gilt $f(tx) = \log(t) + f(x)$ mit einer Konstante $\log(t)$.

Inversion am Kreis. Sei $f^0(x) := f(\frac{x}{\|x\|^2})$ oder $f^0(x_1, \dots, x_n) := f(\frac{x_1}{r^2}, \dots, \frac{x_n}{r^2})$ für Funktionen f auf $\mathbb{R}^n \setminus \{0\}$. Ist f homogen vom Grad α , dann ist f^0 homogen vom Grad $-\alpha$.

Sei nun $U = \mathbb{R}^n \setminus \{0\}$, und $\kappa := n - 2$ und $n \geq 1$. Wir betrachten die **Kelvin Transformation** (Achtung: f^* nicht verwechseln mit dem Pullback)

$$f^*(x) := r^{-\kappa} f\left(\frac{x}{\|x\|^2}\right).$$

Dann ist $f^*(x) \in C^2(U)$ für $f \in C^2(U)$, und es gilt $f^{**} = f$. Ist f homogen vom Grad α , dann ist f^* homogen vom Grad $-\kappa - \alpha$.

Lemma 5.19 (Kelvin). Für $n \neq 2$ und $r^2 = \sum_{i=1}^n x_i^2$ gilt

$$r^2 \Delta(f^*(x)) = (r^2 \Delta f)^*(x).$$

Insbesondere ist $f^*(x)$ harmonisch, wenn $f(x)$ harmonisch ist.

Beweis. Der Beweis beruht auf einer expliziten Rechnung¹² und benutzt die Fußnote¹³ $\square$

¹¹Im Fall $n = 3$ beschreibt diese Funktion bis auf eine Konstante das Coulomb Potential.

¹²Für $f^0(x) := f(\frac{x}{\|x\|^2})$ gilt $\Delta(\frac{f^0}{r^\kappa}) = \Delta(\frac{1}{r^\kappa})f + 2 \sum_i \partial_i(r^{-\kappa}) \partial_i(f^0) + r^{-\kappa} \Delta(f^0)$ mit $\Delta(\frac{1}{r^\kappa}) = 0$. Wegen Formel 1 und 4 ist der Term $2 \sum_i \partial_i(r^{-\kappa}) \partial_i(f^0)$ gleich

$$-\frac{2\kappa}{r^{\kappa+2}} \sum_i \sum_j x_i (f_j)^0 T_{ij} = \frac{2\kappa}{r^{\kappa+4}} \sum_j x_j (f_j)^0.$$

Wegen Formel 1 ist $r^{-\kappa} \Delta(f^0)$ gleich $\frac{1}{r^\kappa} \sum_i \sum_j \partial_i((f_j)^0 T_{ij})$, und aus Formel 2 und 3 folgt daher die Behauptung $\Delta(f^*) = r^{-4}(\Delta(f))^*$ vermöge

$$\frac{1}{r^\kappa} \sum_i \sum_j \sum_k (f_{kj})^0 T_{ik} T_{ij} + \frac{1}{r^\kappa} \sum_j \sum_i (f_j)^0 \partial_i(T_{ij}) = \frac{(\Delta f)^0}{r^{\kappa+4}} - \frac{2\kappa}{r^{\kappa+4}} \sum_j x_j (f_j)^0.$$

¹³Formeln. Beachte $\partial_i(r^\alpha) = \alpha x_i r^{\alpha-2}$. Für $T_{ij} := \partial_j(x_i r^{-2}) = (\delta_{ij} r^2 - 2x_i x_j) r^{-4}$ und $f_j := \partial_j f$ gilt

1. $\partial_i(f^0) = \sum_j (f_j)^0 T_{ij}$ (Kettenregel)
2. $\sum_i T_{ij} T_{ik} = \delta_{jk} r^{-4}$
3. $\sum_i \partial_i(T_{ij}) = -2\kappa x_j r^{-4}$
4. $\sum_i x_i T_{ij} = -x_j r^{-2}$.

wegen $\sum_i (\delta_{ij} r^2 - 2x_i x_j)(\delta_{ik} r^2 - 2x_i x_k) = \delta_{jk} r^4 - 2x_j x_k r^2 - 2x_j x_k r^2 + 4r^2 x_j x_k = \delta_{jk} r^4$ und $\sum_i \partial_i(T_{ij}) = \partial_j(1/r^2) - 2n x_j / r^4 - 2E(x_j / r^4) = (-2 - 2n - 2(-3))x_j / r^4 = -2\kappa x_j / r^4$.

5.10 Harmonische Polynome

Sei $\mathcal{P}_l = \mathcal{P}_l(\mathbb{R}^n)$ der $\mathbb{R}$ -Vektorraum aller Polynome $P(x)$ in n Variablen, welche homogen vom Grad l sind. Sei $\mathcal{H}_l = \mathcal{H}_l(\mathbb{R}^n) \subset \mathcal{P}_l(\mathbb{R}^n)$ der Unterraum aller **harmonischen Polynome**.

Satz 5.20. $\dim_{\mathbb{R}}(\mathcal{H}_l) = \binom{n+l-1}{l} - \binom{n+l-3}{l-2}$ und für $n \geq 3$ auch $\dim_{\mathbb{R}}(\mathcal{H}_l) = \frac{n-2+2l}{n-2} \binom{n+l-3}{l}$.

Beispiel. Im Fall $n = 1$ ist $\mathcal{H}_l = 0$ für $l > 1$. Im Fall $n = 2$ gilt $\dim_{\mathbb{R}}(\mathcal{H}_l) = 2$ für $l \geq 1$, und $\dim_{\mathbb{R}}(\mathcal{H}_0) = 1$. Als Vektorraum wird $\mathcal{H}_l$ von $\operatorname{Re}(z^l)$ und $\operatorname{Im}(z^l)$ aufgespannt ($z = x_1 + ix_2$).

Beispiel. Im Fall $n = 3$ folgt $\dim_{\mathbb{R}}(\mathcal{H}_l) = 2l + 1$ für alle $l \geq 0$.

Lemma 5.21. $\dim_{\mathbb{R}}(\mathcal{P}_l) = \binom{n+l-1}{l}$.

Beweis. Induktion nach der Variablenzahl n . Ersetzt man die Variable x_n durch 1, ergibt dies genau die Polynome in $x_1, \dots, x_{n-1}$ vom Grad $\leq l$. Also

$$\dim_{\mathbb{R}}(\mathcal{P}_l(\mathbb{R}^n)) - \dim_{\mathbb{R}}(\mathcal{P}_{l-1}(\mathbb{R}^n)) = \dim_{\mathbb{R}}(\mathcal{P}_l(\mathbb{R}^{n-1})),$$

wie die Rekursionsformeln $\binom{n+l-1}{l} - \binom{n+(l-1)-1}{(l-1)} = \binom{(n-1)+l-1}{l}$ der **Binomialkoeffizienten**. $\square$

Beispiel. Im Fall $n = 1$ ist $\mathcal{P}_l = \mathbb{R} \cdot x^l$. Im Fall $n = 2$ ist $\mathcal{P}_l = \mathbb{R}x^l + \mathbb{R}x^{l-1}y + \dots + \mathbb{R}y^l$ von der Dimension $l + 1$, und für $n = 3$ gilt $\dim_{\mathbb{R}}(\mathcal{P}_l) = (l + 1)(l + 2)/2$.

Beweis von Satz 5.20. Wir entwickeln $P \in \mathcal{H}_l(\mathbb{R}^n)$ nach der letzten Variable

$$P(x_1, \dots, x_n) = Q_l(x_1, \dots, x_{n-1}) + Q_{l-1}(x_1, \dots, x_{n-1})x_n + \dots$$

Sei $\Delta = \Delta_{n-1} + \partial_n^2$. Wegen $(\Delta_{n-1} + \partial_n^2)P = 0$ bestimmen die Polynome $Q_l \in \mathcal{P}_l(\mathbb{R}^{n-1})$ und $Q_{l-1} \in \mathcal{P}_{l-1}(\mathbb{R}^{n-1})$ das harmonische Polynom P eindeutig, und Q_l und Q_{l-1} vom Grad l resp. $l - 1$ können dabei beliebig vorgegeben (!) werden. Also $\dim_{\mathbb{R}}(\mathcal{H}_l(\mathbb{R}^n)) = \dim_{\mathbb{R}}(\mathcal{P}_l(\mathbb{R}^{n-1})) + \dim_{\mathbb{R}}(\mathcal{P}_{l-1}(\mathbb{R}^{n-1})) = \binom{n-1+l-1}{l} + \binom{n-1+l-2}{l-1} = \frac{n-2+2l}{n-2} \binom{n+l-3}{l}$ für $n \geq 3$. $\square$

Ein Polynom $f(x) = f(x_1, \dots, x_n)$ definiert $f(\partial) := f(\partial_1, \dots, \partial_n)$ als Differentialoperator auf $C^\infty(\mathbb{R}^n)$ (wohldefiniert wegen der Symmetrie der Hessematrix). Die Bilinearform $\mathcal{P}_l \times \mathcal{P}_l \rightarrow \mathbb{R}$, definiert durch

$$\langle f, g \rangle = f(\partial)g(x),$$

ist $\mathbb{R}$ -bilinear und wegen $\langle \prod_{\nu=1}^n x_\nu^{m_\nu}, \prod_{\nu=1}^n x_\nu^{k_\nu} \rangle = \prod_{\nu=1}^n (m_\nu! \delta_{m_\nu, k_\nu})$ symmetrisch positiv definit. Sei nun $r^2 = x_1^2 + \dots + x_n^2$.

Lemma 5.22. Bezüglich obiger Paarung $\langle \cdot, \cdot \rangle$ existiert eine orthogonale Zerlegung

$$\mathcal{P}_l = r^2 \cdot \mathcal{P}_{l-2} \oplus^\perp \mathcal{H}_l, \quad \mathcal{P}_l = \bigoplus_{i \geq 0} r^{2i} \cdot \mathcal{H}_{l-2i}.$$

Die Teilräume $r^2 \cdot \mathcal{P}_{l-2}$ und $\mathcal{H}_l$ (resp. $r^{2i} \cdot \mathcal{H}_{l-2i}$) sind invariant unter den Operatoren $L_{\nu\mu}$.

5 Ausgewählte Themen I

Beweis. Nach Lemma 5.21 und Satz 5.20 gilt $\dim_{\mathbb{R}}(\mathcal{P}_l) = \dim_{\mathbb{R}}(r^2\mathcal{P}_{l-2}) + \dim_{\mathbb{R}}(\mathcal{H}_l)$. Es genügt also für $g(x) = r^2 \cdot f(x)$ im Durchschnitt $r^2\mathcal{P}_{l-2} \cap \mathcal{H}_l$ zu zeigen: $g(x) = 0$. Dies folgt wegen $\Delta g(x) = 0$ aus $\langle g(x), g(x) \rangle = \langle r^2 f(x), g(x) \rangle = \Delta f(\partial)g(x) = f(\partial)\Delta g(x) = \langle f(x), \Delta g(x) \rangle = 0$.

Wegen $L_{\nu\mu}(r^2) = x_\nu \cdot 2x_\mu - x_\mu \cdot 2x_\nu = 0$ gilt $L_{\nu\mu}(r^2 f(x)) = L_{\nu\mu}(r^2)f(x) + r^2 L_{\nu\mu}(f(x)) = r^2 L_{\nu\mu}(f(x))$, also $L_{\nu\mu}(r^2\mathcal{P}_{l-2}) \subseteq r^2\mathcal{P}_{l-2}$. Aus $\Delta L_{\nu\mu}g = L_{\nu\mu}\Delta g = 0$ für $g \in \mathcal{H}_l$ folgt $L_{\nu\mu}(\mathcal{H}_l) \subseteq \mathcal{H}_l$. $\square$

Sei $n \geq 2$. Das Bild des Monoms $(x_1)^l = x_1^l \in \mathcal{P}_l(\mathbb{R}^n)$ unter der harmonischen Projektion

$$\mathcal{P}_l(\mathbb{R}) = r^2 \cdot \mathcal{P}_{l-2}(\mathbb{R}^n) \oplus \mathcal{H}_l(\mathbb{R}^n) \xrightarrow{pr} \mathcal{H}_l(\mathbb{R}^n)$$

(die orthogonale Projektion des letzten Lemmas) ist das **zonal sphärische Polynom**

$$P_{l,0}(x) := pr(x_1^l).$$

Das Monom x_1^l , und damit auch $P_{l,0}(x)$, ist invariant unter orthogonalen Substitutionen, welche den Vektor $(1, 0, \dots, 0)$ fest lassen. $P_{l,0}(x)$ ist nicht trivial und wird von den $L_{\nu\mu}$ mit $\nu \neq 1, \mu \neq 1$ annulliert, und ist folglich (siehe Abschnitt 5.3) als Polynom

$$P_{l,0}(x) = \sum_{i=0}^l a_i \cdot x_1^i r^{l-i}$$

nur abhängig von den Variablen x_1 und r^2 mit $P_{l,0}(-x) = (-1)^l P_{l,0}(x)$.

Lemma 5.23 (Rodrigues Formel). Für $r = 1$ und $n \geq 2$ gilt für eine Konstante¹⁴ $const(l, n)$

$$const(l, n) \cdot P_{l,0}(x) = (t^2 - 1)^{-\frac{n-3}{2}} \left(\frac{d}{dt} \right)^l (t^2 - 1)^{l+\frac{n-3}{2}} \Big|_{t=x_1}.$$

Beweis. Nach Satz 4.18 ist der Lösungsraum der Gleichungen

$$(t^2 - 1)f''(t) + (n - 1)tf'(t) - l(l + n - 2)f(t) = 0$$

und $f(-t) = (-1)^l f(t)$ eindimensional (Anfangsbedingungen im Punkt Null!). Die Funktion $f(t) = \sum_{i=0}^l a_i t^i$ erfüllt diese Bedingungen¹⁵ und die rechte Seite im Lemma ebenfalls (letzteres ist sehr diffizil). $\square$

Lemma 5.24. Der Operator L^2 ist auf $\mathcal{H}_l(\mathbb{R}^n)$ ein Vielfaches der Identität

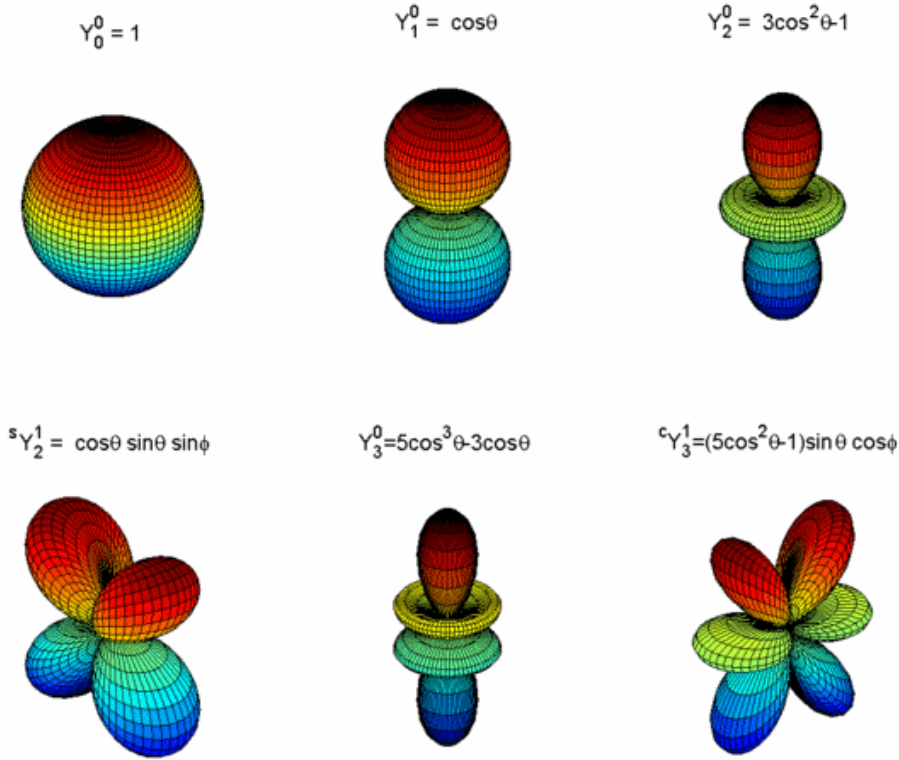
$$L^2 = -l(l + n - 2) \cdot id_{\mathcal{H}_l(\mathbb{R}^n)}.$$

¹⁴Es gilt $const(l, n) = l!(2l+n-3)$.

¹⁵Wir schreiben t oder x anstatt x_1 . Dann ist $\Delta(x^i r^{l-i}) = r^{l-i} \partial_x^2(x^i) + 2\partial_x(x^i) \partial_x(r^{l-i}) + x^i \Delta(r^{l-i})$. Die Summe von $x^i \Delta(r^{l-i}) = (l-i)(n-2+l-i)x^i r^{l-i-2}$ und $2\partial_x(x^i) \partial_x(r^{l-i}) = 2i(l-i)x^i r^{l-i-2}$ ist $-(i^2 + (n-2)i - l(l+n-2)) \cdot x^i r^{l-i-2}$. Also annulliert $\partial_t^2 - (t\partial_t)^2 - (n-2) \cdot t\partial_t + l(l+n-2)$ das Polynom $f(t)$.

Beweis. Der Raum der Polynome $P \in \mathcal{P}_l(\mathbb{R}^n)$, in denen die letzte Variable x_n höchstens linear vorkommt, bildet ein Komplement $V \cong \mathcal{P}_l(\mathbb{R}^{n-1}) \oplus x_n \cdot \mathcal{P}_{l-1}(\mathbb{R}^{n-1})$ von $W = r^2 \cdot \mathcal{P}_{l-2}(\mathbb{R}^n)$ in $\mathcal{P}_l(\mathbb{R}^n)$. [$V \cap W$ ist Null und $\dim(V) = \binom{n+l-2}{l} + \binom{n+l-3}{l-1} = \binom{n+l-1}{l} - \binom{n+l-3}{l-2} = \dim_{\mathbb{R}}(\mathcal{H}_l)$.] Also $\mathcal{H}_l = pr(V)$ nach Lemma 5.22. Aber jedes Monom in V lässt sich linear kombinieren durch Polynome, die durch sukzessives Anwenden von Operatoren $L_{\nu\mu}$ auf das Monom $P(x) = (x_1)^l$ entstehen (Übungsaufgabe). Wendet man daher sukzessive die Operatoren $L_{\nu\mu}$ auf die Funktion $P_{l,0}$ an, erhält man ein Erzeugendensystem von $\mathcal{H}_l$.

Für $1 < \alpha < \beta$ gilt $L_{\alpha\beta}x_1^l = 0$. Wegen $L_{1\alpha}^2(x_1^l) = (x_1\partial_\alpha - x_\alpha\partial_1)(-x_\alpha l x_1^{l-1}) = -l x_1^l + x_\alpha^2 l(l-1)x_1^{l-2}$ liefert Summation über $\alpha \geq 2$ daher $L^2(x_1^l) = -l(n-1)x_1^l - l(l-1)x_1^l + l(l-1)r^2 x_1^{l-2} = -l(l+n-2)x_1^l + l(l-1)r^2 x_1^{l-2}$. Da L^2 mit der harmonischen Projektion vertauscht, folgt aus $pr(l(l-1)r^2 x_1^{l-2}) = 0$ dann $L^2(P_{l,0}) = -l(l+n-2)P_{l,0}$ für $P_{l,0} = pr(x_1^l)$. Da $\mathcal{H}_l$ aus $P_{l,0}$ durch sukzessives Anwenden der $L_{\nu\mu}$ erzeugt wird, die Operatoren $L_{\nu\mu}$ aber mit L^2 vertauschen, folgt die Behauptung. $\square$



5.11 Drehimpuls Operatoren

Im dreidimensionalen Fall $n = 3$ benutzen wir die Abkürzungen $L_1 := L_{32} = z\partial_y - y\partial_z$, $L_2 := L_{13} = x\partial_z - z\partial_x$ und $L_3 := L_{21} = y\partial_x - x\partial_y$. Dann gilt

$$[L_1, L_2] = L_3 \quad , \quad [L_3, L_1] = L_2 \quad , \quad [L_2, L_3] = L_1 .$$

5 Ausgewählte Themen I

Die Operatoren L_1, L_2, L_3 vertauschen mit dem Differentialoperator $L^2 = L_1^2 + L_2^2 + L_3^2$ und es gilt $L^2 = -l(l+1)$ auf $\mathcal{H}_l$ wegen Lemma 5.24.

Aus $L_1 x^l = 0$ und Lemma 5.22 folgt $L_1 P_{l,0} = 0$. Die zonale Eigenfunktion $P_{l,0}$ von L_1 liefert neue Eigenfunktionen $P_{l,k}$ von L_1 . Setze dazu $L_{\pm} = L_2 \mp i \cdot L_3$ für $i = \sqrt{-1} \in \mathbb{C}$. Sei $P_{l,k} := (L_+)^k P_{l,0}$ für $k > 0$ bzw. $P_{l,-k} = \overline{P_{l,k}}$ das konjugiert komplexe. Für Eigenvektoren $L_1 v = m i \cdot v$ gilt wegen $[L_1, L_{\pm}] = \pm i L_{\pm}$ dann $L_1(L_{\pm} v) = L_{\pm}(L_1 v) \pm i L_{\pm} v = (m \pm 1) i \cdot L_{\pm} v$. Also hat $L_{\pm} v$ den L_1 -Eigenwert $(m \pm 1) \cdot i$, wenn $L_{\pm} v$ nicht verschwindet. Dies zeigt

$$L_1 P_{l,k} = k i \cdot P_{l,k} \quad (k \in \mathbb{Z}).$$

Wegen $L^2 = (L_1)^2 + (L_2)^2 + (L_3)^2 = -(i L_1)^2 + (i L_1) + L_- L_+$ folgt aus $L_+(P) = 0$ für einen Eigenvektor $P \in \mathcal{H}_l$ von L_1 zum Eigenwert $k i$ dann $L^2(P) = -k(k+1) \cdot P$. Andererseits ist L^2 gleich $-l(l+1) \cdot id$ auf $\mathcal{H}_l$. Aus $P_{l,k} \neq 0$ und $P_{l,k+1} = L_+ P_{l,k} = 0$ für $k \geq 0$ folgt daher $l(l+1) = k(k+1)$, und damit $k = l$. Also sind die Funktionen $P_{l,k}$ für $k = -l, \dots, 0, \dots, l$ von Null verschieden.

Korollar 5.25. Die $2l+1$ linear unabhängigen **Kugelflächenfunktionen** $P_{l,k}$ für $k = -l, \dots, -1, 0, 1, \dots, l$ definieren wegen $\dim(\mathcal{H}_l(\mathbb{R}^3)) = 2l+1$ eine $\mathbb{C}$ -Basis des Vektorraums $\mathcal{H}_l(\mathbb{R}^3)$ der $\mathbb{C}$ -wertigen harmonischen Polynome, bestehend aus Eigenvektoren von L_1 und L^2 .

In der Physik sind L_1, L_2, L_3 (bis auf Normierung) die **Drehimpuls Operatoren**. In der Quantenmechanik treten sie in anderer Form auf und $l = 0, 1, 2, \dots$ ist der **Spin**. In der Elektrodynamik spielen die Funktionen $P_{l,k}$ eine Rolle bei Radialentwicklungen von Monopolen ($l = 0$), Dipolen ($l = 1$), ... etc. Die Einschränkungen der $P_{l,k}(x, y, z)$ auf die Einheitssphäre im $\mathbb{R}^3$ heissen **Kugelflächenfunktionen** vom Grad l . Man faßt diese oft auf als homogene Funktionen $Y(x, y, z) = P(x, y, z)/r^l$ vom Grad Null (für $P \in \mathcal{H}_l(\mathbb{R}^3)$). Sie treten auf bei der **Radialentwicklung** von elektrostatischen Potentialen $u(x, y, z)$ auf Kugelschalen (siehe Abschnitt 10.5)

$$u(x, y, z) = \sum_{l=0}^{\infty} \sum_{k=-l}^l (a_{lk} r^l + b_{lk} r^{-l-1}) \cdot Y_{l,k}(x, y, z).$$

Hierbei sind $a_{lk}, b_{lk} \in \mathbb{C}$ geeignete Konstanten. Die Polynome $r^l Y_{l,k} = P_{l,k}$ und ihre **Kelvin Transformierten** $r^{-l-1} Y_{l,k} = (P_{l,k})^*$ sind harmonisch. Die Funktionen $r^{-l-1} Y_{l,k} = (P_{l,k})^*$ erhält man auch durch partielles Ableiten der homogenen harmonischen Funktion $\frac{1}{r}$. Diese sind singular bei 0 und klingen im Unendlichen ab im Gegensatz zu den Polynomen $P_{l,k}$. Dies lässt sich (ebenso wie Korollar 5.25) auf Dimensionen $n \geq 3$ verallgemeinern (siehe Kapitel 10).

5.12 Taylor Koeffizienten

Sei $x_0 = 0$ und U eine offene Kugel um $x_0 \in \mathbb{R}^n$ sowie $f \in C^r(U)$. Für $l \leq r$ ist dann der l -te **Taylor Koeffizient** $T_l(f)(x)$ von f im Punkt $x_0 = 0$ (für 'kleine' $t \in \mathbb{R}$ bei festem $x \in \mathbb{R}^n$) definiert durch

$$T_l(f)(x) := \frac{\partial_t^l}{l!} f(tx)_{t=0}.$$

Die so definierte Funktion ist ein homogenes Polynom in x vom Grad l , dessen Koeffizienten bis auf universelle Konstanten¹⁶ partielle Ableitungen von f im Punkt x_0 sind. Dies zeigt man leicht mit Hilfe der Kettenregel und Induktion nach l .

Beispiel 1. Ist $f(x)$ ein homogenes Polynom vom Grad m , d.h. gilt $f(tx) = t^m f(x)$, dann ist $T_l(f)(x) = 0$ für $l \neq m$ und $T_l(f)(x) = f(x)$ für $l = m$.

Beispiel 2. Aus Beispiel 1 folgt für die Funktion $g(x) = f(x) - \sum_{l=0}^r T_l(f)(x)$ sofort $T_l(g)(x) = 0$ für $l = 0, \dots, r$ für beliebiges $f \in C^r(U)$.

Beispiel 3. Ist $f(x) = \sum_{l=0}^{\infty} P_l(x)$ eine Potenzreihe mit homogenen Polynomen $P_l(x)$ vom Grad l und Konvergenzradius $R > 0$. Dann ist $f(tx) = \sum_{l=0}^{\infty} P_l(x)t^l$ für $|t| < 1$, und aus Lemma 4.39 folgt $T_l(f)(x) = P_l(x)$.

Beispiel 4. Ist $f(x)$ eine harmonische Funktion, dann sind die $T_l(f)(x)$ harmonische Polynome, denn es gilt $\Delta T_l(f)(x) = \Delta \frac{\partial_t^m}{m!} f(tx)_{t=0} = \frac{\partial_t^m}{m!} t^2 (\Delta f)(tx)_{t=0} = 0$. (Analog zeigt man: Ist $f(z)$ holomorph, dann sind alle $T_l(f)(z)$ holomorph).

Lemma 5.26. Für $f \in C^r(U)$ gibt es eine stetige Funktion $H : U \rightarrow \mathbb{R}$ mit $H(0) = 0$ und

$$f(x) = \sum_{l=0}^r T_l(f)(x) + \|x\|^r \cdot H(x).$$

Beweis. Wir benutzen Induktion nach r . Der Fall $r = 0$ ist klar. Für $f \in C^r(U)$, $r \geq 1$ ist $f(x) - f(0) = f(tx)|_0^1 = \int_0^1 \frac{d}{dt} f(tx) dx = \sum_{\nu=1}^n x_{\nu} f_{\nu}(x)$ (nach Bemerkung 5.2) mit $f_{\nu} = \int_0^1 (\partial_{\nu} f)(tx) dt$ in $C^{r-1}(U)$ (Satz 4.32). Nach Induktion ist $f_{\nu} = \sum_{l=0}^{r-1} T_l(f_{\nu})(x) + \|x\|^{r-1} \cdot H_{\nu}(x)$, also $f(x) = \sum_{l=0}^r P_l(x) + \|x\|^r H(x)$ für $H(x) = \sum_{\nu} \text{sign}(x_{\nu}) H_{\nu}(x)$ und homogene Polynome $P_l(x)$ vom Grad l . Es gilt $\|x\|^r H(x) \in C^r(U)$ und aus $\lim_{x \rightarrow 0} H_{\nu}(x) = 0$ folgt $\lim_{x \rightarrow 0} H(x) = 0$. Aus $T_l(f)(x) = P_l(x) + T_l(\|x\|^r H(x))$ und¹⁷ $T_l(\|x\|^r H(x)) = 0$ für $l = 0, \dots, r$ folgt $P_l(x) = T_l(f)(x)$. $\square$

5.13 Orthogonale Gruppen

Für eine reelle symmetrische invertierbare $n \times n$ -Matrix S bilden die Matrizen $X \in M_{n,n}(\mathbb{R})$ mit der Eigenschaft¹⁸ (${}^T X$ bezeichne die transponierte Matrix von X)

$${}^T X = -SXS^{-1}$$

¹⁶Durch Reduktion auf $f(x) = \prod_i x_i^{m_i}$ folgt dann $T_l(f)(x) = \sum_{m_1+\dots+m_n=l} \left(\frac{\partial_1^{m_1}}{m_1!} \dots \frac{\partial_n^{m_n}}{m_n!} f \right)(0) \cdot \prod_{i=1}^n x_i^{m_i}$.

¹⁷Aus $\lim_{t \rightarrow 0} h(t) = 0$ und $t^l h(t) \in C^l$ folgt $\partial_t^l (t^l h(t))|_{t=0} = 0$. Benutze Induktion nach l und $\partial_t^l (t^l h(t)) = \partial_t^{l-1} \partial_t (t^{l-1} t h(t)) = \partial_t^{l-1} t^{l-1} [(l-1)h(t) + \tilde{h}(t)]$ für $\tilde{h}(t) = \partial_t (th(t))$, um die Aussage auf $l = 1$ zu reduzieren.

Für $l = 1$ folgt die Aussage aus $th(t) = o(t)$.

¹⁸Beachte ${}^T [X, Y] = {}^T (XY - YX) = {}^T Y^T X - {}^T X^T Y = (-SY S^{-1})(-SXS^{-1}) - (-SXS^{-1})(-SY S^{-1}) = -S(XY - YX)S^{-1} = -S[X, Y]S^{-1}$ für $X, Y \in \mathfrak{so}(S)$.

5 Ausgewählte Themen I

die **orthogonale Lie Algebra** $so(S)$. Die zugehörige **orthogonale Gruppe** $O(S, \mathbb{R})$ besteht aus allen reellen $r \times r$ -Matrizen M mit der Eigenschaft

$$\boxed{^T M S M = S}.$$

Beachte $M, N \in O(S, \mathbb{R}) \Rightarrow M \circ N \in O(S, \mathbb{R})$ und $M^{-1} \in O(S, \mathbb{R})$. Für M in $O(S, \mathbb{R})$ gilt $\det(M)^2 = 1$ und die Matrizen $M \in O(S, \mathbb{R})$ mit der Eigenschaft $\det(M) = 1$ definieren eine Untergruppe $SO(S, \mathbb{R})$, die **spezielle orthogonale Gruppe** zur **quadratischen Form** $q_S(x) = ^T x S x$. Eine Matrix $M \in Gl(n, \mathbb{R})$ liegt genau dann in $O(S, \mathbb{R})$, wenn gilt

$$q_S(M(x)) = q_S(x) \quad , \quad \forall x \in \mathbb{R}^n, .$$

[Es gilt $q_S(M(x)) = ^T (Mx) S (Mx) = ^T x (^T M S M) x = ^T x S x = q_S(x)$ für $M \in O(S, \mathbb{R})$. Wegen $2^T y S x = q_S(x + y) - q_S(x) - q_S(y)$ folgt umgekehrt aus $q_S(M(v)) = q_S(v)$ für alle v sofort $^T y (^T M S M) x = ^T y S x$. Wählt man für x, y die Standardbasisvektoren, ergibt sich $^T M S M = S$]. Wie in Lemma 4.43 zeigt man für reelle Matrizen $X \in M_{n,n}(\mathbb{R})$

$$X \in so(S) \iff \exp(tX) \in O(S, \mathbb{R}) \quad , \quad \forall t \in \mathbb{R} .$$

Sei $S = \text{diag}(\lambda_1, \dots, \lambda_n)$ mit $\lambda_\nu \in \{\pm 1\}$. Wir schreiben dann $x^\nu := \lambda_\nu x_\nu$ und $\partial^\nu := \lambda_\nu^{-1} \partial_\nu$. Die quadratische Form $q(x) = \sum_{\nu=1}^n \lambda_\nu x_\nu^2$ schreibt sich dadurch kurz $q(x) = \sum_{\nu=1}^n x_\nu x^\nu$.

Ist S die Einheitsmatrix $S = E$, schreibt man $so(n)$ anstatt $so(E)$, und $so(n)$ besteht dann aus den antisymmetrischen $n \times n$ -Matrizen. Eine Basis des $\mathbb{R}$ -Vektorraums $so(n)$ bilden die Matrizen $E_{\nu\mu}$ für $1 \leq \nu < \mu \leq n$, die den Eintrag $+1$ bei (ν, μ) und den Eintrag -1 bei (μ, ν) haben und sonst nur Nulleinträge. Man zeigt leicht $[E_{\alpha\beta}, E_{\beta\gamma}] = E_{\alpha\gamma}$. Also kann $so(n)$ mit der Lie Algebra der **Drehfelder** $L_{\nu\mu} = x_\nu \partial_\mu - x_\mu \partial_\nu$ (siehe Abschnitt 5.3) identifiziert werden

$$so(n) \cong \bigoplus_{\nu < \mu} \mathbb{R} \cdot L_{\nu\mu} ,$$

da die $L_{\nu\mu}$ dieselben Kommutator-Relationen erfüllen wie die $E_{\nu\mu}$. Im allgemeinen wird $so(S)$ von den Operatoren $L_\nu^\mu = x_\nu \partial^\mu - x_\mu \partial^\nu$ erzeugt. Wichtige Spezialfälle sind die Gruppen $O(r, s)$ und die Lie Algebren $so(r, s) := so(S)$ für $S = \text{diag}(+1, \dots, +1, -1, \dots, -1)$ mit r mal $+1$ und $s = n - r$ mal -1 als Eintrag. Im Fall $r = 3, s = 1$ erhält man die **Lorentzgruppe** und ihre Lie Algebra $so(3, 1)$. Die affin linearen Abbildungen $f = f(M, b)$ der Gestalt

$$f(x) = M \cdot x + b \quad , \quad M \in O(3, 1) \quad , \quad b \in \mathbb{R}^4$$

$$f(M_1, b_1) \circ f(M_2, b_2) = f(M_1 M_2, M_1(b_2) + b_1)$$

bilden eine Gruppe, die **Poincaregruppe**. Die Liealgebra der Poincaregruppe ist die direkte Summe von $so(3, 1)$ und der Liegruppe der Translationen $\bigoplus_{i=1}^4 \mathbb{R} \cdot \partial_i$

$$so(3, 1) \oplus \bigoplus_{i=1}^4 \mathbb{R} \cdot \partial_i .$$

5.14 Fourier-Graßmann Transformation

Sei $\mathcal{S}_n$ der ‘Polynomring’ in n antikommutierenden Variablen $\theta_1, \dots, \theta_n$. Elemente $f(\theta) \in \mathcal{S}_n$ schreiben sich wegen $\theta_i \theta_j = -\theta_j \theta_i$ (somit $\theta_i^2 = 0$) auf eindeutige Weise in der Form

$$f(\theta) = \sum_I c_I \cdot \theta^I \quad , \quad \theta^I := \theta_{i_1} \cdots \theta_{i_r}$$

für $I = \{i_1, \dots, i_r\} \subset \{1, \dots, n\}$ mit $i_1 < \dots < i_r$ und Koeffizienten c_I in einem Körper $K \subset \mathbb{C}$. Es gibt eine eindeutig bestimmte K -lineare Involution des Ringes $\mathcal{S}_n$ mit den Eigenschaften

$$(f \cdot g)^* = g^* \cdot f^* \quad \text{mit} \quad (\theta_i)^* = \theta_i \quad .$$

Man zeigt dann leicht durch Umordnen $(\theta^I)^* = \varepsilon_r \cdot \theta^I$ für $\varepsilon_r := (-1)^{r(r-1)/2}$ und $r = |I|$.

Berezin Integral. Jedes $f(\theta) \in \mathcal{S}_n$ ist ein Produkt $f(\theta) = f_n(\theta_n) \cdots f_1(\theta_1)$ mit $f_i(\theta_i) = a_i + \theta_i b_i$ und Konstanten a_i, b_i in K . Wir erklären formal

$$\int f(\theta) d\theta := \int (\cdots (\int f(\theta) d\theta_1) \cdots) d\theta_n := \prod_{i=1}^n b_i \quad ,$$

etwa $\int \theta_1 \cdots \theta_n d\theta = \varepsilon_n$. Für $f(\theta) = q(\theta_2, \dots, \theta_n) + r(\theta_2, \dots, \theta_n) \theta_1$ gilt $\int f(\theta) d\theta_1 = r(\theta_2, \dots, \theta_n)$.

Koordinatenwechsel. Für $U \in Gl(n, K)$ und $f(\theta) \in \mathcal{S}_n$ ist $(U^* f)(\theta) = f(U\theta)$ der Pullback in $\mathcal{S}_n$; d.h. man ersetzt jede Variable θ_ν durch $\sum_{\mu=1}^n U_{\nu\mu} \theta_\mu$. Der Pullback vertauscht mit der Involution von $\mathcal{S}_n$, d.h. es gilt $U^*(f^*) = (U^* f)^*$. Die Leibnizformel 4.12 liefert folgende ‘Substitutionsregel’ für das Berezin Integral.

Lemma 5.27. *Es gilt $\int f(U\theta) d\theta = \det(U) \cdot \int f(\theta) d\theta$.*

Achtung: Der Vorfaktor ist nicht $\det(U)^{-1}$ wie man zuerst glauben würde.

Hodge *-Operator. Sei $*\theta^I = \varepsilon_I \cdot \theta^{I^c}$ mit $\varepsilon_I \in \{\pm 1\}$ definiert durch $\theta^I \cdot *\theta^I = \theta_1 \cdots \theta_n$. Hierbei bezeichne I^c die Komplementärmenge von I in $\{1, \dots, n\}$. Ist $q_L(x) = \sum_{\nu=1}^n \lambda_\nu \cdot x_\nu^2$ eine quadratische Form¹⁹ mit $\lambda_\nu \neq 0$ für alle ν , dann definiert man allgemeiner $*_L \theta^I = \pm \theta^{I^c}$ mittels $\theta^I \cdot *_L \theta^I = \sqrt{\prod_{\nu=1}^n |\lambda_\nu|} (\prod_{\nu \in I} \lambda_\nu)^{-1} \cdot \theta_1 \cdots \theta_n$. Setze $\varepsilon_L := \prod_{\nu=1}^n (\lambda_\nu / |\lambda_\nu|)$.

Fourier Transformation. Für weitere (auch mit den θ_i) antikommutierende Variable $\eta_1, \dots, \eta_n$ und $\eta \cdot \theta = \sum_{\nu=1}^n \lambda_\nu \eta_\nu \theta_\nu$ ist $\exp(\eta \cdot \theta) = \sum_{m=0}^n (\eta \cdot \theta)^m / m!$ wohldefiniert und kommutiert mit allen $f \in \mathcal{S}_n$. Man erklärt die Fourier Transformierte²⁰ $\mathcal{F}_L f$ von f durch

$$(\mathcal{F}_L f)(\eta) = \left| \prod_{\nu=1}^n \lambda_\nu \right|^{-1/2} \int f(\theta) \cdot e^{\eta \cdot \theta} d\theta \quad .$$

Wir ersetzen danach oft η_i durch θ_i , und erhalten auf diese Weise eine Abbildung $\mathcal{F}_L : \mathcal{S}_n \rightarrow \mathcal{S}_n$.

Beispiel $n = 1$. Für $f(\theta) = a_1 + b_1 \theta_1$ ist $\mathcal{F} f(\theta) = b_1 + a_1 \theta_1$.

¹⁹Diese definiert eine verallgemeinerte Metrik mit $g^{\nu\mu} = \delta_{\nu\mu} \cdot \lambda_\nu$ für $S = \text{diag}(\lambda_1, \dots, \lambda_n)$. Für $U \in O(S, K)$ gilt $(U^* \eta \cdot U^* \theta) = \sum_{n,\nu,m} U_{\nu n} \eta_n \lambda_\nu U_{\nu m} \theta_m = (\eta \cdot \theta)$ für $(\eta \cdot \theta) = \sum_\nu \eta_\nu \lambda_\nu \theta_\nu$.

²⁰Wir schreiben $\mathcal{F}$ im ‘Euklidischen’ Fall $\lambda_1 = \dots = \lambda_n = 1$, also $(\mathcal{F} f)(\eta) = \int f(\theta) \cdot e^{\eta \cdot \theta} d\theta$.

Lemma 5.28. Es gilt $\mathcal{F}_L(f^*) = \varepsilon_L \cdot *_L f$ für den Hodge *-Operator $f \mapsto *_L f$.

Beweis. OBdA $\lambda_\nu = 1$ für $\nu = 1, \dots, n$. Für $f(\theta) = \theta^I$ mit $|I| = r$ ist $\mathcal{F}(f^*) = \varepsilon_r \mathcal{F}(\theta^I)$ gleich $\varepsilon_r \int \frac{(\eta \cdot \theta)^{n-r}}{(n-r)!} \cdot \theta^I d\theta = \varepsilon_r \int (\eta_{j_1} \theta_{j_1}) \cdots (\eta_{j_{n-r}} \theta_{j_{n-r}}) \cdot \theta^I d\theta$ für $j_1 < \cdots < j_{n-r}$ und $I^c = \{j_1, \dots, j_{n-r}\}$. Dies gibt $\varepsilon_r \varepsilon_{n-r} \int \eta^{I^c} \theta^{I^c} \theta^I d\theta = \varepsilon_r \varepsilon_{n-r} (-1)^{r(n-r)} \varepsilon_I \varepsilon_n \eta^{I^c} = \varepsilon_I \eta^{I^c} = *_L f$ wie behauptet, wegen der trivialen binomischen Formel $\varepsilon_r \varepsilon_{n-r} \varepsilon_n (-1)^{r(n-r)} = 1$. $\square$

Lemma 5.29. q_L -orthogonale Koordinatenwechsel U vertauschen mit dem Hodge Operator $*_L$ bis auf das Vorzeichen $\det(U)$, d.h. für $U \in O(S, K)$ und $S = \text{diag}(\lambda_1, \dots, \lambda_n)$ gilt

$$\mathcal{F}_L(U^* f) = \det(U) \cdot U^*(\mathcal{F}_L f).$$

Beweis. Es gilt $(\mathcal{F}_L(U^* f))(\eta) = c_L \int f(U\theta) \exp(\eta \cdot \theta) d\theta = c_L \int f(U\theta) \exp(U\eta \cdot U\theta) d\theta$. Die Substitutionsregel liefert dann $\det(U) c_L \int f(\theta) \exp(U\eta \cdot \theta) d\theta = \det(U) \cdot (\mathcal{F}_L f)(U\eta)$. $\square$

Bemerkung. $**\eta^I = (-1)^{r(n-r)} \varepsilon_I^2 \eta^I$ für $r = |I|$ sowie $\varepsilon_I^2 = 1$ und Lemma 5.28 implizieren $\mathcal{F}^* \mathcal{F}^* = (-1)^{r(n-r)} \cdot \text{id}$ auf Elementen vom Grad r . Auf solchen Elementen ist $\mathcal{F}^* \mathcal{F}^*$ gleich $\varepsilon_{n-r} \varepsilon_r \mathcal{F}^2$, da $\mathcal{F}$ homogene Elemente vom Grad r in homogene Elemente vom Grad $n-r$ abbildet. Es folgt $\mathcal{F}^2 = \varepsilon_{n-r} \varepsilon_r (-1)^{r(n-r)} \cdot \text{id}$ und damit $\boxed{\mathcal{F}^2 = \varepsilon_n \cdot \text{id}}$. Wir erhalten analog

Satz 5.30. $\boxed{f = \varepsilon_L \varepsilon_n \cdot \mathcal{F}_L(g)}$ für $\boxed{g = \mathcal{F}_L(f)}$.

5.15 Laplace Operatoren

Sei $U \subseteq \mathbb{R}^n$ offen. Formen $\omega = \sum_I \omega_I(x) \cdot dx_I$ in $A^r(U)$ können als C^∞ -Funktionen $\omega : U \rightarrow \mathcal{S}_n$ der Gestalt $\sum_I \omega_I(x) \cdot \theta^I$ mit Werten in $\mathcal{S}_n$ (siehe Abschnitt 5.14) aufgefasst werden. Die Cartan Ableitung $d : A^r(U) \rightarrow A^{r+1}(U)$ ist dann $d = \sum_{\nu=1}^n \frac{\partial}{\partial x_\nu} \theta_\nu$. Für gegebenes $q_L(x) = \sum_{\nu=1}^n \lambda_\nu \cdot x_\nu^2$ definiert die Fourier Transformation $\mathcal{F}_L : \mathcal{S}_n \rightarrow \mathcal{S}_n$ die assoziierte **Koableitung**

$$\boxed{\delta_L = (\mathcal{F}_L)^{-1} \circ d \circ \mathcal{F}_L}.$$

Offensichtlich ist $\delta_L \circ \delta_L = 0$, und in der Sprache der Differentialformen gilt für alle r

$$\delta_L : A^r(U) \rightarrow A^{r-1}(U)$$

wegen $\mathcal{F}_L : A^{n-r+1}(U) \rightarrow A^{r-1}(U)$ sowie $\mathcal{F}_L : A^r(U) \rightarrow A^{n-r}(U)$ nach Lemma 5.28. Für $D_L := d + \delta_L$ beachte $D_L^2 : A^r(U) \rightarrow A^r(U)$, d.h. D_L^2 erhält den Raum der r -Formen, wegen

$$\boxed{D_L^2 = (d + \delta_L)^2 = d \circ \delta_L + \delta_L \circ d}.$$

Lemma 5.31. Man hat $D_L^2 \left(\sum_I \omega_I(x) \cdot dx_I \right) = \sum_I \Delta_L(\omega_I(x)) \cdot dx_I$, d.h. es gilt

$$\boxed{D_L^2 = \Delta_L}$$

für den klassischen Laplace Operator $\Delta_L = \sum_{\nu=1}^n \lambda_\nu^{-1} \frac{\partial^2}{\partial x_\nu^2}$.

Beweis. D_L^2 ist $\sum_{\nu=1}^n \sum_{\mu=1}^n \frac{\partial^2}{\partial x_\nu \partial x_\mu} (\theta_\nu \circ (\mathcal{F}_L)^{-1} \circ \theta_\mu \circ \mathcal{F}_L + (\mathcal{F}_L)^{-1} \circ \theta_\nu \circ \mathcal{F}_L \circ \theta_\mu)$. Wegen $(\mathcal{F}_L)^{-1} \circ \theta_\nu \circ \mathcal{F}_L = \lambda_\nu^{-1} \frac{\partial}{\partial \theta_\nu}$ reduziert²¹ sich die Aussage wegen der Symmetrie der Hessematrix auf die Antikommutator-Identität $\theta_\nu \circ \frac{\partial}{\partial \theta_\mu} + \frac{\partial}{\partial \theta_\mu} \circ \theta_\nu = \delta_{\nu\mu}$ (Fußnote IV.18). $\square$

Der **Hodge-Operator** $d^* := (-1)^r \cdot (*_L)^{-1} \circ d \circ (*_L)$ auf $A^r(U)$, mit $*_L$ wie in Abschnitt 5.14, definiert die **deRham Operatoren** $d^*d + dd^*$, und es gilt²² auf $A^r(U)$

$$d^*d + dd^* = -D_L^2.$$

5.16 Maxwell Gleichungen

Der **Minkowski Raum** $(\mathbb{R}^4, q_L)$ ist $\mathbb{R}^4$ versehen mit der quadratischen **Lorentz Form**

$$q_L(x_1, x_2, x_3, x_4) = -x_1^2 - x_2^2 - x_3^2 + x_4^2;$$

hierbei sei die Lichtgeschwindigkeit $c = 1$ und $(x, y, z, t) = (x_1, x_2, x_3, x_4)$. Die zugehörige orthogonale Gruppe $O(3, 1)$ ist die **Lorentzgruppe**. Der zugehörige Laplace Operator Δ_L ist der **D'Alembert Operator**

$$\square = -\partial_1^2 - \partial_2^2 - \partial_3^2 + \partial_4^2.$$

Für die Lorentz Form $q_L(x) = \sum_{\nu=1}^4 \lambda_\nu x_\nu^2$ ist dann $*_L dx_I = \pm dx_{I^c}$ und das Vorzeichen ist definiert durch $dx_I \wedge *_L dx_I = (\prod_{\nu \in I} \lambda_\nu)^{-1} \cdot dx_1 \wedge \dots \wedge dx_n$.

Notation: $dx_{13} = -dx_{31} = -dx_3 \wedge dx_1$ oder $dx_{134} = -dx_{314}$, etc. Es gilt $dx_{12} \wedge dx_{34} = dx_{31} \wedge dx_{24} = dx_{23} \wedge dx_{14} = dx_{1234}$ sowie $dx_1 \wedge dx_{234} = dx_2 \wedge dx_{314} = dx_3 \wedge dx_{124} = dx_{1234}$, aber beachte $dx_4 \wedge dx_{123} = -dx_{1234}$.

Eine Ladungsverteilung im Raum $\mathbb{R}^3$ ist eine 3-Form $\rho(x_1, x_2, x_3) \cdot dx_1 \wedge dx_2 \wedge dx_3$. Man muß diese Ladungsdichte erst über ein gewisses Volumen, etwa das eines Quaders $Q \subset \mathbb{R}^3$ mitteln, um eine Zahl zu bekommen. D.h. die reelle Zahl

$$(\text{Ladung in } Q) = \int_Q \rho(x_1, x_2, x_3) dx_1 dx_2 dx_3$$

ist die gesamte Ladung, welche in Q enthalten ist (diese hängt ab von der Orientierung!!!)

Ähnlich verhält es sich mit der Strom/Ladungsdichte J in der Minkowski Raum-Zeit. Diese **Strom/Ladungsdichte** ist wiederum eine **3-Form** mit $j_4 =: -\rho(x_1, x_2, x_3, x_4)$, aber nun im $\mathbb{R}^4$; sie ist daher von der Gestalt

$$J = j_1 \cdot dx_{234} + j_2 \cdot dx_{314} + j_3 \cdot dx_{124} - \rho \cdot dx_{123}.$$

²¹für Details dazu siehe auch Abschnitt 11.5.

²² $d^*d + dd^* = -D_L^2$ folgt aus $d^*d + dd^* = d \circ (-1)^r \cdot (*_L)^{-1} d(*_L) + (-1)^{r+1} \cdot (*_L)^{-1} d(*_L) \circ d$ wegen Lemma 5.28 und $(-1)^r \varepsilon_r \varepsilon_{r-1} \varepsilon_{n-r} = (-1)^{r+1} \varepsilon_{r+1} \varepsilon_r = -1$.

5 Ausgewählte Themen I

Die Hodge duale Einsform $j = *_L J = -j_1 dx_1 - j_2 dx_2 - j_3 dx_3 + \rho dx_4$ nennt man den **Strom** (und der Einfachheit werde angenommen²³ die Permeabilität μ sei 1). Die **Ladungserhaltung** zeigt $\int_{\partial Q} J = 0$ für alle 4-dimensionalen Raum-Zeit Würfel Q . Aus Satz 4.35 folgt daher die **erste Maxwell Gleichung** $dJ = (\partial_1 j_1 + \partial_2 j_2 + \partial_3 j_3 + \partial_4 \rho) \cdot dx_{1234} = 0$, d.h.

$$\boxed{dJ = 0}.$$

Seien idealisiert alle Ströme j_i unendlich oft differenzierbar, d.h. $J \in A^3(\mathbb{R}^4)$, so folgt daraus wegen dem Poincare Lemma (Satz 4.30) die Existenz einer 2-Form $\omega \in A^2(\mathbb{R}^4)$ mit

$$\boxed{d\omega = J}.$$

Man schreibt $\omega = -E_1 \cdot dx_{23} - E_2 \cdot dx_{31} - E_3 \cdot dx_{12} + B_1 \cdot dx_{14} + B_2 \cdot dx_{24} + B_3 \cdot dx_{34}$. Das Hodge Dual $F = -*_L \omega$ oder $\omega = *_L F$ definiert die sogenannte **Faraday 2-Form** F

$$\boxed{F = B_1 \cdot dx_{23} + B_2 \cdot dx_{31} + B_3 \cdot dx_{12} + E_1 \cdot dx_{14} + E_2 \cdot dx_{24} + E_3 \cdot dx_{34}}.$$

(E_1, E_2, E_3) und (B_1, B_2, B_3) definieren das **elektrische** respektive **magnetische Feld**. Die **zweite Maxwell Gleichung** lautet dann $dF = 0$. Wegen des Poincare Lemmas gilt deshalb

$$\boxed{F = dA}$$

für eine **1-Form** $A = \sum_i A_i dx_i$, das sogenannte **Vektorpotential** A . Das Vektorpotential A ist durch die Gleichung $F = dA$ nur bestimmt bis auf eine exakte Form $d\varphi$, $\varphi \in C^\infty(\mathbb{R}^4)$. Man kann also A beliebig durch $A + d\varphi$ ersetzen (**Eichfreiheit**, Satz 4.30).

Der $*_L$ -Operator ist nach Lemma 5.28 und 5.29 invariant unter linearen Transformationen der Symmetriegruppe $SO(q_L)$, nimmt bei Spiegelungen aber ein Vorzeichen auf. Der Hodge Operator $d^* = *_L \circ d \circ *_L = (-1)^2 (*_L)^{-1} \circ d \circ *_L$ auf 2-Formen ist dagegen voll invariant unter der **Lorentzgruppe** $O(q_L)$. Daher sind die **Maxwell Gleichungen**

$$\boxed{dF = 0 \quad , \quad d^*F = j}$$

ebenfalls invariant unter der vollen Lorentzgruppe. Die **kombinierten Maxwell Gleichungen**

$$\boxed{DF = j}$$

für $D = d + d^*$ führen wegen $d^2 = (d^*)^2 = 0$ bei geeigneter Eichung²⁴ des Vektorpotentials auf die Bestimmung von A aus der sich dann ergebenden Gleichung $D^2 A = j$

$$D^2 A = (d^*d + dd^*) A = d^*d A = j,$$

also wegen $D^2 = \square$ nach Abschnitt 5.15 bei gegebenem j auf die vier Differentialgleichungen

$$\square A_i = j_i \quad , \quad (i = 1, \dots, 4)$$

²³Im Vakuum ist μ eine Konstante und kann dann in den Maxwell Gleichungen im Prinzip ignoriert werden.

²⁴im Falle der **Lorenz-Eichung** $d^*A = 0$ der Form A .

vom D'Alembertschen Typ. Im elektrostatischen Fall reduzieren sich diese auf $-\Delta u = \rho$ bei Benutzung der *Physiker Notation* $u := -A_4$. Die D'Alembert Differentialgleichungen sind in der Regel lösbar, etwa unter der vereinfachenden Annahme daß j eine 1-Form mit kompaktem Träger im $\mathbb{R}^4$ ist, d.h. im Fall $j_1, \dots, j_4 \in C_c^\infty(\mathbb{R}^4)$. Siehe hierzu Korollar 7.10. Die hierbei gemachte vereinfachende Kompaktheits-Annahme, die in physikalischen Anwendungen natürlich nicht notwendig erfüllt ist, kann aber abgeschwächt werden.

Bemerkung. Für die **Lorenz Eichung** muss man zu A ein $\varphi \in C^\infty(\mathbb{R}^4)$ finden können mit der Eigenschaft $d^*(A + d\varphi) = 0$. Dies führt auf die Gleichung $-d^*d\varphi = d^*A$ und wegen Lemma 5.31 somit erneut auf die D'Alembert Gleichung, jetzt in der Form

$$\square \varphi = d^*A$$

wegen $\square \varphi = -D^2\varphi = -(d^*d + dd^*)\varphi = -d^*d\varphi$; beachte nämlich $d^*\varphi = 0$ aus Gradgründen. Etwa unter den oben genannten Kompaktheits-Bedingungen ist diese Gleichung für φ lösbar.

5.17 Partitionen der Eins

Sei $f(x)$ eine gegebene reellwertige $(r-1)$ -mal stetig partiell differenzierbare Funktion auf dem Intervall $[0, \infty)$ mit der Eigenschaft

$$f^{(\nu)}(0) = 0 \quad , \quad \forall \quad \nu = 0, 1, \dots, r-1 \quad .$$

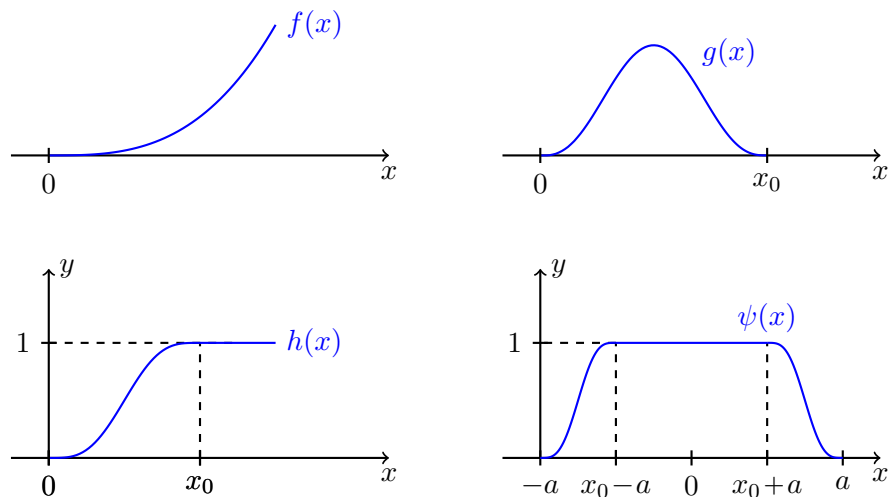
Wir nehmen an $f(x) \geq 0$, sowie $f(x) > 0$ für $x > 0$.

Beispiele. $f(x) = x^r$ für $r < \infty$ und $f(x) = \exp(-\frac{1}{x})$ für $r = \infty$. Sei später obdA $r = \infty$.

Wir setzen für $x < 0$ die Funktion $f(x)$ durch Null zu einer Funktion auf ganz $\mathbb{R}$ fort, und nennen diese Fortsetzung in $C^{r-1}(\mathbb{R})$ wieder $f(x)$.

Für ein beliebiges $x_0 > 0$ ist dann $g(x) = f(x)f(x_0 - x)$ in $C^{r-1}(\mathbb{R})$. Es gilt $g(x) > 0$ genau dann wenn $x \in (0, x_0)$, und $g(x)$ ist Null sonst.

Wegen dem Hauptsatz ist $h(x) = \int_0^x g(t)dt$ dann in $C^r(\mathbb{R})$. Es gilt $h(x) = 0$ für $x \leq 0$, $0 < h(x) < \text{const}$ für $0 < x < x_0$ sowie $h(x) = \text{const}$ für $x \geq x_0$. Hierbei ist $\text{const} = \int_0^{x_0} g(t)dt$ und obdA $\text{const} = 1$ bei geeigneter Normierung von f .



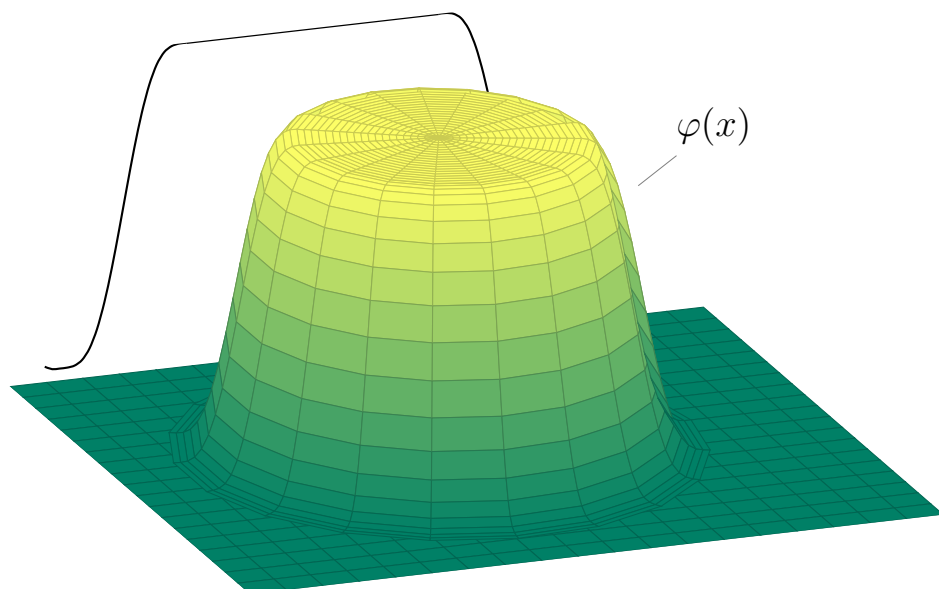
Die Funktion $\psi(x) = h(x+a)h(a-x)$ für $a > x_0$ ist wieder in $C^r(\mathbb{R})$. Es gilt $\psi(x) = 0$ genau dann wenn $x \notin (-a, a)$, und $\psi(x) = 1$ genau dann wenn $x \in [-a+x_0, a-x_0]$, und es gilt $0 < \psi(x) < 1$ sonst. Daraus folgt für $r = \infty$ bei geeigneter Wahl von a und x_0 sofort

Lemma 5.32. *Die Funktionen in $C_c^\infty(\mathbb{R}^n)$ trennen Punkte in $\mathbb{R}^n$.*

Beweis. In der Tat hat die C_c^∞ -Funktion

$$\varphi(x) = \varphi_{\xi,a}(x) = \psi(d_{\mathbb{R}^n}(x, \xi))$$

ihren Träger in einer Kugel $K_a(\xi)$ vom Radius a um den Punkt $\xi \in \mathbb{R}^n$. Sie ‘trennt’ daher ξ von jedem Punkt $x \in \mathbb{R}^n$, der weiter als a von ξ entfernt ist. $\square$



Sei nun M eine beliebige kompakte Teilmenge von $\mathbb{R}^n$ und sei $M = \bigcup_{\xi \in I} M \cap K_r(\xi)$ eine endliche Überdeckung von M durch offene Kugeln um endlich viele Punkte $\xi \in M$ (siehe Satz 11.4). Die Radien $r = a/2 > 0$ mögen hierbei von den Punkten ξ abhängen: $r = r(\xi)$. Sei N die offene Menge

$$N = \bigcup_{\xi \in I} K_{2r}(\xi) .$$

Es gilt $M \subset N$. Für eine gegebene offene Menge im $\mathbb{R}^n$, welche M enthält, kann obdA durch geeignete Wahl der Radien $r = r(\xi)$ angenommen werden, daß N in dieser offenen Menge liegt. Für $x \in N$ ist dann

$$\varphi_\xi(x) = \frac{\varphi_{\xi,a(\xi)}(x)}{\sum_{\xi \in I} \varphi_{\xi,a(\xi)}(x)}$$

eine wohldefinierte C^∞ -Funktion und es gilt

$$\boxed{\sum_{\xi \in I} \varphi_\xi(x) = 1} \quad , \quad \forall x \in \tilde{N}$$

auf der offenen Teilmenge $\tilde{N} = \bigcup_{\xi \in I} K_r(\xi)$ von N , welche M enthält. Man nennt die so konstruierte Funktionenschar φ_ξ (für $\xi \in I$) dann eine **Partition der Eins** auf $\tilde{N}$ (oder M), welche der gegebenen Überdeckung von M zugeordnet ist.

Bemerkung. Für $n \in \mathbb{N}$ und ein Polynom $P(t)$ ist die Funktion

$$f(t) = t^{-n} P(t) \exp\left(-\frac{1}{t}\right)$$

differenzierbar im Bereich $U = (0, \infty]$ mit der Ableitung [benutze Produkt- und Kettenregel]

$$f'(t) = t^{-n-2} Q(t) \exp\left(-\frac{1}{t}\right)$$

für ein geeignetes²⁵ Polynom $Q(t)$. Setzt man formal $f(0) := 0$, gilt

$$\lim_{t \rightarrow 0} \frac{f(t) - f(0)}{t - 0} = \lim_{t \rightarrow 0} t^{-n-1} P(t) \exp\left(-\frac{1}{t}\right) = 0$$

für jede Folge in U , die gegen Null konvergiert²⁶. Somit ist f differenzierbar im Punkt Null mit Ableitung

$$f'(0) = 0 .$$

Dies zeigt rekursiv, daß $f(t) = \exp(-1/t)$ unendlich oft differenzierbar ist auf $[0, \infty)$, und die Ableitungen beliebig hoher Ordnung von f im Punkt Null verschwinden.

²⁵ $Q(t) = P(t) - ntP(t) + t^2P'(t)$.

²⁶Man reduziert auf die Aussage $\lim_{t \rightarrow +\infty} t^m \exp(-t) = 0$. Wegen $t^m \exp(-t) = (mx \exp(-x))^m$ für $x = t/m$ genügt $\lim_{x \rightarrow +\infty} g(x) = 0$ für $g(x) = x \exp(-x)$. Wegen $g'(x) = (1-x) \exp(-x) \leq 0$ für $x \in [1, \infty)$ und dem Mittelwertsatz fällt $g(x)$ monoton in $[1, \infty)$. Nach dem Prinzip der monotonen Konvergenz existiert daher der Limes $\xi = \lim_{x \rightarrow \infty} g(x)$. Für $y = 2x$ folgt $\xi = \lim_{y \rightarrow \infty} g(y) = \lim_{x \rightarrow \infty} (g(x)2 \exp(-x)) = \xi \cdot \lim_{x \rightarrow \infty} 2 \exp(-x) = 0$.

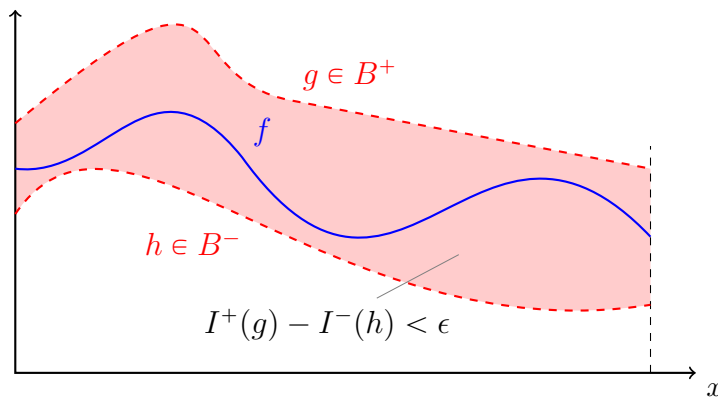
6 Lebesgue Integration

6.1 Übersicht

Für einen Funktionenverband $B(X)$ auf einer Menge X und ein abstraktes Integral I auf $B(X)$ wurde in Satz 3.20 gezeigt, daß sich I eindeutig zu abstrakten Integralen $I^\pm$ der Halbverbände der monotonen Hüllen $B^\pm(X) \supset B(X)$ fortsetzen lässt. Eine beliebige Funktion

$$f : X \rightarrow \hat{\mathbb{R}} := \mathbb{R} \cup \{+\infty\} \cup \{-\infty\}$$

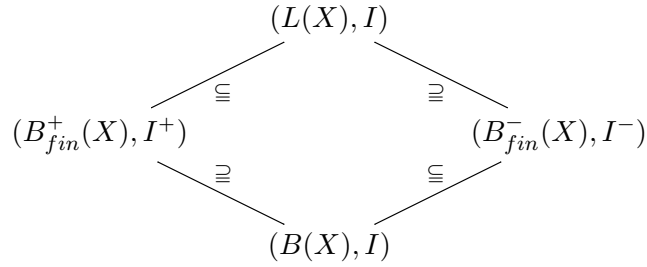
heißt dann **integrierbar** bezüglich $(B(X), I)$, wenn es für jedes $\varepsilon > 0$ geeignete Funktionen $h \in B^-(X)$ und $g \in B^+(X)$ gibt mit $h \leq f \leq g$ und $I^+(g) - I^-(h) < \varepsilon$. Für integrierbare Funktionen kann das Lebesgue Integral $I(f) \in \mathbb{R}$ definiert werden. Dieses stimmt auf $B(X)$ mit dem gegebenen Integral I überein. Sei $\hat{L}(X)$ die Menge der integrierbaren Funktionen.



Wir zeigen, daß die Teilmenge der Punkte $x \in X$, wo eine integrierbare Funktion die Werte $\pm\infty$ annimmt, eine Nullmenge ist, und dass man integrierbare Funktionen auf Nullmengen beliebig abändern kann ohne das Lebesgue Integral zu verändern.

Die $\mathbb{R}$ -wertigen Funktionen f in $B^\pm(X)$ mit endlichem Integral $I^\pm(f) \neq \pm\infty$ bilden einen Halbverband $B_{fin}^\pm(X)$ in $B^\pm(X)$.

Die $\mathbb{R}$ -wertigen Funktionen in $\hat{L}(X)$ bilden einen Verband $L(X)$ und das Lebesgue Integral definiert ein abstraktes Integral auf $L(X)$. Es ist das eindeutig bestimmte abstrakte Integral auf $L(X)$, das auf $B(X) \subset L(X)$ mit dem gegebenen $I : B(X) \rightarrow \mathbb{R}$ übereinstimmt.



Eine Funktion ist Lebesgue integrierbar genau dann wenn sie außerhalb von einer Nullmenge mit einer Funktion in $L(X)$ übereinstimmt. Eine erneute Anwendung des Lebesgue Prozeßes auf $(L(X), I)$ liefert dann keine neuen integrierbaren Funktionen mehr.

6.2 Das Lebesgue Integral

Für eine Menge X , einen Verband von Funktionen $B = B(X)$ auf X und ein abstraktes Integral $I : B(X) \rightarrow \mathbb{R}$ schreiben wir kurz $B^+ = B^+(X)$ und $B^- = B^-(X)$ für die monotonen Hüllen von $B(X)$. Beachte $B^\mp = -B^\pm$. I kann zu abstrakten Integralen $I^+ : B^+ \rightarrow \mathbb{R}^+$ und $I^- : B^- \rightarrow \mathbb{R}^-$ fortgesetzt werden (Satz 3.20). Schließlich gilt für $h \in B^-$

$$I^-(h) = -I^+(-h).$$

[Beweis: Für $h \in B^-$ war $I^-(h) := \liminf_n I(h_n)$ für jede Folge $h_n \searrow h$ mit $h_n \in B(X)$. Benutze $I^+(-h) := \limsup_n I(-h_n) = -\liminf_n I(h_n) = -I^-(h)$ wegen $-h_n \nearrow -h \in B^+$].

Lemma 6.1. Für einen Verband $B(X)$ und $B^-(X) \ni h \leq g \in B^+(X)$ gilt $I^-(h) \leq I^+(g)$.

Beweis. Wegen $-h \in -B^- = B^+$ gilt $g + (-h) \in B^+ + B^+ = B^+$. Aus der Annahme $g + (-h) \geq 0$ folgt $I^+(g) + I^+(-h) = I^+(g + (-h)) \geq 0$ mittels Additivität und Monotonie. Oben wurde gezeigt $I^+(-h) = -I^-(h)$. Daraus folgt $I^+(g) - I^-(h) = I^+(g) + I^+(-h) \geq 0$, d.h. $I^+(g) \geq I^-(h)$. $\square$

Definition 6.2 (Integrierbarkeit). Eine beliebige Funktion

$$f : X \rightarrow \hat{\mathbb{R}} := \mathbb{R} \cup \{\infty\} \cup \{-\infty\}$$

heißt **Lebesgue-integrierbar** bezüglich (B, I) oder nur kurz **integrierbar**, wenn folgendes gilt: Für alle reellen Zahlen $\varepsilon > 0$ existieren¹ Funktionen $h \in B^-$ und $g \in B^+$ mit den Eigenschaften

¹Man sieht dann a posteriori auch, dass es reicht wenn für alle $\varepsilon > 0$ Lebesgue integrierbare Funktionen h und g existieren mit $h \leq f \leq g$ und $I(g) - I(h) < \varepsilon$. [Denn für $g^+ \in B^+$ mit $g \leq g^+$, $I(g) \leq I^+(g^+)$ sowie $I^+(g^+) - I(g) < \varepsilon/2$ und analog $h^- \leq h$, $h^- \in B^-$ mit $I^-(h^-) \leq I(h)$ und $I(h) - I^-(h^-) < \varepsilon/2$ gilt $h^- \leq f \leq g^+$ sowie $I^+(g^+) - I^-(h^-) < \varepsilon + I(g) - I(h) < 2\varepsilon$.]

$h \leq f \leq g$ und²

$$I^+(g) - I^-(h) < \varepsilon .$$

Für eine beliebige Funktion $f : X \rightarrow \hat{\mathbb{R}}$ definieren wir

$$I^b(f) := \sup\{I^-(h) \mid h \in B^- \text{ mit } h \leq f\} \quad , \quad I^\sharp(f) := \inf\{I^+(g) \mid g \in B^+ \text{ mit } f \leq g\}$$

Ist eine der beiden Mengen leer, setzt man $I^b(f) := -\infty$ resp. $I^\sharp(f) := \infty$. Es gilt

$$I^b(f) \leq I^\sharp(f) .$$

[Wegen $I^-(h) \leq I^+(g)$ für $h \leq f \leq g$ (Lemma 6.1) ist $\sup_{h \in B^-, h \leq f} I^-(h) =: I^b(f) \leq I^+(g)$.

Aus $I^b(f) \leq I^+(g)$ folgt $I^b(f) \leq I^\sharp(f) := \inf_{g \in B^+, f \leq g} I^+(g)$.]

Lemma 6.3 (Lebesgue Integral). *Ist $f : X \rightarrow \hat{\mathbb{R}}$ bezüglich (B, I) Lebesgue-integrierbar, dann gilt $I^b(f) = I^\sharp(f)$ und der so definierte Wert liegt in $\mathbb{R}$, d.h. ist verschieden von $\pm\infty$.*

Wir bemerken, daß offensichtlich die Gleichheit $I^b(f) = I^\sharp(f)$ sogar äquivalent ist zur Lebesgue Integrierbarkeit der Funktion f .

Definition 6.4. *Die in Lemma 6.3 definierte reelle Zahl wird das **Lebesgue Integral** $I(f)$ der bezüglich (B, I) integrierbaren Funktion f genannt*

$$I(f) := I^b(f) = I^\sharp(f) .$$

Beweis. Wegen

$$0 \leq I^\sharp(f) - I^b(f) \leq I^+(g) - I^-(h) < \varepsilon$$

für alle $\varepsilon > 0$ folgt $I^b(f) = I^\sharp(f)$ und diese Zahl ist in $\mathbb{R}$ [Aus $I^+(g) - I^-(h) < \varepsilon$ folgt $I^+(g) < \infty$ sowie $I^-(h) > -\infty$ wegen $I^+(g) \in \mathbb{R}^+$, $I^-(h) \in \mathbb{R}^-$. Nach Lemma 6.1 ist damit $\{I^-(h) \mid B^- \ni h \leq f \leq g\}$ nach oben durch $I^+(g) < \infty$ beschränkt, und das Supremum $I^b(f)$ liegt somit in $\mathbb{R}$. Ditto für $I^\sharp(f)$.] $\square$

Lemma 6.5. *$f \in B^+$ (resp. B^-) ist f Lebesgue integrierbar bezüglich (B, I) genau dann, wenn $I^+(f) \neq \infty$ (resp. $I^-(f) \neq -\infty$) gilt. In diesem Fall ist $I^b(f) = I^\sharp(f)$ gleich $I^+(f)$ (resp. $I^-(f)$). Insbesondere sind Funktionen aus $B \subseteq B^+$ Lebesgue integrierbar, und das Lebesgue Integral stimmt auf B mit dem ursprünglich gegebenen Integral $I : B \rightarrow \mathbb{R}$ überein.*

Beweis. Aus $B^+ \ni f \leq g \in B^+$ folgt $I^+(f) \leq I^+(g)$. Das Infimum aller Werte $I^+(g)$, für $g \in B^+$ und $f \leq g$, wird daher bei $g = f \in B^+$ angenommen, d.h. $I^\sharp(f) = I^+(f) < +\infty$. Da $I^b(f) \leq I^\sharp(f)$ ja immer gilt, verbleibt zum Beweis von $I^b(f) = I^\sharp(f) = I^+(f) < \infty$ nur $I^+(f) \leq I^b(f)$ zu zeigen für $f \in B^+$. Wähle dazu Funktionen $h_n \in B$ mit $h_n \nearrow f$. Nach Definition von $I^+(f)$ gilt $I(h_n) \nearrow I^+(f)$. Wegen $I^-(h_n) = I(h_n)$ folgt

$$I^\sharp(f) = I^+(f) := \sup_n I(h_n) = \sup_n I^-(h_n) \leq \sup_{h \in B^-, h \leq f} I^-(h) =: I^b(f) .$$

²Nach Lemma 6.1 gilt ausserdem $0 \leq I^+(g) - I^-(h)$.

Beispiel 6.6. Für *kein* $\alpha \in \mathbb{R}$ ist $f(x) = |x|^\alpha$ Lebesgue integrierbar für das Standardintegral I auf $B(X) = CT(\mathbb{R})$ [anderenfalls wäre wegen $B^- \ni h_n(x) = \chi_{[1/n, n]}(x) \cdot f(x) \leq f(x)$ dann $\lim_n I^-(h_n) \leq I^\sharp(f) < +\infty$ im Widerspruch zu $I^-(h_n) = \int_{1/n}^n x^\alpha dx = \frac{n^{\alpha+1} - n^{-\alpha-1}}{\alpha+1}$ (der Einfachheit halber $\alpha \neq -1$) mit $\sup_n I^-(h_n) = +\infty$.]

Aber dasselbe Argument zeigt mit Hilfe von Lemma 6.5 die *Integrierbarkeit* von $|x|^\alpha$ auf jeder *kompakten Teilmenge* $A \subset \mathbb{R}$ im Fall $\alpha > -1$. [Für $h_n(x) := \chi_{A \setminus (-1/n, +1/n)}(x) \cdot |x|^\alpha$ gilt nämlich $h_n(x) \in B(X) = CT(\mathbb{R})$ und $h_n(x) \nearrow x^\alpha$ (mit $0^\alpha := 0$) und $I^+(h_n) \leq \text{const.} + 2 \frac{x^{\alpha+1}}{\alpha+1} \Big|_{1/n} \leq C < \infty$. Benutze nämlich $\lim_n (\frac{1}{n})^{\alpha+1} = 0$ für $\alpha + 1 > 0$].

6.3 Der Verband $L(X)$

Satz 6.7. Die $\mathbb{R}$ -wertigen (!) bezüglich (B, I) Lebesgue integrierbaren Funktionen

$$\boxed{f : X \rightarrow \mathbb{R}}$$

bilden einen **Verband** $L(X) = \{f : X \rightarrow \mathbb{R} \mid f \text{ ist Lebesgue-integrierbar}\}$, welcher B umfasst.

Sei allgemeiner $\hat{L}(X) = \{f : X \rightarrow \hat{\mathbb{R}} \mid f \text{ ist Lebesgue-integrierbar}\}$. Der letzte Satz ergibt sich aus der folgenden allgemeineren Aussage.

Satz 6.8 (Permanenzeigenschaften). Die Menge aller Funktionen in $\hat{L}(\mathbb{R})$ ist unter reeller Skalarmultiplikation, Addition³ und den Bildungen $\min, \max$ abgeschlossen. Insbesondere gilt

$$\boxed{f \in \hat{L}(X) \implies |f| \in \hat{L}(X)}.$$

Beweis. *Schritt 1.* Für f_1, f_2 gilt $I^\sharp(f_1 + f_2) \leq I^\sharp(f_1) + I^\sharp(f_2)$ sowie $I^\sharp(\lambda f) \leq \lambda I^\sharp(f)$ für $\lambda > 0$, als unmittelbare Folge der Definition von $I^\sharp$. Entsprechendes gilt für $I^\flat$ mit den umgekehrten Ungleichungen [wegen $I^\sharp(-f) = -I^\flat(f)$ und $I^\flat(-f) = -I^\sharp(f)$]. Daraus folgt $I^\flat(f_1) + I^\flat(f_2) \leq I^\flat(f_1 + f_2) \leq I^\sharp(f_1 + f_2) \leq I^\sharp(f_1) + I^\sharp(f_2)$. Sind $f_1, f_2 \in \hat{L}(X)$, folgt daraus leicht $f_1 + f_2 \in \hat{L}(X)$. Aus $I^\flat(f_i) = I^\sharp(f_i)$ folgt ausserdem dann

$$\boxed{I(f_1 + f_2) = I(f_1) + I(f_2)}.$$

Ditto zeigt man $I(\lambda \cdot f) = \lambda \cdot I(f)$ für alle $f \in \hat{L}(X)$ und $\lambda \geq 0$.

Schritt 2. Aus $h \leq f \leq g$ und $h \in B^-, g \in B^+$ sowie $I^+(g) - I^-(h) < \varepsilon$ folgt sofort $-g \leq -f \leq -h$ mit $-g \in B^-, -h \in B^+$ und $I^+(-h) - I^-(-g) = I^+(g) - I^-(h) < \varepsilon$, und analog $I^\sharp(-f) = -I^\flat(f)$ und $I^\flat(-f) = -I^\sharp(f)$. Ist $f \in \hat{L}(X)$, ist daher auch $-f \in \hat{L}(X)$ und es gilt $I(-f) = -I(f)$. Zusammen mit Schritt 1 folgt daher für alle $f \in \hat{L}(X)$ und alle $\lambda \in \mathbb{R}$

$$\boxed{I(\lambda \cdot f) = \lambda \cdot I(f)}.$$

³Hierbei dürfen wir in den Punkten x , wo $f_1(x) = \pm\infty$ und $f_2(x) = \mp\infty$ ist, für $f = f_1 + f_2$ den Funktionswert $f(x) \in \hat{\mathbb{R}}$ beliebig wählen.

Schritt 3. Wegen $\min(f, \tilde{f}) = -\max(-f, -\tilde{f})$ genügt es, dass für $f, \tilde{f} \in \hat{L}(X)$ auch $\max(f, \tilde{f})$ in $\hat{L}(X)$ ist. Wähle dazu $h \leq f \leq g$ resp. $\tilde{h} \leq \tilde{f} \leq \tilde{g}$ mit $I^+(g) - I^-(h) = I^+(g-h) < \varepsilon$ resp. $I^+(\tilde{g}) - I^-(\tilde{h}) = I^+(\tilde{g} - \tilde{h}) < \varepsilon$. Dann gilt automatisch $I^+(\tilde{g}) < \infty$ und $I^+(g) < \infty$ sowie $B^- \ni \max(h, \tilde{h}) \leq \max(f, \tilde{f}) \leq \max(g, \tilde{g}) \in B^+$. Schliesslich gilt

$$I^+(\max(g, \tilde{g})) - I^-(\max(h, \tilde{h})) = I^+(\max(g, \tilde{g}) - \max(h, \tilde{h})) \leq I^+(g-h) + I^+(\tilde{g}-\tilde{h}) < 2\varepsilon$$

unter Benutzung von $g-h \in B^+$ und $\tilde{g}-\tilde{h} \in B^+$ und⁴ $\max(g, \tilde{g}) - \max(h, \tilde{h}) \leq (g-h) + (\tilde{g}-\tilde{h})$. Also ist $\max(f, \tilde{f})$ integrierbar. $\square$

Lemma 6.9. Für Funktionen $f_1, f_2 \in \hat{L}(X)$ und $\lambda_1, \lambda_2 \in \mathbb{R}$ gilt $I(\lambda_1 f_1 + \lambda_2 f_2) = \lambda_1 I(f_1) + \lambda_2 I(f_2)$, und aus $f_1 \leq f_2$ folgt $I(f_1) \leq I(f_2)$.

Beweis. Die erste Behauptung wurde schon gezeigt. Für die Monotonie genügt es $I(f) \geq 0$ zu zeigen für $f \in \hat{L}(X)$ mit $f \geq 0$. Beachte $I(f) = I^\sharp(f) = \inf\{I^+(g)\}$, für das Infimum über alle $g \in B^+$ mit $g \geq f \geq 0$. Aus der Monotonie von I^+ folgt $I^+(g) \geq 0$ wegen $g \geq 0$. Also $\inf\{I^+(g)\} \geq 0$. Es folgt $I(f) \geq 0$, d.h. die Monotonie des Lebesgue Integrals I . $\square$

Aus dem letzten Lemma und Satz 6.11 (Beppo Levi) des nächsten Abschnitts erhält man als

Korollar 6.10. Auf dem Verband $L(X) = L(X, B, I)$ definiert das Lebesgue Integral I ein abstraktes Integral.

6.4 Vertauschungssätze

Satz 6.11 (Beppo Levi). Für eine monotone Folge $f_n \nearrow f$ von integrierbaren Funktionen $f_n \in \hat{L}(X)$ definiert $I(f_n)$ eine monoton wachsende Folge reeller Zahlen. Gilt

$$\boxed{\kappa := \sup_n I(f_n) < \infty},$$

dann ist auch die punktweise Grenzfunktion $f = \sup_n f_n$ in $\hat{L}(X)$ und es gilt

$$\boxed{I(f) = I(\sup_n f_n) = \sup_n I(f_n) = \kappa}.$$

Beweis. Fixiere $\varepsilon > 0$. Für jedes $n \in \mathbb{N}$ gibt es dann $h_n \in B^-$ sowie $g_n \in B^+$ mit

$$h_n \leq f_n \leq g_n$$

sowie (*)

$$I^+(g_n - h_n) = I^+(g_n) - I^-(h_n) < \frac{\varepsilon}{2^n}.$$

OBdA sind die h_n monoton wachsend [ersetze sonst h_{n+1} durch $\max(h_n, h_{n+1})$ usw.]. Aus (*) folgt $I^-(h_n) > -\infty$, und nach Lemma 6.5 damit $h_n \in \hat{L}(X)$ sowie $I^-(h_n) = I(h_n)$. Analog folgt $g_n \in \hat{L}(X)$ mit $I^+(g_n) = I(g_n)$. Wegen $I^-(h_n) = I(h_n) \leq I(f_n) \leq \kappa$ ist die monotone

⁴Ist obdA $g(x) \geq \tilde{g}(x)$, dann ist $\max(g, \tilde{g}) - \max(h, \tilde{h}) \leq g - h \leq (g-h) + (\tilde{g}-\tilde{h})$ im Punkt x .

6 Lebesgue Integration

Folge $I(h_n)$ durch $\kappa < \infty$ beschränkt. Sie konvergiert daher. Also gilt $I(h_m) - I(h_n) < \varepsilon$ für alle $m \geq n \geq N(\varepsilon)$, und damit $I(h_m - h_n) < \varepsilon$ sowie $I(g_i - h_i) < \frac{\varepsilon}{2^i}$ nach Satz 6.8.

Für die monotone Folge $\tilde{g}_n := \max_{i=1}^n (g_i) \in B^+$ gilt $h_n \leq f_n \leq \tilde{g}_n$. Aus Lemma 3.16 folgt daher $\tilde{g}_n \nearrow \tilde{g} \in (B^+)^+ = B^+$ mit

$$B^- \ni h_n \leq f_n \leq \tilde{g} \in B^+.$$

Aus der Halbstetigkeit von I^+ auf B^+ folgt $I^+(\tilde{g} - h_n) = \sup_{m \geq n} I^+(\tilde{g}_m - h_n)$. Weiterhin gilt $I(\tilde{g}_m - h_n) \leq I(h_m - h_n) + \sum_{i=1}^m I(g_i - h_i) < \varepsilon + \sum_{i \geq 1} \frac{\varepsilon}{2^i} \leq 2\varepsilon$ wegen $\tilde{g}_m - h_n \leq (h_m - h_n) + \sum_{i=1}^m (g_i - h_i)$; alle Summanden sind in $\hat{L}(X)$. Es folgt $I^+(\tilde{g}) - I^-(h_n) = I^+(\tilde{g} - h_n) = \sup_m I^+(\tilde{g}_m - h_n) = \sup_m I(\tilde{g}_m - h_n) \leq 2\varepsilon$, und damit sofort die Integrierbarkeit von f . Aus $I(h_n) \leq I(f_n) \leq I(f) \leq I^+(\tilde{g})$ folgt außerdem $0 \leq I(f) - I(f_n) \leq 2\varepsilon$, und damit $I(f) = \lim_n I(f_n)$. $\square$

Bemerkung. Eine analoge Aussage gilt für monoton fallende Funktionenfolgen $f_n \in \hat{L}(X)$ im Fall $\inf_n I(f_n) > -\infty$.

Bemerkung. Wie der Satz von Beppo Levi zeigt, stimmen die monotonen Hüllen $L_{fin}^\pm(X)$ von $(L(X), I)$ mit $L(X)$ überein. Aus Fußnote 1 und Beppo Levi folgt, daß eine Iteration des Lebesgueprozesses, angewendet auf $(L(X), I)$, keine neuen integrierbaren Funktionen mehr liefert.

Satz 6.12 (Satz von Lebesgue). Sei $f_n \rightarrow f$ eine punktweise konvergente Funktionenfolge $f_n \in \hat{L}(X)$, so dass eine von n unabhängige Funktion $F \in \hat{L}(X)$ existiert mit $|f_n| \leq F$ für alle $n \in \mathbb{N}$. Dann ist auch f in $\hat{L}(X)$ und es gilt

$$\boxed{I(\lim_n f_n) = \lim_n I(f_n)}.$$

Man nennt diesen Satz auch **Satz von der dominierten Konvergenz**, da die Funktionen f_n nach Annahme von der fixierten integrierbaren Funktion F dominiert werden.

Beweis. *Schritt 1.* Es gilt

$$\varphi_n(x) := \inf_{i \geq n} (f_i(x)) \nearrow f(x).$$

Für festes n gilt

$$\psi_m(x) := \inf_{m \geq i \geq n} (f_i(x)) \searrow \varphi_n(x).$$

Beachte $\psi_m \in L(X)$ nach Satz 6.8. Wegen $-F \leq \psi_m$ gilt $-\infty < I(F) \leq \lim_m I(\psi_m)$, und somit ist die Grenzfunktion φ_n integrierbar nach Beppo Levi. Wegen $\varphi_n \leq F$ gilt $\lim_n I(\varphi_n) \leq I(F) < \infty$, und somit ist f (erneut nach Beppo Levi) integrierbar mit $I(\varphi_n) \rightarrow I(f)$.

Schritt 2. Analog definiert man $\tilde{\varphi}_n(x) = \sup_{i \geq n} (f_i(x)) \searrow f(x)$ und zeigt die Konvergenz $I(\tilde{\varphi}_n(x)) \rightarrow I(f)$. Nach Definition gilt

$$\varphi_n \leq f_n \leq \tilde{\varphi}_n.$$

Daraus folgt dann wie behauptet die Konvergenz $I(f_n) \rightarrow I(f)$. $\square$

6.5 Anwendungen

Sei $Y = \mathbb{R}^n$ oder $Y = \mathbb{Z}$ oder ein Produkt $\mathbb{R}^n \times \mathbb{Z}^m$ etc. und $B(Y) = C_c(Y)$.

Satz 6.13. Sei Z ein metrischer Raum und $f : Y \times Z \rightarrow \mathbb{R}$ stetig in der Variable $z \in Z$ für festes $y \in Y$. Gilt $|f(y, z)| \leq F(y)$ für ein $F \in L(Y)$ und alle $z \in Z$, und ist $f(y, z)$ in $L(Y)$ für festes $z \in Z$, dann ist $g(z) := \int_Y f(y, z) dy$ definiert und stetig in der Variable z .

Beweis. Für jede Folge $z_n \rightarrow z$ ist $f_n(y) := f(y, z_n) \in L(Y)$. Wegen $|f_n(y)| \leq F(y)$ folgt aus dem Satz 6.12 von der dominierten Konvergenz

$$\lim_{n \rightarrow \infty} \int_Y f_n(y) dy = \int_Y \lim_{n \rightarrow \infty} f_n(y) dy.$$

Also $\lim_n \int_Y f(z_n, y) dy = \int_Y f(z, y) dy$. □

Beweis von Vertauschungssatz 4.32 (siehe Abschnitt 4.13). Für jede Folge $x_n \rightarrow \xi$ ($x_n \neq \xi$) gibt es wegen dem Mittelwertsatz ein $\theta_{y,n} \in [a, b]$ zwischen ξ und x mit

$$f(x_n, y) - f(\xi, y) = (x_n - \xi) \cdot \partial_x f(\theta_{y,n}, y) =: (x_n - \xi) \cdot f_n(y).$$

Nach Annahme ist die linke Seite in $L(Y)$. Also gilt $f_n(y) \in L(Y)$, und nach der Definition der Ableitung $\lim_n f_n(y) = f'(\xi, y)$. Wegen $|f_n(y)| = |\partial_x f(\theta_{y,n}, y)| \leq F(y)$ folgt $\int_Y f_n(y) dy \rightarrow \int_Y \partial_x f(\xi, y) dy$ aus dem Satz von der dominierten Konvergenz 6.12. Nach Definition von $f_n(y)$ ist der Limes der $\int_Y f_n(y) dy$ dann $\lim_n (g(x_n) - g(\xi)) / (x_n - \xi) = g'(\xi)$. □

Korollar 6.14 (Fubini). Für das Produkt $X = Y \times Z$ zweier nicht entarteter beschränkter Quader Y, Z wie in Beispiel 2.25 definiert daher $I(f) := \int_Z f(y, z) dz$ eine lineare Abbildung

$$I : B = C(X) \rightarrow \tilde{B} = C(Y).$$

Es gilt $\int_X f dy dz = \int_Y I(f) dy$ (da dies richtig ist für Treppenfunktionen f und da I , $\int_Y$ und $\int_X$ $\mathbb{R}$ -linear, monoton und halbstetig sind).

Eine reelle Folge $f : \mathbb{N} \rightarrow \mathbb{R}$ definiert eine sogenannte **absolut konvergente Reihe** $\sum_{n \in \mathbb{N}} f(n)$, wenn f in $L(\mathbb{N})$ liegt. In der Tat: Da $L(\mathbb{N})$ ein Verband ist, ist mit f auch die Folge $|f(n)|$ der Absolutbeträge in $L(\mathbb{N})$, und aus Satz 6.12 folgt

$$\lim_{N \rightarrow \infty} \sum_{n=1}^N |f(n)| < \infty.$$

Die Umkehrung gilt auch: Aus $\lim_{N \rightarrow \infty} \sum_{n=1}^N |f(n)| < \infty$ folgt $|f(n)| \in L(\mathbb{N})$ nach dem Satz von Beppo Levi oder Lemma 6.5, sowie die absolute Konvergenz im Sinne von $f \in L(\mathbb{N})$ wegen des Satzes 6.12 von der dominierten Konvergenz (mit der Majorante $F = |f|$).

Satz 6.15 (Umordnungssatz). *Liegt die Folge $f : \mathbb{N} \rightarrow \mathbb{R}$ in $L(\mathbb{N})$ (d.h., ist die Reihe $\sum_{n=1}^{\infty} f(n)$ absolut konvergent), dann gilt für jede Bijektion $\sigma : \mathbb{N} \rightarrow \mathbb{N}$*

$$I(f) = \lim_{n \rightarrow \infty} \sum_{i=0}^n f(\sigma(i)) .$$

Beweis. Sei $F(x) = |f(x)|$. Da $f_n(x) := f(x)\chi_{[0,1,\dots,n]}(\sigma^{-1}(x))$ punktweise gegen $f(x)$ konvergiert, konvergiert $I(f_n) = \sum_{i=0}^n f(\sigma(i))$ gegen $I(f)$ nach Satz 6.12. $\square$

6.6 Nullmengen

Eine Teilmenge $Y \subseteq X$ heisst **integrierbar** (bezüglich $(B(X), I)$), wenn die charakteristische Funktion χ_Y integrierbar ist. Man nennt die reelle Zahl $\text{vol}(Y) = I(\chi_Y) \geq 0$ das **Volumen** von Y . Sind Y_1, Y_2 integrierbar, dann nach Satz 6.7 auch $Y_1 \cap Y_2$ und $Y_1 \cup Y_2$ und es gilt

$$\boxed{\text{vol}(Y_1) + \text{vol}(Y_2) = \text{vol}(Y_1 \cup Y_2) + \text{vol}(Y_1 \cap Y_2)} .$$

Dies folgt aus $\chi_{Y_1} + \chi_{Y_2} = \chi_{Y_1 \cup Y_2} + \chi_{Y_1 \cap Y_2}$ und $\max(\chi_{Y_1}, \chi_{Y_2}) = \chi_{Y_1 \cup Y_2}$ sowie analog $\min(\chi_{Y_1}, \chi_{Y_2}) = \chi_{Y_1 \cap Y_2}$. Insbesondere gilt $\text{vol}(Y_1 \cup Y_2) \leq \text{vol}(Y_1) + \text{vol}(Y_2)$.

Beispiel. Wegen Beispiel 2.25 und dem Satz von Beppo Levi ist jede *kompakte* Teilmenge des $\mathbb{R}^n$ integrierbar im obigen Sinne (bezüglich des standard Lebesgue Integrals des $\mathbb{R}^n$).

Eine Nullmenge ist eine integrierbare Menge vom Volumen Null.

Für eine Nullmenge Y gilt $I(n \cdot \chi_Y) = 0$ für alle natürlichen Zahlen n . Im Limes $n \rightarrow \infty$ folgt daher wegen Beppo Levi: Die Funktion, die $+\infty$ auf Y ist, und Null sonst, ist integrierbar mit Integral Null. Ditto für $-\infty$ anstelle von $+\infty$. Aus $-\infty \cdot \chi_Y \leq f \leq +\infty \cdot \chi_Y$ folgt daher $f \in \hat{L}(X)$ und $I(f) = 0$. Somit gilt:

Jede Funktion f mit Werten in $\mathbb{R} \cup \{\infty\} \cup \{-\infty\}$ und Träger in einer Nullmenge ist integrierbar mit Integral Null. Insbesondere ist jede Teilmenge einer Nullmenge eine Nullmenge.

Lemma 6.16 (ε -Kriterium für Nullmengen). *$Y \subseteq X$ ist eine Nullmenge, wenn für jedes $\varepsilon > 0$ eine abzählbare Überdeckung von Y durch integrierbare Mengen U_i existiert mit*

$$\sum_{i=1}^{\infty} \text{vol}(U_i) < \varepsilon .$$

Beweis. Mittels Beppo Levi für $f_n = \chi_{U_1 \cup \dots \cup U_n} \nearrow \chi_U$ zeigt man: $U := \bigcup_{i=1}^{\infty} U_i$ ist integrierbar mit $\text{vol}(U) < \varepsilon$. Für $\varepsilon = 1/n$ sei $U(n)$ diese Vereinigung. Dann bilden die Durchschnitte $Z_n := \bigcap_{i=1}^n U(i)$ eine absteigende Kette integrierbarer Mengen mit $\text{vol}(Z_n) < \frac{1}{n}$. Daher ist $Z = \bigcap_{i=1}^{\infty} Z_n$ eine Nullmenge (Beppo Levi). Aus $Y \subseteq Z$ folgt die Behauptung. $\square$

Ein Spezialfall:

Eine abzählbare Vereinigung Y von Nullmengen Y_i , $i = 1, 2, \dots$ ist eine Nullmenge.

Lemma 6.17. Sei $f : X \rightarrow \mathbb{R} \cup \{\infty\} \cup \{-\infty\}$ integrierbar. Dann ist die Menge der Unendlichkeitsstellen $f^{-1}(\{\pm\infty\})$ eine Nullmenge in X .

Beweis. Es genügt, daß $\Sigma = f^{-1}(+\infty)$ Nullmenge ist [Ersetze f durch $-f$]. OBdA gilt $f \geq 0$ [Ersetze f durch $\max(0, f)$]. OBdA $f \in B^+(X)$ mit $I(f) \in \mathbb{R}$ [Benutze die Definition der Integrierbarkeit]. Wegen $f \in B^+(X)$ gilt $\varphi_n \nearrow f$ für gewisse $\varphi_n \in B(X)$; oBdA $\varphi_n \geq 0$. Offensichtlich gilt $\Sigma \subset \bigcup_{n=1}^{\infty} \{x \in X \mid \varphi_n(x) > 0\}$, also $\Sigma = \bigcup_{n=1}^{\infty} \{x \in \Sigma \mid \varphi_n(x) > 0\}$. Damit ist Σ eine Nullmenge, denn $\{x \in \Sigma \mid \varphi(x) > 0\}$ ist für jedes $\varphi = \varphi_n$ eine Nullmenge! [Für fixes $\varphi(x) \geq 0$ in $B(X)$ und reelles $c > 0$ gilt $\Sigma \subset \Sigma(c) := \{x \in X \mid f(x) > c\varphi(x)\}$. Aus $h := 0 \leq \chi_{\Sigma} \cdot \varphi \leq c^{-1}f \cdot \chi_{\Sigma(c)} \leq c^{-1}f =: g$ und $I(g) - I(h) \leq I(f)/c$ folgt im Limes $c \rightarrow \infty$ dann $\chi_{\Sigma} \cdot \varphi \in L(X)$ mit $I(\chi_{\Sigma} \cdot \varphi) = 0$ (Fußnote zu Definition 6.2). Wegen $B^-(X) \ni 0 \leq \chi_{\Sigma_m} \leq m(\chi_{\Sigma} \cdot \varphi) \in B^+(X)$ ist damit $\Sigma_m := \{x \in \Sigma \mid \varphi(x) \geq 1/m\}$ eine Nullmenge für alle $m \in \mathbb{N}$, also auch die abzählbare Vereinigung $\{x \in \Sigma \mid \varphi(x) > 0\} = \bigcup_{m=1}^{\infty} \Sigma_m$.] $\square$

6.7 Messbare Funktionen

Sei $L(X) = L(X, B, I)$ der Verband der reellwertigen Lebesgue integrierbaren Funktionen. Wir machen jetzt folgende

Annahme: $B(X) \supset C_c(X)$.

Unter dieser Annahme machen wir folgende

Definition 6.18. Eine Funktion $f : X \rightarrow \mathbb{R}$ heisst **messbar**, wenn eine Folge $f_n \in C_c(X)$ existiert, welche **punktweise fast überall** (d.h. **ausserhalb einer Nullmenge**) gegen die Grenzfunktion f konvergiert

$$f_n \rightarrow f \quad (\text{f\"u}).$$

Wir betrachten später die Räume der $\mathbb{C}$ -wertigen Funktionen $L(X, \mathbb{C})$ resp. $C_c(X, \mathbb{C})$ etc. Damit sei immer gemeint, dass sowohl Real- als auch Imaginärteil der Funktion im entsprechenden Raum $L(X)$ resp. $C_c(X)$ liegen. Setze $I(u + iv) := I(u) + iI(v)$ für komplexes $f = u + iv$ aus $L(X, \mathbb{C})$.

Satz 6.19. Es gilt

1. Die messbaren Funktionen bilden einen Verband $M(X)$.
2. $M(X)$ ist eine $\mathbb{R}$ -Algebra und $M(X, \mathbb{C})$ eine $\mathbb{C}$ -Algebra (nicht notwendig mit 1).

6 Lebesgue Integration

3. $f \in M(X, \mathbb{C}) \implies \bar{f}$ und $|f| \in M(X, \mathbb{C})$.
4. Ist $f = \sup_{n \in \mathbb{N}} (f_n)$ f.ü. reellwertig und alle $f_n \in M(X)$, dann ist $f \in M(X)$ (nach Abänderung auf einer Nullmenge).
5. Ist $f: X \rightarrow \mathbb{C}$ ein punktwiser Limes $f_n \rightarrow f$ (f.ü) für $f_n \in M(X, \mathbb{C})$, so ist $f \in M(X, \mathbb{C})$.
6. Gilt $|f| \leq g$ fast überall und $f \in M(X, \mathbb{C})$ sowie $g \in L(X)$, dann ist $f \in L(X, \mathbb{C})$.

Beweis. Eigenschaft 1., 2. und 3. folgt sofort aus den Permanenzsätzen der Konvergenz und der entsprechenden Eigenschaft von $C_c(X)$. Punkt 4. benutzt die Abgeschlossenheit 1. unter endlichen Suprema (Infima) und folgt mittels des Diagonalfolgentricks und der Tatsache, daß abzählbare Vereinigungen von Nullmengen wieder Nullmengen sind. Eigenschaft 5. folgt dann unmittelbar aus 4. mit der Beweismethode von Satz 6.12.

Zur letzten Aussage 6. OBdA ist f reellwertig. Für $f = f_+ + f_-$ genügt es f_+ zu betrachten. Das heisst oBdA $f \geq 0$. Wähle $f_n \rightarrow f$ (f.ü) mit $f_n \in C_c(X)$. Ersetzt man f_n durch $\max(0, f_n)$, gilt oBdA $0 \leq f_n$. Dann gilt $\min(f_n, g) \rightarrow \min(f, g) = f$ (f.ü). Wegen $|\min(f_n, g)| \leq g$ ist g eine Majorante. Aus dem Satz von der dominierten Konvergenz folgt $f \in L(X)$, denn g und alle f_n und damit alle $\min(f_n, g)$ sind in $L(X)$ wegen der Verbandseigenschaft von $L(X)$. $\square$

Beispiel. Charakteristische Funktionen χ_A kompakter Teilmengen $A \subset \mathbb{R}^n$ sind messbare Funktionen in $M(\mathbb{R}^n, \mathbb{C})$ (siehe Beispiel 2.25).

Bemerkung. Für $B(X) = C_c(X)$ kann man leicht zeigen $L(X) \subset M(X, \mathbb{C})$.

Definition 6.20. Sei $\mathcal{L}^1(X)$ der Raum der messbaren und integrierbaren Funktionen auf X . Der Raum $L^1(X)$ ist der Quotient $\mathcal{L}^1(X)/(\text{Nullfunktionen})$ von $\mathcal{L}^1(X)$ nach dem Unterraum der Nullfunktionen, d.h. nach dem Raum der Funktionen die auf X (f.ü) verschwinden.

Definition 6.21. $L^1_{loc}(\mathbb{R}^n)$ ist der Raum aller messbaren Funktionen auf $\mathbb{R}^n$, für die $\chi_A \cdot f$ ein Element in $L^1(\mathbb{R}^n)$ definiert für alle kompakten Teilmengen $A \subset \mathbb{R}^n$.

Beispiel 6.22. Nach Satz 6.19(2)(6) gilt

$$f \in L^1_{loc}(\mathbb{R}^n), \varphi \in C_c(\mathbb{R}^n) \implies f \cdot \varphi \in L^1(\mathbb{R}^n) \subset L^1_{loc}(\mathbb{R}^n).$$

Beispiel 6.23. Wegen der Substitutionsformel⁵ gilt für die Funktion $f(x) = r^\alpha$

$$r^\alpha \in L^1_{loc}(\mathbb{R}^n) \iff \alpha > -n.$$

⁵In Polarkoordinaten (siehe auch Abschnitt 11.1) ist $\lim_{\varepsilon \rightarrow 0} \text{vol}(S^{n-1}) \int_\varepsilon^R r^\alpha r^{n-1} dr$ genau dann endlich, wenn gilt $\alpha + n - 1 > -1$.

7 Verallgemeinerte Funktionen

7.1 Basics

Für offene Teilmengen $M \subseteq \mathbb{R}^n$ nennen wir $C_c^\infty(M)$ den Raum¹ der **Testfunktionen** auf M . Eine **verallgemeinerte Funktion**² auf M ist eine $\mathbb{R}$ -Linearform (ein sogenanntes Funktional)

$$F : C_c^\infty(M) \longrightarrow \mathbb{R} .$$

Sei $\mathcal{F}_{gen}(M)$ der $\mathbb{R}$ -Vektorraum der verallgemeinerten Funktionen auf M .

Beispiele. a) $\mathbb{R}$ -Linearformen $I : C_c(M) \rightarrow \mathbb{R}$ definieren wegen $C_c^\infty(M) \subset C_c(M)$ verallgemeinerte Funktionen. b) Funktionen $f \in C(M)$ oder allgemeiner³ in $L_{loc}^1(\mathbb{R}^n)$ definieren wegen Beispiel 6.22 verallgemeinerte Funktionen $F = F_{f \cdot dx} \in \mathcal{F}_{gen}(M)$, vermöge

$$F_{f \cdot dx}(\varphi) := \int_M f(x) \varphi(x) dx \quad , \quad f \in L_{loc}^1(\mathbb{R}^n) .$$

Die so definierte verallgemeinerte Funktion $F = F_{f \cdot dx}$ heißt **glatt** im Fall $f \in C^\infty(M)$.

Träger. Wir sagen $F \in \mathcal{F}_{gen}(M)$ hat Träger in A , wenn das Komplement U von A offen in M ist und wenn gilt $F(\varphi) = 0$ für alle $\varphi \in C_c^\infty(U)$.

Multiplikation. Für $F \in \mathcal{F}_{gen}(M)$ und $f \in C^\infty(M)$ liegt $f \cdot F(\varphi) := F(f \cdot \varphi)$ in $\mathcal{F}_{gen}(M)$. Für $f \in C^\infty(M)$ gilt offensichtlich $f \cdot F_{g \cdot dx} = F_{(fg) \cdot dx}$. Hat F Träger in A , dann auch $f \cdot F$.

Lokalität. Ist U offen in M , definiert $C_c^\infty(U) \subseteq C_c^\infty(M)$ eine **Einschränkungsabbildung** $res_U^M : \mathcal{F}_{gen}(M) \rightarrow \mathcal{F}_{gen}(U)$. Für Überdeckungen $M = \bigcup_{i \in I} U_i$ durch offene Mengen U_i gilt

$$res_{U_i}^M(F) = 0 \quad , \quad \forall i \in I \quad \implies \quad F = 0 .$$

¹ $C_c^\infty(M)$ ist der Raum der unendlich oft differenzierbaren Funktionen auf M mit kompaktem Träger.

²diese Bezeichnung ist weit verbreitet; besser wäre die Bezeichnung verallgemeinerte Dichte! Die $\mathbb{R}$ -Linearformen auf den i -Formen $A_c^i(M)$ mit kompaktem Träger (Testformen) sind die $(n-i)$ -Ströme $\mathfrak{D}^{n-i}(M)$ auf M .

³Elemente in $L_{loc}^1(M)$ werden definiert durch Funktionen, für die gilt $\chi_K \cdot f \in \hat{L}(\mathbb{R}^n)$ für alle Kompakta $K \subset M$ (bezüglich des Standard Lebesgue Integrals).

7 Verallgemeinerte Funktionen

[Wähle eine C^∞ -Partition der Eins $\sum_i \psi_i = 1$ zu dieser Überdeckung (deren Existenz soll hier für $\#I = \infty$ nicht bewiesen werden). Dann ist $F(\varphi) = F(\sum_i \varphi \psi_i) = \sum_i F(\varphi \psi_i) = 0$ für alle $\varphi \in C_c^\infty(M)$, da nur endlich viele $\varphi \psi_i \in C_c^\infty(U_i)$ nicht Null sind.]⁴

Ableitungen. Die Ableitung $\partial_i F$ einer verallgemeinerten Funktion $F \in \mathcal{F}_{gen}(M)$ ist

$$\partial_i F(\varphi) := -F(\partial_i \varphi).$$

Hat F Träger in A , dann auch $\partial_i F$. Es gilt $\partial_i(f \cdot F) = (\partial_i f) \cdot F + f \cdot (\partial_i F)$. [Benutze dazu $\partial_i(f \cdot F)(\varphi) = -F(f \partial_i \varphi)$ und $(f \cdot (\partial_i F))(\varphi) = -F(\partial_i(f \varphi))$]. Für **glattes** $F = F_{f \cdot dx}$ gilt

$$\partial_i(F_{f \cdot dx}) = F_{(\partial_i f) \cdot dx}$$

[Überdecke M durch offene Quader U ; wegen Lokalität kann $M = U$ angenommen werden. Auf dem offenen Quader U benutze partielle Integration in dx_i ; der Randterm verschwindet!].

Direkte Bilder. Eine *stetige* Abbildung $g : M \rightarrow N$ heisst **eigentlich**, wenn Urbilder von kompakten Mengen wieder kompakt sind. Sei I monoton und $\mathbb{R}$ -linear auf $C_c(M)$ und g *stetig* und *eigentlich*. Dann definiert $(g_* I)(\varphi) := I(g^*(\varphi))$ eine verallgemeinerte Funktion $F = g_* I$ auf N wegen $g^*(C_c(N)) \subseteq C_c(M)$. Ist g *eigentlich* und *glatt*, d.h. eine C^∞ -Abbildung zwischen offenen Teilmengen M und N von Euklidischen Räumen, dann gilt $g^*(\varphi) = \varphi(g(x)) \in C_c^\infty(M)$ für $\varphi(y) \in C_c^\infty(N)$. Jede verallgemeinerte Funktion F auf M definiert dann eine verallgemeinerte Funktion $g_* F$ auf N vermöge $g_* F(\varphi) := F(g^*(\varphi))$.

Lemma 7.1. Für eigentliche glatte Abbildungen g gilt
$$g_*(\partial_i F) = \sum_j \frac{\partial g_j}{\partial x_i} \cdot \partial_j(g_* F).$$

Beweis. Beachte $\partial_j(g_* F)(\varphi) = -g_* F(\partial_j \varphi) = -F(\partial_j(\varphi)(g(x)))$. Aus der Kettenregel (Lemma 4.7) folgt daher $\sum_j \frac{\partial g_j}{\partial x_i} \cdot \partial_j(g_* F) = -F(\partial_i(\varphi(g(x))) = g_*(\partial_i F)$. $\square$

7.2 Distributionen

Sei $M \subseteq \mathbb{R}^n$ offen. Für $F \in \mathcal{F}_{gen}(M)$ und $\xi \in M$ betrachten wir dann die Menge aller $\alpha \in \mathbb{R}$, für die zu jedem gegebenem $\psi, \varphi \in C_c^\infty(\mathbb{R}^n)$ eine Konstante $C = C(\psi, \varphi)$ und ein $t_0 > 0$ existiert, so dass für alle $t \in (0, t_0)$ gilt⁵

$$|F\left(\psi(x)\varphi\left(\frac{x-\xi}{t}\right)\right)| \leq C(\psi, \varphi) \cdot t^\alpha.$$

⁴Seien umgekehrt $F_U \in \mathcal{F}_{gen}(U)$ gegeben für alle $U = U_i, i \in I$ einer Überdeckung $M = \bigcup_{i \in I} U_i$ so dass gilt $res_{U \cap V}^U(F_U) = res_{U \cap V}^V(F_V), \forall U, V \in I$. Dann existiert ein $F \in \mathcal{F}_{gen}(M)$ mit der Eigenschaft $res_U^M(F) = F_U, \forall U \in I$. [Setze $F(\varphi) := \sum_{i=1}^\infty F_{U_i}(\varphi \cdot \psi_i)$ für eine C^∞ -Partition der Eins $\sum_i \psi_i = 1$ mit $supp(\psi_i) \subset U_i \in I$. Offensichtlich ist F eine verallgemeinerte Funktion auf M mit $res_U^M(F) = F_U$].

⁵Beachte: $\psi(x)\varphi(\frac{x-\xi}{t}) \in C_c(\mathbb{R}^n)$ hat kompakten Träger in $\xi + t \cdot supp(\varphi)$. Dieser ist für $\xi \in M$ und genügend kleine t in M enthalten.

Ist die Menge leer, setzen wir $d_\xi(F) = -\infty$. Ansonsten sei $d_\xi(F) \in \{-\infty\} \cup \mathbb{R} \cup \{+\infty\}$ das Supremum der Menge all dieser α . Wir nennen $d_\xi(F) \in \hat{\mathbb{R}}$ die **lokale Ordnung** von F in ξ .

Definition. Eine verallgemeinerte Funktion auf M nennen wir **Distribution**⁶ auf M , wenn für jede kompakte Teilmenge $K \subset M$ und alle $\psi \in C(M)$ gilt

$$\inf_{\xi \in K} (d_\xi(F)) > -\infty.$$

Der Raum aller so definierten Distributionen ist ein $\mathbb{R}$ -Untervektorraum $\mathcal{F}(M)$ von $\mathcal{F}_{gen}(M)$.

Lemma 7.2. Jede monotone $\mathbb{R}$ -Linearform F auf $C_c(M)$ sowie jede Funktion $f \in L^1_{loc}(M)$ definiert eine Distribution $F_{f \cdot dx}$ auf M der Ordnung ≥ 0 . Ist f stetig, dann ist $F_{f \cdot dx}$ von der Ordnung $\geq n$. Ist F eine Distribution auf M , dann auch $\partial_i F$ sowie $g \cdot F$ für $g \in C^\infty(M)$. Hat F Ordnung $d_\xi(F) \geq d$ in ξ , dann hat $\partial_i F$ resp. $g \cdot F$ Ordnung $\geq d - 1$ resp. $\geq d$ in ξ .

Beweis. Für monotonen $\mathbb{R}$ -lineares F auf $C_c(M)$ gilt $d_\xi(F) \geq 0$, denn für $\psi, \varphi \in C_c^\infty(M)$ gilt $|F(\psi(x)\varphi(\frac{x-\xi}{t}))| \leq \|\varphi\| \cdot F(|\psi|) \leq C$, und C ist unabhängig von t . Sei nun $F = F_{f \cdot dx}$ und $f \in L^1_{loc}(\mathbb{R}^n)$. Für $f \geq 0$ ist F monoton und wir sind fertig. Allgemein gilt $f = f_+ - f_-$ mit $f_\pm \geq 0$ in $L^1_{loc}(\mathbb{R}^n)$. Aus $F = F_{f_+ \cdot dx} - F_{f_- \cdot dx}$ folgt daher⁷ $d_\xi(F) \geq 0$ für F in $L^1_{loc}(\mathbb{R}^n)$. Es gilt $d_\xi(G) \geq d_\xi(F)$ für $G = \psi \cdot F$ sowie $d_\xi(G) \geq d_\xi(F) - 1$ für $G = \partial_i F$. Für letzteres benutze $\partial_i(\psi(x)\varphi(\frac{x-\xi}{t})) = t^{-1} \cdot \psi(x)(\partial_i \varphi)(\frac{x-\xi}{t}) + (\partial_i \psi)(x)\varphi(\frac{x-\xi}{t})$. $\square$

Homogenität. Für homogenes $f \in L^1_{loc}(\mathbb{R}^n)$ mit⁸ $f(t \cdot x) = t^\alpha \cdot f(x)$ gilt $d_\xi(F_{f \cdot dx}) \geq \alpha + n$ für $\xi = 0$ wegen $|F_{f \cdot dx}(\psi(x)\varphi(\frac{x-\xi}{t}))| = |\int_{\mathbb{R}^n} f(t \cdot x)\psi(tx)\varphi(x)t^n dx| \leq \|\psi\| \|\varphi\| I(|f|) \cdot t^{\alpha+n}$.

Dirac Distribution. Für $\xi \in M$ definieren wir die **Dirac Distribution** $F = \delta_\xi$ durch

$$\varphi \mapsto \varphi(\xi).$$

Die Dirac Distributionen definieren abstrakte Integrale auf $C_c(M)$; siehe Definition 3.18. Der Träger der Dirac Distribution $F = \delta_\xi \in \mathcal{F}(M)$ ist in $\{\xi\}$ enthalten, und es gilt $d_\xi(F) = 0$.

Lemma 7.3. Ist der Träger von $F \in \mathcal{F}(M)$ in $\{\xi\}$ enthalten und $d_\xi(F) > 0$, dann ist $F = 0$.

Beweis. OBdA $\xi = 0$. Für $\psi \in C_c^\infty(\mathbb{R}^n)$ wähle $\varphi \in C_c^\infty(\mathbb{R}^n)$ mit $\varphi(x) = 1$ für alle $x \in \text{supp}(\psi)$. Dann gilt $\psi = \psi\varphi$ und $F(\psi) = F(\psi\varphi)$. Andererseits ist $F(\psi\varphi) = F(\psi(x)\varphi(\frac{x}{t}))$ für alle $t \in (0, 1)$, denn $\psi(x)\varphi(x) - \psi(x)\varphi(\frac{x}{t})$ hat Träger in $U = \mathbb{R}^n \setminus \{\xi\}$ und $\text{res}_U^{\mathbb{R}^n}(F) = 0$. Aus $d_\xi(F) > 0$ folgt daher $F(\psi) = F(\psi(x)\varphi(\frac{x}{t})) = \lim_{t \rightarrow 0} F(\psi(x)\varphi(\frac{x}{t})) = 0$. $\square$

⁶Vielleicht nennt man dann F besser ein **Feld** auf M , denn der hier betrachtete Begriff der Distribution ist möglicher Weise allgemeiner als der übliche Begriff der Distributionen im Sinne von Schwartz.

⁷Für $F = F_{f \cdot dx}$ und stetiges f gilt $d_\xi(F) \geq n$ wegen Lemma 3.9.

⁸Sei $M = \mathbb{R}^n$ oder $\mathbb{R}^n \setminus \{0\}$ und $g_t(x) = t^{-1}x$ für $t \in \mathbb{R}_{>0}$. Man sagt F ist *homogen vom Grad* $\alpha \in \mathbb{R}$, wenn $g_{t*}F = t^\alpha \cdot F$ gilt für alle $t > 0$ in $\mathbb{R}$. Ist F homogen vom Grad α , dann ist $\partial_i F$ homogen vom Grad $\alpha - 1$. Die Distribution δ_0 ist homogen vom Grad 0. Die Distribution $F_{r^\alpha \cdot dx}$ ist homogen vom Grad $n + \alpha$. Allgemeiner: Für $f \in L^1_{loc}(\mathbb{R}^n)$ mit $f(t \cdot x) = t^\alpha \cdot f(x)$ ist $F = F_{f \cdot dx}$ homogen vom Grad $\alpha + n$.

7 Verallgemeinerte Funktionen

Sei $P \in \mathcal{P}_\ell(\mathbb{R}^n)$ ein homogenes Polynom und $G = P(x - \xi) \cdot F$. Dann gilt (*)

$$d_\xi(G) \geq d_\xi(F) + \ell.$$

Zum Beweis ist oBdA $\xi = 0$. Benutze dann $|(x_i \cdot F)(\psi(x)\varphi(\frac{x}{t}))| = t \cdot |F(\psi(x)(x_i\varphi(\frac{x}{t})))| \leq t \cdot C(\psi, x_i\varphi)t^\alpha$ für alle $\alpha < d_0(F)$.

Korollar 7.4. Sei $\ell \in \mathbb{N}$. Jede Distribution F mit Träger in $\{\xi\}$ und Ordnung $d_\xi(F) \geq -\ell$ ist eine $\mathbb{R}$ -Linearkombination von Ableitungen vom Grad $\leq \ell$ der Dirac Distribution δ_ξ .

Beweis. OBdA ist $\xi = 0$. Aus Lemma 7.3 und (*) sowie $d_\xi(F) + \ell + 1 > 0$ folgert man (**): $P \cdot F = 0$ für alle $P \in \mathcal{P}_{\ell+1}(\mathbb{R}^n)$. Der Einfachheit⁹ halber sei $\ell = 0$. Wähle $\psi \in C_c^\infty(M)$ konstant 1 in einer kleinen Kugel um ξ . Für $\varphi \in C_c^\infty(M)$ gilt $\varphi(x)\psi(x) = \varphi(0) \cdot \psi(x) + \sum_{\nu=1}^n x_\nu \cdot \varphi_\nu(x)$ für gewisse $\varphi_\nu \in C_c(\mathbb{R}^n)$ [nach Lemma 5.26; multipliziere noch mit der Funktion $\psi \in C_c(M)$]. Weiterhin gilt $\varphi(x) - \varphi(x)\psi(x) = x_1 \cdot \varphi_0(x)$ für ein $\varphi_0 \in C_c^\infty(M)$, da die linke Seite in einer Kugel um ξ verschwindet. Es folgt $F(\varphi) = a_0 \cdot \delta_0(\varphi)$ für $a_0 := F(\psi)$, da denn im Fall $\ell = 0$ gilt wegen (**) $F(x_i \cdot \varphi_\nu(x)) = 0$ für $\nu = 0, 1, \dots, n$ und alle i . $\square$

7.3 Faltung

Seien $g, f \in C^\infty(\mathbb{R}^n)$. Besitzt eine der Funktionen f, g kompakten Träger auf $\mathbb{R}^n$, dann ist die **Faltung** $g * f \in C^\infty(\mathbb{R}^n)$ wohldefiniert durch

$$(g * f)(y) := \int_{\mathbb{R}^n} g(x)f(y-x)dx.$$

Die Substitution $x \mapsto y - x$ zeigt $f * g = g * f$. Analog zeigt man $(f * g) * h = f * (g * h)$, falls mindestens zwei der drei Funktionen f, g, h kompakten Träger besitzen. Aus dem Satz von der dominierten Konvergenz folgert man leicht $\partial_i(g * f) = g * (\partial_i f) = (\partial_i g) * f$.

Für $F \in \mathcal{F}_{gen}(\mathbb{R}^n)$ und $f \in C_c^\infty(\mathbb{R}^n)$ ist die Faltung $F * f$ erklärt als *Funktion* auf $\mathbb{R}^n$ (nota bene *nicht* als verallgemeinerte Funktion !) durch

$$(F * f)(y) := F(h_y) \quad , \quad h_y(x) := f(y-x).$$

Insbesondere gilt dann $F_{g \cdot dx} * f = \int_{\mathbb{R}^n} g(x)f(y-x)dx$ oder $F_{g \cdot dx} * f = g * f$.

Beispiel. Für die Dirac Distribution $F = \delta_0$ gilt $\delta_0 * f = f$.

Lemma 7.5. Für monotone $\mathbb{R}$ -Linearformen F auf $C_c(\mathbb{R}^n)$ und $f \in C_c^\infty(\mathbb{R}^n)$ gilt $F * f \in C^\infty(\mathbb{R}^n)$ sowie

$$\partial_i(F * f) = F * (\partial_i f).$$

Hat F ausserdem kompakten Träger, dann ist $F * f \in C_c^\infty(\mathbb{R}^n)$.

⁹Analog zeigt man $\varphi(x) = Q(x) \cdot \psi(x) + \sum_P P(x) \cdot \varphi_P(x)$ für ein Polynom Q vom Grad $\leq \ell$ und homogene Polynome P vom Grad $\ell + 1$ und gewisse $\varphi_P \in C_c^\infty(M)$. Es gilt $F(\varphi) = F(Q \cdot \psi)$ wegen (*). Der Raum der Distributionen mit der Eigenschaft (**) hat daher höchstens die Dimension $d = \sum_{i=0}^\ell \dim_{\mathbb{C}}(\mathcal{P}_i(\mathbb{R}^n))$. Der Unterraum aller partiellen Ableitungen von δ_0 bis zum Grad ℓ hat (!) die Dimension d . Daraus folgt die Behauptung.

Beweis. Ist der Träger von F in K' und der Träger von f in A enthalten, ist der Träger von $F * f$ in $A + K'$ enthalten. Der Einfachheit¹⁰ nehmen wir an K' sei kompakt. Wähle M offen mit $A + K' \subset M$ und K kompakt mit $M - A \subset K$. Wegen $F * f(y) = 0$ für $y \notin M$ ist die Differenzierbarkeit von $F * f(y)$ eine lokale Aussage in M . Wähle $\psi \in C_c^\infty(\mathbb{R}^n)$ identisch 1 auf K und sei $c_1 := F(\psi)$. Da F monoton und $\mathbb{R}$ -linear ist, gilt dann

$$|F(\varphi)| \leq c_1 \cdot \|\varphi\|_{C(K)} \quad , \quad \text{für alle } \varphi \in C_c^\infty(\mathbb{R}^n) \text{ mit Träger in } K .$$

Für $\xi, y \in M$ sei $a := y - \xi$. Dann gilt $|(F * f)(y) - (F * f)(\xi) - \sum_i (y - \xi)_i \cdot (F * \partial_i f)(\xi)| = |F(g(x))| \leq c_1 \cdot \|g\|_{C(K)}$ für $g(x) := f(y - x) - f(\xi - x) - a \cdot df(\xi - x)$, denn die Träger von $f(y - x)$, $f(\xi - x)$ und $a \cdot df(\xi - x)$ und damit $g(x)$ sind in $M - A \subset K$ enthalten. Es gilt¹¹

$$\|g\|_{C(K)} = \sup_{x \in K} \{ |f(y - x) - f(\xi - x) - a \cdot df(\xi - x)| \} \leq c_2(f) \cdot \|a\|_{\mathbb{R}^n}^2$$

und damit $|(F * f)(y) - (F * f)(\xi) - \sum_i (y - \xi)_i \cdot (F * \partial_i f)(\xi)| \leq c_1 c_2(K, f) \cdot \|a\|_{\mathbb{R}^n}^2$. Die rechte Seite ist von der Ordnung $o(a)$. Also ist $F * f$ differenzierbar im Punkt $\xi \in M$ mit den partiellen Ableitungen $F * \partial_i f$. Dies zeigt $\partial_i(F * f)(\xi) = (F * \partial_i f)(\xi)$ und rekursiv $F * f \in C^\infty(\mathbb{R}^n)$. $\square$

Aus $(\partial_i F) * f := F(-\frac{\partial}{\partial x_i} h_y)$ für $h_y(x) = f(y - x)$ sowie der Kettenregel $-\frac{\partial}{\partial x_i} h_y(x) = -\frac{\partial}{\partial x_i} f(y - x) = (\frac{\partial}{\partial y_i} f)(y - x) = g(y - x)$ für $g := \partial_i f$ folgt $(\partial_i F) * f = F * g = F * (\partial_i f)$.

Korollar 7.6. Für monotone $\mathbb{R}$ -lineare Funktionale¹² F auf $C_c(\mathbb{R}^n)$ und $f \in C_c^\infty(\mathbb{R}^n)$ gilt

$$\partial_i(F * f) = F * (\partial_i f) = (\partial_i F) * f ,$$

also $D(F * f) = D(F) * f$ für Differentialoperatoren D mit konstanten Koeffizienten. Gilt $DF = \delta_0$ (d.h. ist F eine **Fundamentallösung**), löst für $f \in C_c^\infty(\mathbb{R}^n)$ die Funktion

$$u := F * f \in C^\infty(\mathbb{R}^n)$$

die Differentialgleichung $D(u(x)) = f(x)$.

Beweis. $D(u) = D(F * f) = D(F) * f = \delta_0 * f = f$. $\square$

In den beiden nächsten Abschnitten wenden wir Korollar 7.6 an im Fall der zwei wichtigen Beispiele $D = \Delta$ (**Laplace Operator**) und $D = \square$ (**D'Alembert Operator**). Wir konstruieren in beiden Fällen Distributionen F (genauer gesagt sogar abstrakte Integrale) mit der Eigenschaft

$$D(F) = \delta_0 ,$$

sogenannte **Fundamentallösungen** F der durch D definierten Differentialgleichung.

¹⁰Jede Distribution F ist eine Summe $F = \sum_{i=0}^\infty F_i$ von Distributionen $F_i = \psi_i \cdot F$ mit Trägern in $K'_i := \{x \mid i \leq \|x\| \leq i + 2\}$ [benutze Partition der Eins $\sum_i \psi_i = 1$ mit $\text{supp}(\psi_i) \subset K'_i$]. Für alle $y \in K$ und ein festes Kompaktum K gilt $(F_i * f)(y) = 0$ für fast alle i und daher $(F * f) = \sum_i (F_i * f)(y)$, denn der Träger von $(F_i * f)(y)$ ist in dem Kompaktum $A + K'_i$ enthalten.

¹¹Beachte $r(t) := f(z + ta) - f(z) - ta \cdot df(z) = {}^T a \text{Hess}_\zeta(f) a$ für $\zeta = z + \eta \cdot a \in M$ mit $0 < \eta < t$, was man leicht mit Hilfe von $r(0) = r'(0) = 0$ durch zweimaliges Anwenden des Mittelwertsatzes zeigt (Lemma 4.10). Setze $z = \xi - x$ und $t = 1$. Beachte $c_2(f) = \max_{\|a\|_{\mathbb{R}^n}=1, \zeta \in \mathbb{R}^n} (|{}^T a \text{Hess}_\zeta(f) a|) < \infty$, da $\text{obdA } \zeta \in A$ ist.

¹²oder $F = F_{f \cdot dx}$ für $f \in L^1_{loc}(\mathbb{R}^n)$.

7.4 Coulomb Distribution

Beachte $f(x) = \frac{1}{r^{n-2}} \in L^1_{loc}(\mathbb{R}^n)$ (Beispiel 6.22) ist homogen vom Grad $2 - n$. Also ist $G = F_{r^{2-n}dx}$ eine Distribution (Lemma 7.2). Die Einschränkung von G auf $U = \mathbb{R}^n \setminus \{0\}$ ist glatt sowie homogen vom Grad $(2 - n) + n = 2$. Die Laplace Ableitung $F = -\Delta G$ ist daher homogen vom Grad Null. Es folgt $d_\xi(F) \geq 0$ für $\xi = 0$. Wegen $\Delta \frac{1}{r^{n-2}} = 0$ auf U (siehe Kapitel 5.9) gilt $\text{supp}(F) \subseteq \{0\}$. Dies zeigt $-\Delta G = a_0 \cdot \delta_0$ nach Korollar 7.4. Also

Lemma 7.7. Mit den Konstanten $a_0 = (n - 2) \cdot \text{vol}(S^{n-1})$ gilt in $\mathcal{F}(\mathbb{R}^n)$ für alle $n \geq 1$

$$-\Delta F_{r^{2-n}dx} = a_0 \cdot \delta_0$$

Beweis. Wir wissen bereits $a_0 \cdot \varphi(0) = \int r^{2-n}(-\Delta \varphi(x))dx$ für alle $\varphi \in C_c(\mathbb{R}^n)$. Durch einen Limeschluß gilt dies auch¹³ für $\varphi(x) = \exp(-r^2)$ mit $\Delta \varphi = (4r^2 - 2n)\varphi$. Dies gibt in Polarkoordinaten (Lemma 9.8) $a_0 = \text{vol}(S^{n-1}) \int_{r>0} (2n - 4r^2)e^{-r^2}rdr = (n - 2)\text{vol}(S^{n-1})$. $\square$

7.5 Wellengleichung

Wir betrachten den D'Alembert Operator auf $\mathbb{R}^{n+1}$ mit $t = x_{n+1}$

$$\square = -\partial_1^2 - \partial_2^2 - \partial_3^2 - \dots - \partial_n^2 + \partial_{n+1}^2, \quad (n \geq 2).$$

Definiere auf $C_c(\mathbb{R}^{n+1})$ monotone $\mathbb{R}$ -lineare Abbildungen I_+, I_- durch die Komposition $I_\pm = a \circ b_\pm$ der monotonen $\mathbb{R}$ -linearen Abbildungen a resp. $b = b_\pm$

$$C_c(\mathbb{R}^{n+1}) \xrightarrow{b_\pm} C_c(\mathbb{R}^n) \xrightarrow{a} \mathbb{R},$$

für $b_\pm : \varphi(x_1, \dots, x_n, x_{n+1}) \mapsto \varphi(x_1, \dots, x_n, \pm r)$ mit $r = \sqrt{x_1^2 + \dots + x_n^2}$, sowie dann $a : \psi \mapsto \int_{\mathbb{R}^n} \frac{\psi(x_1, \dots, x_n)}{(n-1)\text{vol}(S^{n-1})r^{n-2}} dx_1 \dots dx_n$. Nach Lemma 7.2 sind $I_\pm$ Distributionen auf $\mathbb{R}^{n+1}$, in symbolischer Notation

$$I_\pm = \frac{\delta(t \mp r)}{(n-1)\text{vol}(S^{n-1})r^{n-2}}.$$

Lemma 7.8 (Avancierte/Retardierte Potentiale). Für $n = 3$ gilt $\square I_\pm = \frac{1}{2} \cdot \delta_0$.

Der Beweis ist dem von Lemma 7.7 ähnlich, denn $I_\pm$ ist homogen vom Grad 2. Also ist $\square I_\pm$ homogen vom Grad Null. Der Beweis zerfällt wieder in die zwei Schritte:

¹³ Sei $\psi \in C_c^\infty(\mathbb{R}^n)$ mit $0 \leq \psi(x) \leq 1$ konstant 1 für alle Punkte x mit $\|x\| \leq 1$. Dann konvergiert $\varphi_m(x) = \psi(x/m) \exp(-r^2)$ gegen $\exp(-r^2)$. Es gilt $\int r^{2-n}(\Delta \exp(-r^2))dx = \lim_{m \rightarrow \infty} \int r^{2-n}(\Delta \varphi_m(x))dx$ und $\int r^{2-n}(\Delta \varphi_m(x))dx = a_0$ wegen $\varphi_m(0) = 1$ und $\varphi_m \in C_c^\infty(\mathbb{R}^n)$. Beachte dabei $|\int r^{2-n} \Delta((1 - \varphi(x/m)) \exp(-r^2))dx| \leq \sum_{|i| \leq 2} m^2 |\int_{|x| \geq m} b_i(x) \exp(-r^2)dx|$ für beschränkte stetige Funktionen $b_i(x)$ auf $\mathbb{R}^n$. Benutze nun $\lim_{m \rightarrow \infty} m^2 \int_{|x| \geq m} \exp(-r^2)dx = 0$.

1. $\square I_{\pm}$ hat Träger in $\{0\}$, und wie in Lemma 7.7 folgt daraus $\square I_{\pm} = a_0 \cdot \delta_0$.
2. Die Bestimmung der Konstante a_0 (in Polarkoordinaten).

Beweis. Die Abbildungen $g^{\varepsilon}(x_1, \dots, x_n) = (x_1, \dots, x_n, \varepsilon \cdot r)$ für $\varepsilon = \pm 1$

$$g^{\pm} : U = \mathbb{R}^n \setminus \{0\} \longrightarrow V = \mathbb{R}^{n+1} \setminus \{0\}$$

sind **eigentlich** und **glatt** auf U . Die Distribution

$$F := \frac{1}{(n-1)\text{vol}(S^{n-1})} F_{r^{2-n}.dx}$$

ist glatt auf $U = \mathbb{R}^n \setminus \{0\}$ und nach Definition ist $\text{res}_V^{\mathbb{R}^{n+1}}(I_{\pm})$ auf $V = \mathbb{R}^{n+1} \setminus \{0\}$ das direkte Bild $(g^{\pm})_*(F)$ unter $g = g^{\pm}$

$$\text{res}_V^{\mathbb{R}^{n+1}}(I_{\pm}) = g_*(F) .$$

Nach Lemma 7.1 gilt $\partial_i(g_*F) = g_*(\partial_i F) - \varepsilon \cdot \frac{x_i}{r} \cdot \partial_t(g_*F)$ für $i = 1, \dots, n$ [Beachte $\frac{\partial}{\partial x_i} g_j = 1$ für $i=j=1, \dots, n$ sowie $\frac{\partial}{\partial x_i} g_{n+1} = \varepsilon \cdot \frac{x_i}{r}$ für $1 \leq i \leq n$, und $\frac{\partial}{\partial x_i} g_j$ ist Null sonst].

Folgerung. Eine kleine Nebenrechnung¹⁴ zeigt daher auf $V = \mathbb{R}^{n+1} \setminus \{0\}$

$$\square g_*(F) = \varepsilon \cdot \frac{(3-n)}{(n-1)\text{vol}(S^{n-1})} \cdot \partial_t g_*(F_{r^{1-n}.dx}) \quad , \quad \varepsilon = \pm 1 .$$

Im Fall $n = 3$ verschwindet $\square g_*(F)$ auf V , d.h. $\square I_{\pm}$ hat Träger im Punkt 0. Dies zeigt Schritt 1. Wertet man $a_0 \cdot \delta_0 = \square \left(\frac{\delta(t \mp r)}{(n-1)\text{vol}(S^{n-1})r^{n-2}} \right)$ aus für die Funktion $\varphi(x, t) = \exp(-r^2 - t^2)$, mit einem Limeschluß wie im Beweis von Lemma 7.7, folgt $a_0 = 1/2$. $\square$

Auf Grund der Vorzeichen $\varepsilon = \pm 1$ heben sich bei der Addition von I_+ und I_- für alle $n \geq 2$ die im obigen Beweis auftretenden Terme $\pm \square g_*(F)$ weg. Es folgt

Lemma 7.9. Für alle $n \geq 2$ gilt $\boxed{\square(I_+ + I_-) = \delta_0}$.

Korollar 7.10. Für Funktionen $\rho \in C_c^\infty(\mathbb{R}^n)$ resp. $g \in C_c^\infty(\mathbb{R}^{n+1})$ sind die Gleichungen

$$\boxed{-\Delta u = \rho \quad \text{resp.} \quad \square u = g}$$

lösbar durch Funktionen $u \in C^\infty(\mathbb{R}^n)$.

¹⁴Für $H = g_*F$ folgt $\square H = \partial_t^2 H - \sum_{i=1}^n \partial_i(g_*(\partial_i F) - \varepsilon \frac{x_i}{r} \cdot \partial_t H)$. Dies ist die Summe von $\partial_t^2 H$ und $-g_*(\Delta F) + \varepsilon \sum_{i=1}^n \frac{x_i}{r} \cdot \partial_t g_*(\partial_i F)$ (beachte $\Delta F = 0$ auf $U = \mathbb{R}^n \setminus \{0\}$ nach Lemma 7.7) sowie $\varepsilon \sum_{i=1}^n \partial_i(\frac{x_i}{r} \cdot \partial_t H) = \varepsilon \sum_{i=1}^n \partial_i(\frac{x_i}{r}) \cdot \partial_t H + \varepsilon \sum_{i=1}^n \frac{x_i}{r} \cdot \partial_t \partial_i H$. Nun ist $\varepsilon \sum_{i=1}^n \frac{x_i}{r} \cdot \partial_t \partial_i H = \varepsilon \sum_{i=1}^n \frac{x_i}{r} \cdot \partial_t g_*(\partial_i F) - \varepsilon^2 \sum_{i=1}^n (\frac{x_i^2}{r^2}) \cdot \partial_t^2 H = \varepsilon \sum_{i=1}^n \frac{x_i}{r} \cdot \partial_t g_*(\partial_i F) - \partial_t^2 H$. Zusammenfassen aller Terme gibt damit auf $V = \mathbb{R}^{n+1} \setminus \{0\}$

$$\square H = \varepsilon \cdot \partial_t \left(2 \sum_{i=1}^n \frac{x_i}{r} \cdot g_* \partial_i F + \sum_{i=1}^n \partial_i \left(\frac{x_i}{r} \right) \cdot g_* F \right) = \varepsilon \cdot \partial_t g_* \left(2 \sum_{i=1}^n \frac{x_i}{r} \cdot \partial_i F + \sum_{i=1}^n \partial_i \left(\frac{x_i}{r} \right) \cdot F \right) .$$

Beachte, $F|_U$ ist $F_{r^{2-n}.dx}$ bis auf die Konstante $(n-1)\text{vol}(S^{n-1})$. Benutze dann zur Berechnung der Klammer $2 \sum_{i=1}^n \frac{x_i}{r} \cdot \partial_i(r^{2-n}) + \sum_{i=1}^n \partial_i(\frac{x_i}{r}) \cdot r^{2-n} = (3-n)r^{1-n}$ auf U .

8 Hilberträume

8.1 Vorbemerkung

In der **Quantentheorie** betrachtet man (in der Regel unendlich dimensionale) $\mathbb{C}$ -Vektorräume V , welche mit einer *positiv definiten Hermiteschen Bilinearform* versehen sind: Das heisst, auf V existiert eine $\mathbb{R}$ -bilineare Form

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{C}$$

mit $\langle \lambda v, \mu w \rangle = \bar{\lambda} \mu \langle v, w \rangle$ für $\lambda, \mu \in \mathbb{C}$ und $v, w \in V$ sowie

$$\langle v, w \rangle = \overline{\langle w, v \rangle},$$

so daß weiterhin $\langle v, v \rangle$ reell und > 0 ist für alle $v \neq 0$. Außerdem wird angenommen, daß V *vollständig* ist bezüglich der Metrik $d(v, w) = \|v - w\|$ auf V , welche durch $\|v\|^2 = \langle v, v \rangle$ definiert wird. Einen solchen Vektorraum V nennt man einen **Hilbertraum**.

Physikalisch betrachtet definiert V den sogenannten **Zustandsraum**: Zustände sind dabei die komplexen Geraden $\mathbb{C} \cdot v$ (komplex lineare Unterräume der Dimension 1 in V) mit $v \neq 0$. Da sich der Beweis der Schwarz Ungleichung überträgt, erfüllt die reelle Zahl

$$W(v, w) = \frac{|\langle v, w \rangle|^2}{\|v\|^2 \|w\|^2}$$

die Ungleichung

$$0 \leq W(v, w) \leq 1.$$

Diese Zahl hängt nur von den Zuständen $\mathbb{C}v$ und $\mathbb{C}w$ ab. Physikalisch wird $W(v, w)$ gedeutet als Wahrscheinlichkeit dafür, daß der Zustand $\mathbb{C}v$ den Zustand $\mathbb{C}w$ enthält (oder in ihn übergeht). Eine Bijektion der Menge aller Zustände, welche alle Übergangswahrscheinlichkeiten erhält, nennt man einen Automorphismus des Zustandsraumes. Jede **unitäre** lineare Abbildung $L : V \rightarrow V$ liefert einen solchen Automorphismus, wobei $\mathbb{C}$ -lineare Abbildungen L unitär genannt werden wenn $\langle L(v), L(w) \rangle = \langle v, w \rangle$ für alle Vektoren v, w aus V gilt.

Von besonderer **mathematischer Bedeutung** sind **anti-hermitesche** $\mathbb{C}$ -lineare Abbildungen $X : V \rightarrow V$, also Abbildungen¹ mit der Eigenschaft $\langle X(v), w \rangle = -\langle v, X(w) \rangle$, denn sie definieren eine **Lie Algebra**: Sind X und Y anti-hermitesch, dann ist auch $[X, Y] := XY - YX$

¹Ist V endlich dimensional, so daß $U_t = \exp(tX)$ definiert ist, dann ist U_t unitär nach Lemma 4.43.

anti-hermitesch. Ein antihermitescher Operator X kann physikalisch wie folgt mit **Messungen** in Verbindung gebracht werden. Ist X anti-hermitesch, dann ist $\frac{X}{2\pi i}$ ein **hermitescher** Operator

$$X \text{ anti-hermitesch} \iff \frac{X}{2\pi i} \text{ hermitesch,}$$

und hermitesche Operatoren haben reelle Eigenwerte. Die **hermiteschen** Operatoren $\frac{X}{2\pi i}$ bilden keine Liealgebra, aber ihre Eigenwerte haben **physikalische Bedeutung**, in konkreten Fällen etwa als Impuls, Ort, Energie etc. (jeweils abhängig vom Operator X). Zustände $\mathbb{C}v$, die Eigenräume des Operators $\frac{X}{2\pi i}$ zum Eigenwert λ definieren, haben den wohldefinierten Messwert λ (Impuls, Ort, Energie resp.). Beliebige Zustände $\mathbb{C}v$ haben keinen wohldefinierten Messwert, aber der statistische **Erwartungswert** der Messung ist die reelle (!) Zahl

$$E(X, v) := \frac{\langle \frac{X}{2\pi i} v, v \rangle}{\|v\|^2}.$$

Die **Standardabweichung** $\sigma(X, v)$ vom Erwartungswert $E(X, v)$ der Messung ist dann

$$\sigma(X, v) := \frac{\|\frac{X}{2\pi i} v\|}{\|v\|} = \frac{\|\frac{X}{2\pi i} v - E(X, v)v\|}{\|v\|}$$

für den normalisierten (auf den Erwartungswert zentrierten) Operator

$$\tilde{X} := X - 2\pi i E(X, v) \cdot id_V.$$

Beachte $[\tilde{X}, \tilde{Y}] = [X, Y]$. Ein System von antihermiteschen Operatoren X_ν (Messungen) für $\nu = 1, \dots, n$ heisst kohärent, wenn alle Operatoren X_ν und damit auch alle normalisierten Operatoren $\tilde{X}_\nu$ miteinander kommutieren.

Die **Heisenberg Unschärferelation**. Seien X und Y zwei 'Messungen' zu Operatoren X und Y , die nicht miteinander kommutieren sondern die Heisenberg Kommutatorrelation² erfüllen

$$[X, Y] = 2\pi i \cdot h \cdot id_V$$

für eine Konstante h (das **Wirkungsquantum**). Sei $\hbar = \frac{h}{2\pi}$ (bei unserer Normierung ist $h = 1$). Dann gilt für jeden Zustand $\mathbb{C}v$ (oBdA mit $\|v\| = 1$) die Unschärferelation

$$\sigma(X, v) \cdot \sigma(Y, v) \geq \frac{1}{2} \hbar.$$

Beweis. Durch Übergang zu den normalisierten Operatoren kann man obdA annehmen $X = \tilde{X}$ und $Y = \tilde{Y}$. Es folgt $2\pi h = |\langle 2\pi i h v, v \rangle| = |\langle (XY - YX)v, v \rangle| = |-\langle Y(v), X(v) \rangle + \langle X(v), Y(v) \rangle| \leq \|Y(v)\| \|X(v)\| + \|X(v)\| \|Y(v)\| = 2\|X(v)\| \|Y(v)\| = 2(2\pi)^2 \sigma(X, v) \sigma(Y, v).$

□

²Für dieses berühmte Beispiel siehe Heisenberg, Physikalische Prinzipien der Quantentheorie, Seite 14

$$X(f) = h \cdot \frac{d}{dx} f(x)$$

$$Y(f) = 2\pi i x \cdot f(x),$$

wobei $\frac{X}{2\pi i}$ und $\frac{Y}{2\pi i}$ zu Impuls p resp. Ort q korrespondieren und $\frac{[X, Y]}{2\pi i}$ zur Poissonklammer $\{p, q\} = 1$ in Planckkoordinaten, d.h. für $h = 1$ (nicht $\hbar = 1$). Siehe auch Abschnitt 8.9.

8.2 L^2 -Räume

Für komplexwertige Funktionen $f, g : X \rightarrow \mathbb{C}$ gelten die trivialen Abschätzungen

1. $|\bar{f} \cdot g| \leq \max(|f|, |g|)^2 = \max(|f|^2, |g|^2)$
2. $|f + g|^2 \leq (|f| + |g|)^2 \leq 4 \cdot \max(|f|, |g|)^2 = 4 \cdot \max(|f|^2, |g|^2).$

Sei X sei ein metrischer Raum. Wir nehmen an $C_c(X, \mathbb{C}) \subset L(X, \mathbb{C})$ gelte für ein abstraktes Integral I auf $B(X)$. Wir definieren dann $\mathcal{L}^2(X, I) = \mathcal{L}^2$ als einen Teilraum aller *messbaren Funktionen* $M(X, \mathbb{C})$ wie folgt

$$\mathcal{L}^2(X, I) = \left\{ f \in M(X, \mathbb{C}) \mid |f|^2 \in L(X) \right\}.$$

Lemma 8.1. $\mathcal{L}^2$ ist ein $\mathbb{C}$ -Vektorraum.

Beweis. $M(X, \mathbb{C})$ ist ein $\mathbb{C}$ -Vektorraum. Daher folgt die Behauptung aus der zweiten obigen Abschätzung mit Hilfe der Eigenschaften 6.19.6 und 6.19.2. $\square$

Analog zeigt dieses Argument mit Hilfe der ersten trivialen Abschätzung die Aussage : $\bar{f}g \in L(X)$ für $f, g \in \mathcal{L}^2$. Somit definiert

$$\langle f, g \rangle = I(\bar{f} \cdot g)$$

für $f, g \in \mathcal{L}^2$ eine positiv semidefinite hermitesche Bilinearform

$$\langle \cdot, \cdot \rangle : \mathcal{L}^2 \times \mathcal{L}^2 \rightarrow \mathbb{C}.$$

Wie in Satz 1.7 folgen daraus die beiden nächsten Lemmata für

$$\|f\| := \sqrt{\langle f, f \rangle}$$

Lemma 8.2. Es gilt die Schwarzungleichung und die Dreiecksungleichung. Gleichheit wird nur für proportionale Vektoren angenommen.

Lemma 8.3. $\|f\| = 0 \iff \text{Träger von } f \text{ ist eine Nullmenge.}$

Beweis. $\Leftarrow$ klar. $\Rightarrow$: Für $n \in \mathbb{N}$ gilt $\text{vol}\{x \in X \mid |f|^2 \geq 1/n\} = 0$ wegen der Abschätzung $\frac{1}{n} \cdot \text{vol}(\cdot) \leq \|f\|^2$ und $\|f\| = 0$. Nach Beppo Levi folgt im Limes $\text{vol}\{x \in X \mid |f| > 0\} = 0$. $\square$

Die Nullfunktionen (Funktionen in $\mathcal{L}^2$, die ausserhalb einer Nullmenge verschwinden) bilden einen $\mathbb{C}$ -linearen Untervektorraum von $\mathcal{L}^2$. Der Quotientenraum sei $L^2 = \mathcal{L}^2 / \{\text{Nullfunktionen}\}$. Die Werte $\|f\| = \|f\|_{L^2}$ und $\langle f, g \rangle$ hängen offensichtlich nur von der Äquivalenzklasse von f und g in L^2 ab. Aus Lemma 8.2 und Lemma 8.3 und dem nächsten Paragraphen folgt

Korollar 8.4. $(L^2(X, I), \|\cdot\|_{L^2})$ ist ein Hilbertraum.

8.3 Satz von Fischer-Riesz

Satz 8.5. $(L^2(X, I), \|\cdot\|_{L^2})$ ist vollständig und damit ein Hilbertraum.

Beweis. Alle nun auftretenden Funktionen f_n, f, g_n, g sind messbar, somit diskutieren wir nur die Integrierbarkeit. Wir müssen zeigen, dass jede Cauchyfolge f_n in $L^2(X)$ konvergiert. *Präparation.* Durch Übergang zu einer Teilfolge gilt oBdA

$$\|f_n - f_m\|_{L^2} \leq 2^{-\min(n,m)}.$$

Majorantenfolge. Wir betrachten die monotone Folge von Hilfsfunktionen $g_n \in L^2(X)$

$$g_1 = 0, \quad g_n := |f_1| + |f_1 - f_2| + \cdots + |f_n - f_{n-1}| \quad \text{für } n \geq 2$$

und ihren Limes $g_n \nearrow g$ (mit Werten in $\mathbb{R}^+$). Wir zeigen nun die Beschränktheit von $I(|g_n|^2)$ und damit $g \in \mathcal{L}^2$ (wegen Beppo Levi) mit Hilfe der Dreiecksungleichung und unserer Präparation:

$$I(|g_n|^2)^{\frac{1}{2}} = \|g_n\|_{L^2} \leq \|f_1\|_{L^2} + \sum_{n \geq 2} \|f_n - f_{n-1}\|_{L^2} \leq \|f_1\|_{L^2} + 1 < \infty.$$

Punktweise Konvergenz. Wegen $g_n(x) \rightarrow g(x)$ punktweise f.ü. (im Komplement der Unendlichkeitsstellen der g_n, g , die nach Lemma 6.17 eine Nullmenge bilden) ist $g_n(x) \in \mathbb{R}$ für festes $x \in X$ eine reelle Cauchyfolge f.ü. Dasselbe gilt für $f_n(x)$ wegen

$$|f_m(x) - f_n(x)| \leq \sum_{k=n+1}^m |f_k(x) - f_{k-1}(x)| = g_m(x) - g_n(x) \quad , \quad m > n.$$

Es existiert daher eine Funktion $f : X \rightarrow \mathbb{R}$ mit $f_n \rightarrow f$ punktweise f.ü.

Majorante. Aus $|f_n| \leq |f_1| + \sum_{k=2}^n |f_k - f_{k-1}| \leq g_n \leq g$ für alle $n \geq 2$ folgt $|f_n|^2 \leq g^2$ und damit $|f|^2 \leq g^2$ sowie dann $|f - f_n|^2 \leq (2g)^2$. Aus $M(X) \ni |f - f_n|^2 \leq (2g)^2 \in L(X)$ folgt $|f|^2 \in L(X)$ und $|f - f_n|^2 \in L(X)$ nach Satz 6.19.6. Wegen $|f_n - f|^2 \rightarrow 0$ punktweise f.ü. und $L(X) \ni |f - f_n|^2 \leq (2g)^2 \in L(X)$ folgt aus dem Satz 6.12 von der dominierten Konvergenz

$$\lim_n I(|f - f_n|^2) = I(\lim_n |f - f_n|^2) = 0.$$

Dies zeigt $\lim_n \|f - f_n\|_{L^2} = 0$ und damit unsere Behauptung. □

Wir nehmen im Rest dieses Abschnittes an: $B(X) \supseteq C_c(X)$, und für jede Funktion $h \in B(X)$ und jedes $\varepsilon > 0$ existiere ein $g \in C_c(X)$ mit $I(|h - g|) < \varepsilon$. (In der Tat gilt in den uns interessierenden Fällen sogar immer $B(X) = C_c(X)$.) Unter diesen Annahmen gilt

Lemma 8.6. $C_c(X, \mathbb{C})$ liegt dicht in $(L^2(X, I), \|\cdot\|_{L^2})$.

Beweis. *Schritt 1.* Nach Beppo Levi gilt $\lim_n (\min(n, f)) \rightarrow f$ in $L^2(X)$. Man kann also f oBdA durch $\tilde{f} = \min(n, f)$ ersetzen. Somit ist oBdA f durch eine Konstante C nach oben beschränkt, und dann analog auch durch $-C$ nach unten. Sei also oBdA $|f| \leq C$.

Schritt 2. Für $\varepsilon > 0$ existieren nach Definition $h \in B^-(X)$ und $h_n \in B(X) = C_c(X)$ mit $h_n \searrow h$ und $I(|f - h|) \leq I(|g - h|) < \varepsilon$ so dass für $n \gg 0$ gilt $I(|h - h_n|) < \varepsilon$ und damit $I(|f - h_n|) < 2\varepsilon$. Dies bleibt richtig, wenn man h, h_n nach oben und unten bei $\pm C$ abschneidet. Sei also oBdA $|h_n| \leq C$ und damit $|f - h_n| \leq 2C$. Es folgt

$$\|f - h_n\|_{L^2}^2 = I(|f - h_n|^2) \leq (2C) \cdot I(|f - h_n|) < (2C)2\varepsilon.$$

□

8.4 Der Folgenraum $L^2(\mathbb{Z})$

Die Teilmenge $\mathbb{Z}$ der reellen Zahlen $\mathbb{R}$ versehen mit der Einschränkung der Euklidischen Metrik hat folgende Eigenschaften: 1) Jede Funktion $f: \mathbb{Z} \rightarrow \mathbb{C}$ ist stetig, 2) Eine Teilmenge $A \subseteq \mathbb{Z}$ ist genau dann kompakt, wenn sie endlich ist. Somit gilt

$$C_c(X, \mathbb{C}) = \{a: \mathbb{Z} \rightarrow \mathbb{C} \mid a(n) = 0 \text{ für fast alle } n\}.$$

Offensichtlich definiert daher auf $C_c(\mathbb{Z})$ die normale (endliche !) Summe

$$I(a) = \sum_{n \in \mathbb{Z}} a(n)$$

eine $\mathbb{R}$ -lineare monotone Abbildung $C_c(\mathbb{Z}, \mathbb{R}) \rightarrow \mathbb{R}$. Wegen Satz 3.17 ist I sogar halbstetig, d.h. ein abstraktes Integral auf $C_c(X)$. Der Raum der bezüglich $C_c(\mathbb{Z}), I$ messbaren Funktionen $M(X, \mathbb{C})$ ist dann der Raum aller Folgen $a: \mathbb{Z} \rightarrow \mathbb{C}$, der Raum $L(X, \mathbb{C})$ ist der Unterraum aller *absolut konvergenten* Folgen. Weiterhin ist die leere Menge die einzige Nullmenge. Daher gilt

$$L^2(X, I) = \{a: \mathbb{Z} \rightarrow \mathbb{C} \mid \sum_{n \in \mathbb{Z}} |a(n)|^2 < \infty\}$$

sowie

$$\langle a, b \rangle = \sum_{n \in \mathbb{Z}} \overline{a(n)} \cdot b(n).$$

Für $\nu \in \mathbb{Z}$ definieren wir die Funktionen $f_\nu \in C_c(\mathbb{Z})$ durch

$$f_\nu(n) = \delta_{\nu n} \quad (\text{Kronecker-Delta}).$$

Diese haben die Eigenschaften

- $\langle f_\nu, f_\mu \rangle = \delta_{\nu\mu}$ (**Orthogonalität**).
- Für alle $f \in L^2(\mathbb{Z})$ und für alle $\varepsilon > 0$ existiert eine endliche $\mathbb{C}$ -Linearkombination $g \in C_c(X, \mathbb{C})$ der f_ν mit $\|f - g\|_{L^2} < \varepsilon$ (**L^2 -Dichtigkeit**).

8.5 Orthonormalbasen

Wir erinnern: Ein **Hilbertraum** $(H, \langle \cdot, \cdot \rangle)$ ist vollständig als metrischer Raum, wobei die Metrik $d(v, w) = \|v - w\|$ durch das positiv definite hermitesche Skalarprodukt $\langle \cdot, \cdot \rangle : H \times H \rightarrow \mathbb{C}$ gegeben ist

$$\|v\|^2 = \langle v, v \rangle.$$

Definition 8.7. Eine (abzählbare) **Hilbertraum-Basis**³ v_n ($n \in \mathbb{Z}$) eines Hilbertraumes H ist eine Folge von Vektoren $v_n \in H$ mit der Eigenschaft

- **Orthonormalität:** $\langle v_\nu, v_\mu \rangle = \delta_{\nu\mu}$ (Kronecker Delta)
- **Dichtigkeit:** Für alle $f \in H$ und für alle $\varepsilon > 0$ existiert eine endliche $\mathbb{C}$ -lineare Kombination g der v_ν so dass gilt $\|f - g\| < \varepsilon$.

Die erste Eigenschaft impliziert die lineare Unabhängigkeit der Vektoren v_ν : Verschwindet $g = \sum_n a(n)v_n$ (endliche Summe), dann folgt $a(m) = \langle g, v_m \rangle = 0$ für alle m .

Satz 8.8. Ist v_ν eine abzählbare Hilbertraum-Basis eines Hilbertraumes $(H, \langle \cdot, \cdot \rangle)$, dann induziert $L^2(\mathbb{Z}) \ni a \mapsto \sum_n a(n) \cdot v_n$ einen isometrischen Isomorphismus von Hilberträumen

$$i : L^2(\mathbb{Z}) \cong H, \quad \langle a, b \rangle_{L^2(\mathbb{Z})} = \langle i(a), i(b) \rangle.$$

Die Umkehrabbildung ordnet einem Vektor $v \in H$ die Folge $a(n) = \langle v_n, v \rangle$ in $L^2(\mathbb{Z})$ zu.

Beweis. Da die v_n linear unabhängig sind, ist die Abbildung i auf dem Teilraum $C_c(\mathbb{Z}) \subseteq L^2(\mathbb{Z})$ wohldefiniert (fast alle $a(n)$ sind Null), und für $a \in C_c(\mathbb{Z})$ gilt wegen der Orthonormalität

$$\|i(a)\|^2 = \left\langle \sum_n a(n) \cdot v_n, \sum_m a(m) \cdot v_m \right\rangle = \sum_n |a(n)|^2 = \|a\|_{L^2(\mathbb{Z})}^2.$$

Wegen dieser Isometrieeigenschaft lässt sich i durch Limesbildung auf ganz $L^2(\mathbb{Z})$ fortsetzen, denn $C_c(\mathbb{Z})$ liegt dicht in $L^2(\mathbb{Z})$: Für $a \in L^2(\mathbb{Z})$ wähle eine im L^2 -Sinn konvergente Folge $a_n \rightarrow a$ mit $a_n \in C_c(\mathbb{Z})$. Wegen $\|i(a_n) - i(a_m)\| = \|i(a_n - a_m)\| = \|a_n - a_m\|$ ist dann $i(a_n)$ eine Cauchyfolge in H , besitzt also einen Grenzwert und dieser definiert $i(a)$. Die so definierte Abbildung $i : L^2(\mathbb{Z}) \rightarrow H$ ist linear [$i(a+b) := \lim_n i(a_n + b_n) = \lim_n i(a_n) + \lim_n i(b_n) = i(a) + i(b)$ für $L^2(\mathbb{Z})$ -konvergente Folgen $a_n \rightarrow a$ und $b_n \rightarrow b$.] Wegen $\|i(a)\| = \lim_n \|i(a_n)\| = \lim_n \|a_n\| = \|a\|$ ist i eine Isometrie, also injektiv: $i(a) = 0 \implies \|a\| = \|i(a)\| = 0 \implies a = 0$. Ausserdem folgt, daß das Bild $B := i(L^2(\mathbb{Z}))$ abgeschlossen in H ist. [Sei $B \ni v_n \rightarrow v$ eine konvergente Folge und sei $v_n = i(w_n)$. Dann definiert w_n wegen $\|v_n - v_m\| = \|i(w_n) - i(w_m)\| = \|i(w_n - w_m)\| = \|w_n - w_m\|$ eine Cauchy Folge, welche wegen der Vollständigkeit von $L^2(\mathbb{Z})$ konvergiert: $w_n \rightarrow w$. Es folgt sofort $i(w) = v$.]

³Es handelt sich hierbei nicht um eine Basis im Sinne der Linearen Algebra.

Wegen der Dichtigkeits-Annahme enthält das Bild B einen dichten Teilraum, d.h. für $f \in H$ existiert ein $g \in B$ mit $\|f - g\| < \varepsilon$. Somit gibt es eine Folge von $g_n \in B$ mit $\|f - g_n\| < 1/n$, die also gegen f konvergiert. Da B abgeschlossen ist, folgt $f \in B$. Daher ist i surjektiv. $\square$

Korollar 8.9. Ist $\{v_n\}_{n \in \mathbb{Z}}$ eine (abzählbare) Hilbertraum-Basis eines Hilbertraumes V , dann gilt für jedes $v \in V$ (die Summe ist dabei definiert im Sinne der Hilbertraum-Konvergenz)

$$\boxed{v = \sum_{n \in \mathbb{Z}} a(n) \cdot v_n} \quad , \quad \boxed{a(n) = \langle v_n, v \rangle} .$$

Bemerkung. Physiker schreiben dies oft symbolisch in der bra-cket Form $|v\rangle = \sum_n |v_n\rangle \cdot \langle v_n | v \rangle$.

8.6 Fourier Reihen

Mittels der Parametrisierung $t \mapsto \exp(2\pi it)$ entsprechen Funktionen g auf dem Einheitskreis $X = S^1$ periodischen Funktionen $f(t) = g(\exp(2\pi it))$ auf $\mathbb{R}$ mit der Periode 1 und umgekehrt. Der Raum $C_c(X, \mathbb{C}) = C(X, \mathbb{C})$ kann dabei mit dem Raum der **periodischen** stetigen Funktionen $f : \mathbb{R} \rightarrow \mathbb{C}$ mit Periode $f(t+1) = f(t)$ identifiziert werden. Das Integral

$$I(f) = \int_0^1 f(t) dt$$

definiert ein abstraktes Integral auf $C_c(X)$. Sei I das zugehörige Lebesgue Integral. In diesem Sinne gilt

$$L^2(S^1, I) \cong L^2([0, 1], \mathbb{C}) \cong L^2_{\text{periodisch}}(\mathbb{R}, \mathbb{C}) .$$

Satz 8.10. Die Funktionen $\chi_n(t) = \exp(2\pi int)$ definieren eine Hilbertraum-Basis von $L^2(S^1, I)$. Das heisst: Jede Funktion $f \in L^2(S^1)$ schreibt sich als L^2 -Limes

$$\boxed{f(t) = \sum_{n \in \mathbb{Z}} a(n) \exp(2\pi int)}$$

mit den **Fourierkoeffizienten**

$$\boxed{a(n) = \int_0^1 f(t) \exp(-2\pi int) dt}$$

und es gilt die **Plancherel Formel**

$$\boxed{\sum_{n \in \mathbb{Z}} |a(n)|^2 = \int_0^1 |f(t)|^2 dt = \|f\|_{L^2}^2 < \infty} .$$

Die Fourierreihe $\sum_{n \in \mathbb{Z}} a(n) \exp(2\pi int)$ konvergiert **punktweise** gegen $f(t)$ für alle Funktionen $f \in C^2([0, 1], \mathbb{C})$ mit der Eigenschaft $f(0) = f(1)$ und $f'(0) = f'(1)$.

Beweis. Nach 8.5 genügt es zu zeigen, dass die Funktionen $\chi_n(t) = \exp(2\pi int)$ für $n \in \mathbb{Z}$ eine Hilbertraumbasis von $L^2(S^1, I)$ bilden. Die Plancherel Formel folgt dann aus der Existenz des isometrischen Isomorphismus $i : L^2(\mathbb{Z}) \cong L^2(S^1, I)$ nach Satz 8.8.

8 Hilberträume

Orthonormalität.

$$\langle \chi_n, \chi_m \rangle = \int_0^1 \exp(2\pi i k t) dt$$

für $k = m - n$. Für $k = 0$ ist das Integral 1, und für $k \neq 0$ gleich $(2\pi i k)^{-1} \exp(2\pi i k t)|_0^1 = 0$.

Dichtigkeit. Der von endlichen Linearkombinationen $g(t)$ der Funktionen $\chi_n(t)$ aufgespannte $\mathbb{C}$ -Untervektorraum definiert wegen $\chi_n(t) \cdot \chi_m(t) = \chi_{n+m}(t)$ eine $\mathbb{C}$ -Algebra A in $C(X, \mathbb{C})$. Offensichtlich trennt die Funktion $\chi_1(t) = \exp(2\pi i t)$ die Punkte von S^1 im Sinne von Abschnitt 8.7. Also gibt es für jede Funktion $f \in C(X, \mathbb{C})$ und jedes $\varepsilon > 0$ ein $g \in A$ mit

$$\|f - g\|_{L^2}^2 \leq \text{vol}(S^1) \cdot \sup_{t \in S^1} |f(t) - g(t)|^2 < \varepsilon$$

wegen des Satzes von Stone-Weierstraß (nächster Paragraph). Andererseits liegt $C(X, \mathbb{C})$ dicht in $L^2(X, \mathbb{C})$ nach Abschnitt 8.6, und damit liegt auch A dicht in $L^2(X, \mathbb{C})$.

Punktweise Konvergenz. Beachte $|\exp(2\pi i n t)| = 1$ für $t \in [0, 1]$. Wegen $f \in C^2([0, 1], \mathbb{C})$ und $f(0) = f(1)$ sowie $f'(0) = f'(1)$ folgt durch zweimalige partielle Integration

$$a(n) = \int_0^1 f(t) \exp(-2\pi i n t) dt = \frac{1}{-4\pi^2 n^2} \int_0^1 f''(t) \exp(-2\pi i n t) dt.$$

Also $|a(n)| \leq \frac{C}{n^2}$ für eine Konstante C , da die stetige Funktion $f''(t)$ auf dem Kompaktum $[0, 1]$ beschränkt ist. Wegen $\sum_{n=1}^{\infty} \frac{1}{n^2} \leq 1 + \int_1^{\infty} \frac{dt}{t^2} < +\infty$ konvergiert $g(x) := \sum_{n \in \mathbb{Z}} a_n \exp(2\pi i n t)$ absolut und gleichmäßig auf $[0, 1]$, definiert also nach Satz 2.24 eine stetige Funktion auf $[0, 1]$. Aus Schritt 1 (Übergang zu einer Teilreihe!) und dann Schritt 2 und 3 im Beweis des Satzes von Fischer-Riesz folgt, daß die Fourier Reihe punktweise (fü.) gegen $f(x)$ konvergiert. Daher sind $f(x)$ und $g(x)$ fast überall gleich. Da beide Funktionen stetig sind, folgt aus dem nächsten Lemma $f(x) = g(x)$. $\square$

Lemma 8.11. Ist h stetig auf $[0, 1]$ und gilt $\|h\|_{L^2}^2 = \int_0^1 |h(t)|^2 dt = 0$, dann ist $h = 0$.

Beweis. Wäre $h(t_0) \neq 0$, gäbe es wegen der Stetigkeit ein $\delta > 0$ mit $|h(t)| > \frac{1}{2}|h(t_0)|$ für $|t - t_0| < \delta$. Ein Widerspruch wegen $I(|h|^2) \geq \frac{1}{4}|h(t_0)|^2 \cdot I(\chi_{[-\delta+t_0, t_0+\delta]}) > 0$. $\square$

Warnung. Für $f(x) = x - \frac{1}{2}$ und $a_\nu = \int_0^1 (t - \frac{1}{2}) \exp(-2\pi i \nu t) dt$ gilt $a_0 = 0$, und mittels partieller Integration zeigt man $a_\nu = (t - \frac{1}{2}) \frac{\exp(-2\pi i \nu t)}{-2\pi i \nu} \Big|_0^1 = \frac{-1}{2\pi i \nu}$. Es folgt für $N \rightarrow \infty$

$$\sum_{\nu=1}^N \frac{\sin(2\pi \nu x)}{\pi \nu} \longrightarrow f(x) = x - \frac{1}{2}$$

im L^2 -Sinn wegen $f \in L^2([0, 1], \mathbb{C})$. Beachte $f(0) \neq f(1)$. Tatsächlich konvergiert die Reihe aber nicht punktweise. Im Punkt $x = 0$ ist der Fourier-Limes Null, aber es gilt $f(0) = -1/2$.

8.7 Stone-Weierstraß

Zur Erinnerung: Ein **Verband** B auf X ist ein $\mathbb{R}$ -Vektorraum von Funktionen $X \rightarrow \mathbb{R}$ mit der Eigenschaft $f(x) \in B \implies |f(x)| \in B$. Wegen $\max/\min(f, 0) = \frac{1}{2}(f \pm |f|)$ sowie $\max/\min(f, g) = g + \max/\min(f - g, 0)$ ist B unter der Bildung endlicher Minima und Maxima abgeschlossen.

Beispiel. Sei X ein metrischer Raum und $K \subseteq X$ eine Teilmenge. Ist B ein Verband auf X , dann bilden die Funktionen $B_K \subseteq B$ mit Träger in K wieder einen Verband.

Sei $K \subseteq X$, und B ein Verband auf X mit folgender Eigenschaft der **Punktetrennung** (*):

- Für je zwei Punkte $x \neq y$ in K und reelle Zahlen $a, b \in \mathbb{R}$ gibt es eine Funktion $f_{x,y} \in B_K$ mit der Eigenschaft $f_{x,y}(x) = a$ und $f_{x,y}(y) = b$, so daß $f_{x,y}$ stetig ist in einer Umgebung von x und y .

Ist B in $C(X)$ enthalten, ist natürlich $f_{x,y}$ automatisch stetig auf ganz X .

Satz 8.12. Sei K kompakt in X , B ein Verband auf X mit der Eigenschaft (*) und sei g eine stetige Funktion auf X mit Träger in K . Dann existiert für jedes $\varepsilon > 0$ eine beschränkte Funktion $f \in B_K$ mit $\text{Supremumsnorm}^4 \|g - f\| < \varepsilon$. Ist $g \geq 0$, kann $f \geq 0$ gewählt werden.

Zum Beweis dieses Satzes benutzen wir die folgende nichttriviale Aussage (Satz 11.4) von Heine-Borel für folgenkompakte metrische Räume X : Sei $X = \bigcup_{i \in I} U_i$ eine Überdeckung durch offene Teilmengen U_i von X , dann existiert eine endliche Teilmenge J der Indexmenge I mit der Eigenschaft $X = \bigcup_{i \in J} U_i$.

Beweis. Für $x \neq y \in K$ existiert ein $f_{x,y} \in B_K$ mit $f_{x,y}(x) = g(x)$ und $f_{x,y}(y) = g(y)$. Für festes x existiert zu jedem y eine offene Umgebung $V(y)$ mit $\sup_{y' \in V(y)} |f_{x,y}(y') - g(y')| < \varepsilon$, da g und $f_{x,y}$ bei y stetig sind. Endlich viele $V(y_1), \dots, V(y_m)$ der $V(y)$ überdecken K . Das Infimum $f_x := \inf(f_{x,y_1}, \dots, f_{x,y_m})$ ist in B_K (wegen der Verbandseigenschaft) und stetig in einer Umgebung von x mit

$$f_x(x) = g(x) \quad , \quad f_x(y) < g(y) + \varepsilon \quad (\forall y \in K) .$$

Für jedes $x \in K$ gibt es analog eine offene Umgebung $U(x)$ mit $\sup_{x' \in U(x)} |g(x') - f_x(x')| < \varepsilon$. Endlich viele $U(x_1), \dots, U(x_n)$ überdecken K . Für $f = \sup(f_{x_1}, \dots, f_{x_n}) \in B_K$ folgt

$$f(x) > g(x) - \varepsilon \quad , \quad f(y) < g(y) + \varepsilon \quad (\forall x, y \in K) ,$$

also $d_\infty(f, g) < \varepsilon$. Ersetze f durch $\max(f, 0)$ im Fall $g = \max(g, 0) \geq 0$. □

Satz 8.13 (Stone-Weierstraß). Sei X folgenkompakt und $A \subseteq C(X)$ eine $\mathbb{R}$ -Unteralgebra von $C(X)$ mit 1. Existiert für jedes Paar $x \neq y$ in X ein $f \in A$ mit $f(x) \neq f(y)$, dann liegt A dicht in $C(X)$. (Im Fall $A \subseteq C(X, \mathbb{C})$ fordert man noch $f \in A \implies i \cdot f \in A$ und $\bar{f} \in A$).

⁴siehe Abschnitt 2.7

Beweis. Der Abschluß $\overline{A}$ von A in $C(X)$ ist wieder eine Algebra und enthält A , ist also punktstetig. OBdA ist daher A abgeschlossen in $C(X)$. Wegen Satz 8.12 genügt es zu zeigen: A ist ein Verband, d.h. $f \in A \implies |f| \in A$. Also genügt es aus $0 \leq h = f^2 \in A$ eine positive Wurzel $\sqrt{h} \in A$ ziehen zu können. Im Fall $0 < c_1 \leq h \leq c_2 < 1$ ist $h = 1 - g$ mit $\|g\|_\infty < 1$, so daß nach Lemma 4.40 die Potenzreihe $\sqrt{h} = 1 - \frac{1}{2}g - \frac{1}{8}g^2 + \dots$ in $A \subseteq C(X)$ konvergiert. Ersetzt man also die beschränkte Funktion h durch $c \cdot (h + d)$ für kleine positive Konstanten c, d , kann man daraus in A die Wurzel ziehen. Da c eine Wurzel besitzt, folgt $\sqrt{h + d} \in A$. Im Limes $d \rightarrow 0$ folgt $\sqrt{h} \in A$. $\square$

8.8 Fourier Transformation

Eine **Schwartz-Funktion** $f : \mathbb{R}^N \rightarrow \mathbb{C}$ ist eine unendlich oft differenzbare Funktion auf $\mathbb{R}^N$, so daß für jede Ableitung $f^{(n)}(x)$ von $f(x)$ und jedes Polynom $P(x)$ eine von n und $P(x)$ abhängige Konstante $C = C(n, P(x))$ existiert mit

$$|P(x) \cdot f^{(n)}(x)| \leq C.$$

Sei $\mathcal{S}(\mathbb{R}^N)$ der Raum der Schwartz-Funktionen und $\mathcal{S} = \mathcal{S}(\mathbb{R})$.

Beispiel 8.14. Die **Gauß-Funktionen** $f(x) = \exp(-ax^2 - bx - c)$ für $b, c \in \mathbb{C}$ und reellem Exponenten $a > 0$ liegen in $\mathcal{S}$.

Der Raum der Schwartz-Funktionen $\mathcal{S}(\mathbb{R}^N)$ ist ein $\mathbb{C}$ -Untervektorraum von $L^2(\mathbb{R}^N)$. Für Schwartz-Funktionen $f \in \mathcal{S}(\mathbb{R}^N)$ und $y \in \mathbb{R}^N$ ist die **Fourier Transformierte** $\mathcal{F}f$ erklärt durch

$$(\mathcal{F}f)(y) = \int_{\mathbb{R}^N} f(x) \cdot \exp(2\pi i x \cdot y) dx.$$

Da $f(x) \cdot \exp(2\pi i x \cdot y)$ stetig in x und damit messbar ist, existiert das Integral nach Satz 6.19.6 auf Grund der Existenz einer auf $\mathbb{R}^N$ integrierbaren Majorante, denn $|f(x) \cdot \exp(2\pi i x y)| = |f(x)| \leq g(x) := \min(c_0, \frac{c_1}{\|x\|^{N+1}})$ mit $g \in L(\mathbb{R}^N)$. [Setze $c_0 = C(0, 1)$ und $c_1 = C(0, r^{2(N+1)})$.] Insbesondere ist daher $(\mathcal{F}f)(y)$ eine durch $const := \int_{\mathbb{R}^N} g(x) dx < +\infty$ beschränkte Funktion der Variable y . Es gilt $\mathcal{F}(\mathcal{S}(\mathbb{R}^N)) \subset \mathcal{S}(\mathbb{R}^N)$. Der Einfachheit halber sei nun $N = 1$.

Lemma 8.15. 1. $\mathcal{F}$ definiert eine $\mathbb{C}$ -lineare Abbildung $\mathcal{F} : \mathcal{S} \rightarrow \mathcal{S}$.

2. $\mathcal{F}$ bildet dabei $x^n \cdot f(x)$, für $f \in \mathcal{S}$, ab auf $(\frac{\partial y}{\partial x})^n (\mathcal{F}f)(y)$.

3. $\mathcal{F}$ bildet $(\frac{\partial x}{\partial y})^n f(x)$, für $f \in \mathcal{S}$, ab auf $(-y)^n \cdot (\mathcal{F}f)(y)$.

4. Es gilt $\mathcal{F}f = f$ für $f(x) = \exp(-\pi x^2)$.

Beweis. Wegen $\mathcal{F}(x^n f(x))(y) = \int_{\mathbb{R}} x^n f(x) \exp(2\pi i x y) dx = \int_{\mathbb{R}} (\frac{\partial y}{\partial x})^n f(x) \exp(2\pi i x y) dx$ folgt Aussage 2 aus Satz 4.32 [verifiziere die Voraussetzungen!]. Durch partielle Integration folgt Aussage 3. Daraus folgt $|P(-y)(\frac{\partial y}{\partial x})^n \mathcal{F}f(y)| = |\mathcal{F}(P(\frac{\partial x}{\partial y})^n f(x))(y)| \leq const$, und dies zeigt $\mathcal{F}(f) \in \mathcal{S}$ für $f \in \mathcal{S}$ und damit Aussage 1.

Aussage 4 ist am schwierigsten: Mit Hilfe des Hauptsatzes und Satz 4.32 zeigt man zuerst, daß $c(y) = \exp(\pi y^2) \cdot \mathcal{F}f(y) = \int_{\mathbb{R}} \exp(-\pi(x-iy)^2) dx$ konstant als Funktion von y ist, denn

$$\begin{aligned} \frac{d}{dy} c(y) &= \frac{d}{dy} \int_{\mathbb{R}} e^{-\pi(x-iy)^2} dx = \int_{\mathbb{R}} \frac{d}{dy} e^{-\pi(x-iy)^2} dx = 2\pi i \int_{\mathbb{R}} (x-iy) e^{-\pi(x-iy)^2} dx \\ &= -i \int_{\mathbb{R}} \frac{d}{dx} e^{-\pi(x-iy)^2} dx = \lim_{n \rightarrow \infty} -i e^{-\pi(x-iy)^2} \Big|_{-n}^{+n} = 0. \end{aligned}$$

Aus der Substitutionsregel und dem Satz von Fubini folgt dann für $c = c(0)$

$$c^2 = \int_{\mathbb{R}^2} e^{-\pi(x_1^2+x_2^2)} dx_1 dx_2 = \lim_{N \rightarrow \infty} \lim_{\varepsilon \rightarrow 0} \int_{[\varepsilon, N]} \int_0^{2\pi} e^{-\pi r^2} r dr d\theta = \frac{-e^{-x}}{2\pi} \Big|_0^\infty \cdot 2\pi = 1.$$

Wegen $c \geq 0$ gilt daher $c = 1$. Es folgt $(\mathcal{F}f)(y) = \exp(-\pi y^2)$ sowie

$$\boxed{\int_{\mathbb{R}} e^{-\pi x^2} dx = 1}.$$

□

Die **Gauß-Funktionen** $\exp(-ax^2 - bx - c)$ für $b, c \in \mathbb{C}$ mit reellen **Exponenten** $a > \pi$ spannen einen $\mathbb{C}$ -Untervektorraum $\mathcal{G}$ von $\mathcal{S}$ auf.

Lemma 8.16. $\mathcal{G}$ (und damit $\mathcal{S}$) liegt dicht in $L^2(\mathbb{R})$.

Beweis. Wegen Satz 8.13, Lemma 5.32 sowie Lemma 8.6 liegt $C_c^2(\mathbb{R}, \mathbb{C})$ dicht⁵ in $L^2(\mathbb{R})$. Es genügt daher Funktionen $\tilde{f}(x) \in C_c^2(\mathbb{R}, \mathbb{C})$ durch Funktionen $\tilde{g}_m \in \mathcal{G}$ in der L^2 -Norm zu approximieren. Fixiere $\tilde{f}(x) \in C_c^2(\mathbb{R}, \mathbb{C})$ und $\varepsilon > 0$; fixiere dann ein $t_0 = t_0(\tilde{f}) \geq 1$ und ein x_0 mit $|x_0| < \frac{1}{2}$ derart daß der Träger von $f(x) = \tilde{f}(t_0 x)$ in $[-x_0, x_0] \subset [-\frac{1}{2}, \frac{1}{2}]$ liegt. Wir benutzen dann: Approximationen von $\tilde{f}(x)$ durch Funktionen $\tilde{g}_m(x)$ in $\mathcal{G}$ (d.h. mit **Exponenten** $> \pi$) entsprechen Approximationen von $f(x)$ durch Summen $g_m(x) = \tilde{g}_m(t_0 x)$ von Gaußfunktionen mit **Exponenten** $> t_0^2 \cdot \pi$ vermöge $t_0^{1/2} \|f - g\|_{L^2(\mathbb{R})} = \|\tilde{f} - \tilde{g}\|_{L^2(\mathbb{R})}$. Für

$$h(x) := e^{tx^2} \cdot f(x) \quad , \quad t \geq t_0^2 \cdot \pi$$

gilt $\text{supp}(h) \subseteq [-x_0, x_0]$ und $h(x) \in C_c^2(\mathbb{R}, \mathbb{C})$. Wegen $h^{(\nu)}(-1/2) = h^{(\nu)}(1/2) = 0$ für $\nu = 0, 1$ konvergiert daher auf $[-\frac{1}{2}, \frac{1}{2}]$ die Fourierentwicklung $\sum_{n \in \mathbb{Z}} a(n) e^{2\pi i n x} \rightarrow h(x)$ gleichmäßig (Satz 8.10). Wegen $\exp(-tx^2) \leq 1$ konvergiert auf $[-\frac{1}{2}, \frac{1}{2}]$ dann ebenfalls

$$g_m(x) = e^{-tx^2} \cdot p_m(x) = \sum_{|n| \leq m} a(n) e^{-tx^2 + 2\pi i n x} \rightarrow f(x) = e^{-tx^2} \cdot h(x)$$

gleichmäßig gegen $f(x)$. Es folgt für $m \geq m_0(\varepsilon, t_0, t)$

$$t_0 \int_{|x| \leq \frac{1}{2}} |f(x) - g_m(x)|^2 dx < \varepsilon^2/2.$$

⁵Allgemeinen liegen Produkte $\prod_{i=1}^N f_i(x_i)$ von Gaußfunktionen $f_i(x_i)$ dicht in $L^2(\mathbb{R}^N)$. Man reduziert diese Aussage mit Lemma 8.6 leicht auf den hier behandelten Fall $N = 1$, da nach Satz 8.12, Satz 8.13 die Algebra erzeugt von allen Produkte $\prod_{i=1}^N f_i(x_i)$ mit $f_i \in C_c^2(\mathbb{R}, \mathbb{C})$ dicht in $C_c(\mathbb{R}^N, \mathbb{C})$ liegt.

8 Hilberträume

Im Bereich $\frac{1}{2} \leq |x|$ ist $f(x) = 0$. Also lässt sich $\int_{\frac{1}{2} \leq |x|} |f - g_m|^2 dx = \int_{\frac{1}{2} \leq |x|} |g_m(x)|^2 dx$ durch $2 \sum_{\nu=1}^{\infty} \exp(-2t(\frac{2\nu+1}{2})^2) \cdot \int_{-1/2}^{1/2} |p_m(x)|^2 dx$ abschätzen, da $|p_m(x)|^2$ periodisch ist und e^{-2tx^2} monoton fallend ist. Wegen $\int_{-1/2}^{1/2} |p_m(x)|^2 dx = \sum_{|n| \leq m} |a_n|^2 \leq \|h\|_{L^2([-1/2, 1/2])}^2 \leq C_1 e^{2tx_0^2}$ (Plancherelformel) und $\sum_{\nu=1}^{\infty} \exp(-t(\frac{2\nu+1}{2})^2) \leq \exp(-\frac{t}{2}) / (1 - \exp(-\frac{t}{2})) \leq C_2 \exp(-\frac{t}{2})$ für genügend große t , folgt daher aus $x_0^2 < 1/4$ für $t \geq \max(t_0(\varepsilon), t_0^2 \cdot \pi)$

$$t_0 \int_{\frac{1}{2} \leq |x|} |f(x) - g_m(x)|^2 dx \leq t_0 C_1 C_2 \cdot e^{2t(x_0^2 - \frac{1}{4})} < \varepsilon^2 / 2.$$

Zusammen mit der Abschätzung für $|x| \leq \frac{1}{2}$ folgt für genügend großes $t = t(\varepsilon)$

$$t_0^{1/2} \|f - g_m\|_{L^2(\mathbb{R})} < \varepsilon \quad \text{für} \quad m \geq m_0(\varepsilon, t_0, t).$$

Die durch $g_m(x) = \tilde{g}_m(t_0 x)$ definierten Funktionen $\tilde{g}_m$ liegen dann nach Konstruktion in $\mathcal{G}$ und approximieren $\tilde{f}(x)$ in der L^2 -Metrik: $\|\tilde{f} - \tilde{g}_m\|_{L^2(\mathbb{R})} = t_0^{1/2} \|f - g_m\|_{L^2(\mathbb{R})} < \varepsilon$. $\square$

Lemma 8.17. *Der von den $\mathbb{C}$ -linear unabhängigen Funktionen $x^n \cdot \exp(-\pi x^2)$ für natürliche $n \in \mathbb{N}$ aufgespannte $\mathbb{C}$ -Vektorraum $\mathcal{H}$ liegt dicht in $L^2(\mathbb{R})$.*

Beweis. Punktweise Konvergenz $f_n(x) \rightarrow f(x)$ bzw. $f(x) - f_n(x) \rightarrow 0$ auf $\mathbb{R}$ für f, f_n in $L^2(\mathbb{R})$ zusammen mit $|f - f_n|^2 \leq F$ für ein $F \in L(\mathbb{R})$ impliziert nach Satz 6.12

$$\lim_{n \rightarrow \infty} \|f - f_n\|_{L^2(\mathbb{R})}^2 = \lim_{n \rightarrow \infty} \int_{\mathbb{R}} |f(x) - f_n(x)|^2 dx \rightarrow 0.$$

Schritt 1. $f_n(x) = Q_n(x) \cdot \exp(-\alpha x^2) := \exp(-\alpha x^2) P(x) \sum_{m=0}^{m=n} \frac{(-1)^m}{m!} (\rho x^2 + bx + c)^m$ konvergiert wegen Satz 4.42 punktweise auf $\mathbb{R}$ gegen $f(x) := P(x) \exp(-ax^2 - bx - c)$ für

$$a = \alpha + \rho.$$

Es gilt $|f - f_n|^2 \leq |P(x)|^2 \exp(-2\alpha x^2) \exp(2(\rho x^2 + |b||x| + |c|)) = F(x)$. Wegen $a > \pi$ ist die Majorante $F(x) := |P(x)|^2 \exp(-\lambda x^2 + 2|b||x| + 2|c|)$ mit $\lambda = 2(a - 2\rho)$ in $L(\mathbb{R})$ für alle $0 \leq \rho \leq \frac{\pi}{2}$. Die Funktionen $f_n(x) = Q_n(x) \cdot \exp(-\alpha x^2)$ approximieren in diesem Fall nach Satz 6.12 im L^2 -Sinn die Funktion $f(x) = P(x) \cdot \exp(-ax^2 - bx - c)$, d.h.

$$\|f - f_n\|_{L^2}^2 < \varepsilon^2 \quad \text{für} \quad n \geq n_0(\varepsilon).$$

Schritt 2. Nach Lemma 8.16 liegt der Aufspann der Funktionen $f(x) = P(x) \cdot e^{-ax^2 - bx - c}$ mit $b, c \in \mathbb{C}$ und reellem Exponent $a > \pi$ und Polynomen $P(x)$ dicht in $L^2(\mathbb{R})$. Wendet man Schritt 1 sukzessive auf die approximierenden Funktionen $f(x) = P(x) \cdot e^{-ax^2 - bx - c}$ an, kann man den Exponent $a > \pi$ um ρ auf $\alpha = a - \rho$ verkleinern; wähle $0 \leq \rho \leq \min(\frac{\pi}{2}, a - \pi)$. Nach endlich vielen Schritten folgt, daß alle Exponenten gleich $\alpha = \pi$ gewählt werden können. Mit anderen Worten: Für jedes $\varepsilon > 0$ existiert ein komplexes Polynom $Q(x)$ mit

$$\|P(x) \cdot e^{-ax^2 - bx - c} - Q(x) \cdot e^{-\pi x^2}\|_{L^2(\mathbb{R})} < \varepsilon.$$

Das heisst, die Funktionen $Q(x) \cdot \exp(-\pi x^2)$ für Polynome $Q(x)$ liegen dicht in $L^2(X)$.

Schritt 3. Die Funktionen $x^n e^{-\pi x^2}$ sind $\mathbb{C}$ -linear unabhängig, da Monome $\mathbb{C}$ -linear unabhängig sind (Lemma 4.39), und bilden eine Basis des von den $Q(x) \cdot \exp(-\pi x^2)$ für Polynome $Q(X)$ aufgespannten Vektorraums. $\square$

Fourier Transformation erhält die Teilräume $V_N = \bigoplus_{n=0}^N \mathbb{C} x^n \cdot e^{-\pi x^2}$ wegen Lemma 8.15. Mittels Induktion nach n zeigt man durch orthogonale Projektion (**Gram-Schmidt**): Es gibt Polynome $H_n(x)$ vom Grad n mit der Eigenschaft: 1) Die Funktionen $f_n(x) = H_n(x) e^{-\pi x^2}$ bilden eine ON-Basis des Hilbertraums. 2) Die $f_n(x)$ für $n = 0, \dots, N$ definieren eine Basis von

$$V_N = \bigoplus_{n=0}^N \mathbb{C} \cdot x^n e^{-\pi x^2}.$$

Die dadurch eindeutig bestimmten Polynome $H_n(x)$ mit positivem reellem höchstem Koeffizienten sind die sogenannten **Hermiteischen Polynome**. In der Tat gilt bis auf geeignete Normierungskonstanten $c(n) \in \mathbb{R}$

$$H_n(x) := c(n) \cdot e^{2\pi x^2} \partial_x^n (e^{-2\pi x^2}),$$

denn $\int_{\mathbb{R}} \overline{P(x)} \exp(-\pi x^2) \cdot H_n(x) \exp(-\pi x^2) dx = 0$ und $\int_{\mathbb{R}} \overline{P(x)} c(n) (\partial_x^n \exp(-2\pi x^2)) dx = 0$ für alle Polynome $P(x)$ vom Grad $\leq n-1$ (für letzteres benutze partielle Integration!). Also liegen $e^{\pi x^2} \partial_x^n (e^{-2\pi x^2})$ und $H_n(x) e^{-\pi x^2}$ in V_n und sind orthogonal zu dem Unterraum V_{n-1} von V_n der Kodimension 1, und damit proportional.

Aus der obigen expliziten Formel für $H_n(x)$ folgt $H_n(-x) = (-1)^n H_n(x)$.

Satz 8.18. Die Funktionen $f_n(x) = H_n(x) \exp(-\pi x^2) \in \mathcal{S}$ sind Eigenfunktionen der Fourier Transformation $\mathcal{F} : \mathcal{S} \rightarrow \mathcal{S}$ zu den Eigenwerten i^n und definieren eine ON-Basis des Hilbertraumes $L^2(\mathbb{R})$. Die Fourier Transformation lässt sich daher eindeutig fortsetzen zu einer **unitären $\mathbb{C}$ -linearen Transformation des Hilbertraumes $L^2(\mathbb{R})$**

$$\mathcal{F} : L^2(\mathbb{R}) \cong L^2(\mathbb{R}) \quad , \quad (\mathcal{F}f, \mathcal{F}g) = (f, g).$$

Beweis. $(\mathcal{F}f_n)(y) = c(n) \int_{\mathbb{R}} e^{\pi x^2 + 2\pi ixy} \partial_x^n e^{-2\pi x^2} dx = c(n) e^{\pi y^2} \int_{\mathbb{R}} e^{\pi(x+iy)^2} \partial_x^n e^{-2\pi x^2} dx$ ist wegen partieller Integration dasselbe wie $c(n) e^{\pi y^2} (-1)^n i^{-n} (\partial_y)^n \int_{\mathbb{R}} e^{-2\pi x^2} e^{\pi(x+iy)^2} dx$, und wie $i^n c(n) e^{\pi y^2} \partial_y^n e^{-\pi y^2} (\mathcal{F}f_0)(y) = i^n f_n(y)$. Daher ist $\mathcal{F} : \mathcal{H} \rightarrow \mathcal{H}$ eine $\mathbb{C}$ -lineare Isometrie. Diese kann man auf ganz $L^2(\mathbb{R})$ fortsetzen: Für $v \in L^2(\mathbb{R})$ existiert eine Folge v_n aus $\mathcal{H}$, welche gegen v konvergiert. Da v_n eine Cauchyfolge in $\mathcal{H} \subseteq \mathcal{S}$ ist, ist $\mathcal{F}(v_n)$ eine Cauchyfolge in $\mathcal{H}$, denn $\mathcal{F}$ ist auf $\mathcal{H} \subseteq \mathcal{S}$ definiert und eine Isometrie! Setze $\mathcal{F}(v) := \lim_n \mathcal{F}(v_n)$. Man zeigt nun leicht, dass $\mathcal{F} : L^2(\mathbb{R}) \rightarrow L^2(\mathbb{R})$ wohldefiniert, $\mathbb{C}$ -linear und eine Isometrie ist. $\square$

Korollar 8.19 (Fourier Inversion). Für alle f in $L^2(\mathbb{R})$ gilt

$$(\mathcal{F}\mathcal{F}f)(x) = f(-x).$$

Wir bemerken: $(\mathcal{F}\mathcal{F}f)(x) = f(-x)$ gilt für alle Basisfunktionen $f(x) = H_n(x)e^{-\pi x^2}$ von $\mathcal{H}$ und somit dann auch für beliebiges $f \in L^2(\mathbb{R})$ (zumindestens fast überall auf $X = \mathbb{R}$). Ist f und damit auch $(\mathcal{F}\mathcal{F}f)(x)$ in $\mathcal{S}$, stimmen daher $f(-x)$ und $(\mathcal{F}\mathcal{F}f)(x)$ als stetige Funktionen in allen Punkten $x \in \mathbb{R}$ überein. Es folgt

Korollar 8.20 (Fourier Inversion). Für alle $f \in \mathcal{S}(\mathbb{R})$ gilt

$$\boxed{f(x) = (\mathcal{F}g)(-x) = \int_{\mathbb{R}} g(y)e^{-2\pi i y x} dy} \quad \text{für} \quad \boxed{g(y) = (\mathcal{F}f)(y) = \int_{\mathbb{R}} f(x)e^{2\pi i y x} dx}.$$

Die Aussagen und Argumente übertragen sich verbatim auf den N -dimensionalen Fall.

8.9 Der harmonische Oszillator

Die von uns konstruierte ON-Basis des Hilbertraumes $L^2(\mathbb{R})$ ist definiert durch

$$f_n(x) = H_n(x)e^{-\pi x^2} = c(n) \cdot e^{\pi x^2} (e^{\pi x^2} \partial_x e^{-\pi x^2})^n (e^{-\pi x^2}) = c(n) \cdot A_+^n (e^{-\pi x^2}).$$

Hierbei ist A_+ der Differentialoperator $A_+ := e^{\pi x^2} \partial_x e^{-\pi x^2} = \partial_x - 2\pi x = X + iY$ mit den Abkürzungen $X = \partial_x$ und $Y = 2\pi i x$. Ist analog $A_- := X - iY$, wird $f_0(x) = e^{-\pi x^2}$ von A_- annulliert, wobei die Gleichung $A_-(f(x)) = \partial_x f(x) - 2\pi x f(x) = 0$ die Funktion $f_0(x)$ bis auf eine Normierungskonstante c eindeutig bestimmt (Satz 4.18).

$[A_+, A_-] = (X + iY)(X - iY) - (X - iY)(X + iY) = -2[X, iY] = 4\pi[\partial_x, x] = 4\pi$ zeigt, daß $A_-(A_+^n) - (A_+)^n A_- = [A_-, A_+]A_+^{n-1} + A_+[A_-, A_+]A_+^{n-2} + \dots + A_+^{n-1}[A_-, A_+]$ gleich $-4\pi n \cdot A_+^{n-1}$ ist, und $A_+ A_- f_n(x) = c(n) A_+ A_- A_+^n(x) f_0(x) = -4\pi n c(n) A_+^n f_0(x) = -4\pi n \cdot f_n(x)$. Also sind alle $f_n(x)$ Eigenfunktionen des Operators $A_+ A_-$.

$$(A_+ A_-) f_n(x) = -4\pi n \cdot f_n(x).$$

Lemma 8.21. Für $n = 0, 1, 2, \dots$ gilt $\boxed{c(n)^2 = \frac{c(0)^2}{(4\pi)^n n!}}$ und oBdA $c(0) = 1$.

Beweis. Offensichtlich gilt $A_+ f_n(x) = c_n \cdot f_{n+1}(x)$ für gewisse andere Konstanten c_n mit $c_n^2 = \|c_n f_{n+1}\|^2 = \|A_+ f_n\|^2$. Partielle Integration zeigt $\|A_+ f_n\|^2 = -\langle f_n(x), A_- A_+ f_n(x) \rangle$. Benutzt man $A_- A_+ = -4\pi + A_+ A_-$, folgt $c_n^2 = -\langle f_n(x), -4\pi(n+1)f_n(x) \rangle = 4\pi(n+1)$. Wegen $c_n \cdot c(n) = c(n+1)$ ergibt sich dann die Behauptung durch Induktion nach n . $\square$

Der Operator $\frac{1}{2}(A_+ A_- + A_- A_+) = A_+ A_- - 2\pi$ entspricht bis auf den Faktor $(2\pi i)^2$ dem Operator $(\partial_x - 2\pi x)(\partial_x + 2\pi x) - 2\pi = \partial_x^2 + (2\pi i x)^2 = X^2 + Y^2$. Es folgt⁶

$$\boxed{\frac{1}{2}(X^2 + Y^2)f_n(x) = -2\pi(n + \frac{1}{2}) \cdot f_n(x)}.$$

⁶Wir bemerken, der Operator $\frac{1}{2}(X^2 + Y^2)$ ist bis auf Normierungen der **Hamilton Operator** des quantenmechanischen harmonischen Oszillators. Zur Erinnerung $\frac{X}{2\pi i} = h \cdot \frac{1}{2\pi i} \partial_x$ entspricht dem Impulsoperator p und $\frac{Y}{2\pi i} = x$ dem Ortsoperator q (bei uns ist $h = 1$).

9 Integration auf Mannigfaltigkeiten

9.1 Untermannigfaltigkeiten mit Rand

Sei $f : \mathbb{R}^n \rightarrow \mathbb{R}$ eine C^∞ -Funktion mit den Eigenschaften

- $f(\xi) = 0 \implies df(\xi) \neq 0$
- $M = \{x \in \mathbb{R}^n \mid f(x) \leq 0\}$ sei kompakt.

Man nennt dann $\partial M = \{x \in \mathbb{R}^n \mid f(x) = 0\}$ den **Rand** von M , und M eine **komakte Untermannigfaltigkeit** von $\mathbb{R}^n$ mit Rand ∂M .

Beispiel. Für $f(x) = -1 + \sum_{i=1}^n x_i^2$ ist M die abgeschlossene Einheitskugel E im $\mathbb{R}^n$ und ∂M ist die Einheitssphäre S der Dimension $n - 1$ im $\mathbb{R}^n$.

Lokale Beschreibung des Randes. Für jeden Randpunkt $\xi \in \partial M$ gilt $\partial_\nu f(\xi) \neq 0$ für ein $\nu = 1, \dots, n$ nach Annahme. Durch Umbenennen sei obdA $\nu = 1$ (für gegebenes ξ) und damit $\partial_1 f(\xi) \neq 0$. Aus Stetigkeitsgründen ist dann die Determinante der Jacobi Matrix

$$\begin{pmatrix} \partial_1 f(x) & 0 \\ * & E \end{pmatrix}$$

der folgenden Abbildung [Beachte: Die Jacobimatrix ist eine Dreiecksmatrix und hat daher die Determinante $\partial_1 f(x)$] von Null verschieden für alle $x \in K_{2r}(\xi)$ nahe bei ξ

$$(x_1, \dots, x_n) \mapsto (f(x), x_2, \dots, x_n)$$

bei geeigneter Wahl von $r = r(\xi) > 0$. Für spätere Zwecke bezeichne

$$\varepsilon_\xi = \text{sign}(\partial_1 f(\xi))$$

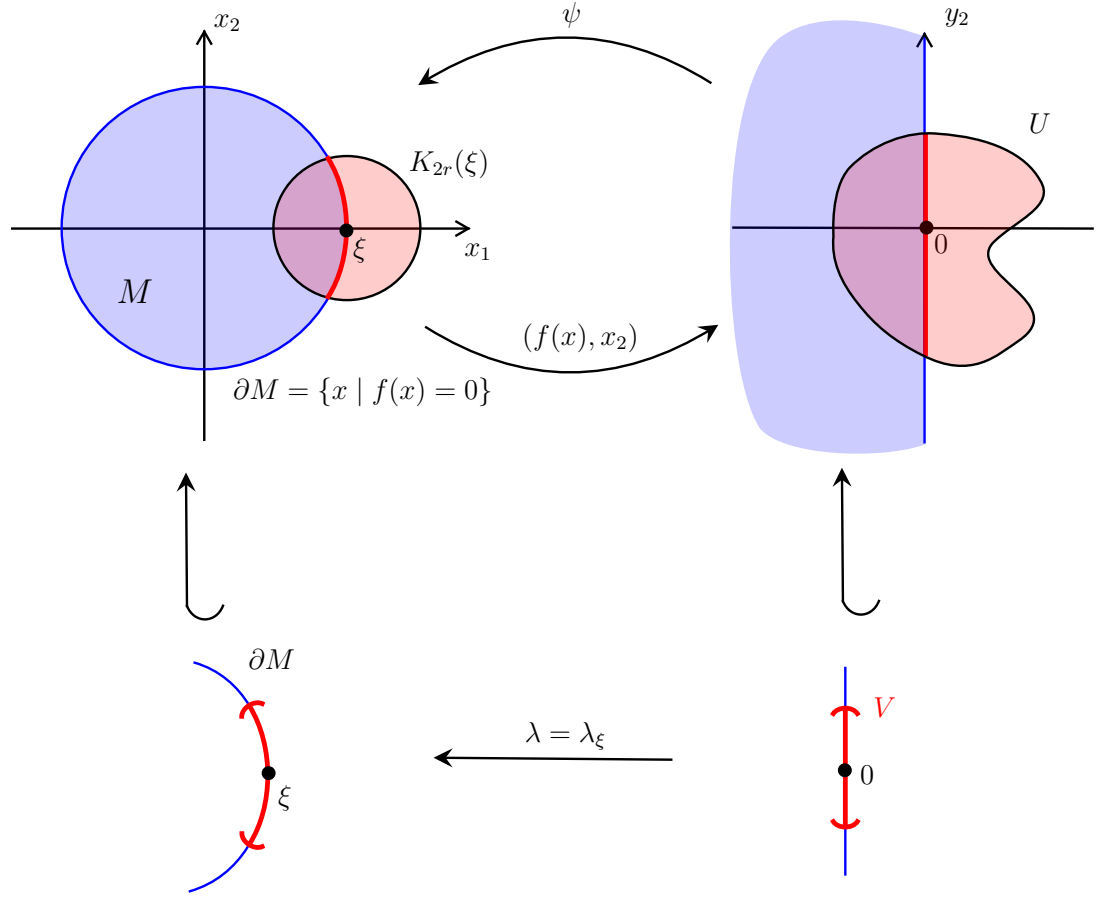
das **Vorzeichen** dieser Determinante im Punkt ξ . Aus dem Satz von der Umkehrfunktion 4.25 folgt dann die Existenz einer lokalen C^∞ -Umkehrfunktion ψ der obigen Funktion. Diese lokale Umkehrfunktion hat dann notwendiger Weise die Gestalt

$$\psi(y) = \psi(y_1, \dots, y_n) = (g(y), y_2, \dots, y_n).$$

und ψ stiftet eine bijektive Abbildung zwischen den offenen Mengen

$$U := \psi^{-1}(K_{2r}(\xi)) \cong K_{2r}(\xi);$$

(r hängt von ψ etc. ab).



Somit definiert ψ Bijektionen

$$U \cap (\mathbb{R}_{\leq 0} \times \mathbb{R}^{n-1}) \cong K_{2r}(\xi) \cap M,$$

$$V := U \cap (\{0\} \times \mathbb{R}^{n-1}) \cong K_{2r}(\xi) \cap \partial M.$$

Die Gleichung $y_1 = 0$ beschreibt daher in $U = \psi^{-1}(K_{2r}(\xi))$ den Rand ∂M von M . Wegen $x_2 = y_2, \dots, x_n = y_n$ (für die Abbildung ψ und ihr Inverses) ist daher

$$\lambda: \mathbb{R}^{n-1} \supset V \ni (y_2, \dots, y_n) \mapsto (g(0, y_2, \dots, y_n), y_2, \dots, y_n)$$

eine **lokale C^∞ -Parametrisierung** des Randes ∂M von M in der Nähe des Punktes $\xi \in \partial M$.

Für Punkte $\xi \in M \setminus \partial M$ findet man ein $r = r(\xi) > 0$ derart, daß gilt $K_{2r}(\xi) \subseteq M \setminus \partial M$. In der Tat ist $M \setminus \partial M$ eine offene Teilmenge des $\mathbb{R}^n$, denn das Komplement $\{x \in \mathbb{R}^n \mid f(x) \geq 0\}$ ist eine abgeschlossene Teilmenge des $\mathbb{R}^n$. Da M kompakt ist, überdecken bereits endliche viele der Kugeln $K_r(\xi)$, $\xi \in M$ die Menge M (Satz 11.4). Sei I die endliche Menge der zugehörigen Mittelpunkte $\xi \in M$.

9.2 Randintegrale

Annahmen. Sei $f : \mathbb{R}^n \rightarrow \mathbb{R}$ eine C^∞ -Funktion mit der Eigenschaft

$$f(\xi) = 0 \implies df(\xi) \neq 0,$$

und $\partial M := \{x \in \mathbb{R}^n \mid f(x) = 0\}$ sei kompakt. Sei U offen in $\mathbb{R}^n$ mit $\partial M \subseteq U$.

Wir definieren nun für eine Differentialform $\eta \in A^{n-1}(U)$ das Randintegral $\int_{\partial M} \eta$.

Warnung. Dies ist kein Integral im $\mathbb{R}^n$, denn dort ist ∂M eine Menge vom Maß Null! Wir machen hierfür folgende

Für endliche viele $\xi \in I$ auf dem Rand ∂M überdecken die Kugeln $K_r(\xi)$, $r = r(\xi)$ die Randmenge ∂M . ObdA gilt $K_r(\xi) \subset U$. Für jedes $\xi \in I$ haben wir eine **lokale Parametrisierung** λ_ξ durch eine offene Teilmenge $V = V_\xi$ des $\mathbb{R}^{n-1}$ gefunden

$$\lambda_\xi : \mathbb{R}^{n-1} \supseteq V_\xi \longrightarrow \partial M \cap K_{2r}(\xi) \subseteq K_{2r}(\xi) \subseteq \mathbb{R}^n.$$

Die Abbildung $\lambda = \lambda_\xi$ ist vom Typ C^∞ . Sei $\{\varphi_\xi \mid \xi \in I\}$ eine **zugeordnete Partition der Eins** auf der in Abschnitt 5.17 definierten offenen Teilmenge $\tilde{N} = \bigcup_{\xi \in I} K_r(\xi) \subseteq U$ von $\mathbb{R}^n$.

Definition 9.1. Sei $\eta \in A^{n-1}(\tilde{N})$ und damit $\varphi_\xi(x) \cdot \eta \in A_c^{n-1}(V_\xi)$. Dann setzen wir

$$\boxed{\int_{\partial M} \eta := \sum_{\xi \in I} \varepsilon_\xi \cdot \int_{V_\xi} \lambda_\xi^*(\varphi_\xi(x) \cdot \eta)} \quad \varepsilon_\xi \in \{\pm 1\}.$$

Obwohl dies auf den ersten Blick hochgradig von der Wahl der Stützpunkte $\xi \in I$, der Wahl der lokalen Parametrisierungen (V_ξ, λ_ξ) von ∂M und der Wahl einer Partition der Eins φ_ξ abzuhängen scheint, gilt erstaunlicherweise

Lemma 9.2. Das in Definition 9.1 definierte Randintegral ist unabhängig von allen hierbei getroffenen Wahlen $\lambda_\xi, \varphi_\xi, V_\xi, I$.

Beweis. Zur Unabhängigkeit von der Wahl der Partition der Eins, der Stützpunkte $\xi \in I$ und der Radien $r = r(\xi)$: Für eine andere Wahl bekommt man eine neue Partition der Eins $\varphi'_{\xi'}$ (für endlich viele neue Stützpunkte $\xi' \in I'$). Man betrachtet zum Vergleich die neue Partition der Eins $\varphi_\xi(x) \cdot \varphi_{\xi'}(x)$ für $(\xi, \xi') \in J = I \times I'$ mit Trägern in $K_{2r}(\xi) \cap K_{2r'}(\xi')$ (Übungsaufgabe!). Zur Unabhängigkeit von der Parametrisierung: Hierzu beachte, daß nach unserer Konstruktion für zwei verschiedene Parametrisierungen die Zusammensetzung $h = \lambda^{-1} \circ \lambda'$

$$V \xrightarrow{\lambda} \partial M \cap K_{2r}(\xi) \xleftarrow{\lambda'} V'$$

sowie die Umkehrung h^{-1} beide C^∞ -Abbildungen sind; insbesondere daher $\det(D(h)(x)) \neq 0$. Die Behauptung folgt dann aus der Substitutionsregel 4.27, denn den Übergang vom Integral über V

$$\int_V \varepsilon_\xi \cdot \lambda_\xi^*(\eta)$$

zum Integral über V'

$$\int_{V'} \varepsilon_{\xi'} \cdot \lambda_{\xi'}^*(\eta)$$

erhält man durch eine Substitution mittels der Abbildung h , wenn man

$$h^*(\lambda_{\xi'}^*(\eta)) = \lambda_{\xi}^*(\eta)$$

berücksichtigt sowie $\text{sign}(\det(Dh(x))) \cdot \varepsilon_{\xi'} = \varepsilon_{\xi}$. Letztere Aussage über Vorzeichen folgt aus einer der Übungsaufgaben auf den Übungsblättern und wird benötigt aus folgendem Grund: In der Substitutionsregel für mehrdimensionale Integrale tritt der Absolutbetrag der Determinante der Jacobimatrix $D(h(x))$ auf, beim Pullback dagegen nur die Determinante der Jacobimatrix $D(h(x))$ selbst. Siehe dazu Seite 70.

□

9.3 Der Satz von Stokes

Satz 9.3. Sei M eine kompakte Untermannigfaltigkeit des $\mathbb{R}^n$ mit Rand ∂M und $\eta \in A^{n-1}(\tilde{N})$ eine Differentialform auf einer offenen Menge $\tilde{N} \subseteq \mathbb{R}^n$, welche M enthält. Dann gilt

$$\boxed{\int_M d\eta = \int_{\partial M} \eta}.$$

Beweis. Sei $\{\varphi_{\xi}(x) \mid \xi \in I\}$ eine Partition der Eins zu einer Überdeckung von M durch offene Kugeln $K_{2r}(\xi)$. Dann gilt für das Integral $\int_M d\eta := \int_{\mathbb{R}^n} \chi_M(x) \cdot d\eta$ wegen der $\mathbb{R}$ -Linearität sowohl des Lebesgue Integrals auf $\mathbb{R}^n$ als auch der Cartan Ableitung d

$$\int_M d\eta = \int_M d\left(\sum_{\xi \in I} \varphi_{\xi}(x) \cdot \eta\right) = \sum_{\xi \in I} \int_M d(\varphi_{\xi}(x) \cdot \eta).$$

Es genügt daher $\int_M d(\varphi_{\xi}(x) \cdot \eta) = \varepsilon_{\xi} \cdot \int_{V_{\xi}} \lambda_{\xi}^*(\varphi_{\xi}(x) \cdot \eta)$ für jeden Summanden zu zeigen; Summation über die $\xi \in I$ liefert dann den Satz von Stokes. Jetzt ist $\rho = \varphi_{\xi}(x) \cdot \eta \in A_c^{n-1}(K_{2r}(\xi))$ eine Differentialform mit kompaktem Träger in einer lokalen Kartenmenge $K_{2r}(\xi)$ von M . Zu zeigen bleibt dafür (*)

$$\int_M d\rho = \varepsilon_{\xi} \cdot \int_{V_{\xi}} \lambda_{\xi}^*(\rho).$$

Die Substitutionsregel für die Substitution $\psi = \psi_{\xi}$ (siehe Abschnitt 9.1) sowie Trägergründe implizieren für die linke Seite von (*) und die offene Menge $U_{\xi} := \psi^{-1}(K_{2r}(\xi)) \subseteq \mathbb{R}^n$

$$\int_M d\rho := \int_{K_{2r}(\xi)} \chi_M(x) \cdot d\rho = \varepsilon_{\xi} \cdot \int_{U_{\xi}} \psi^*(\chi_M(x) \cdot d\rho) = \varepsilon_{\xi} \cdot \int_{U_{\xi}} \chi_{y_1 \leq 0} \cdot d(\psi^*(\rho)),$$

denn es gilt $\psi^*(\chi_M(x)) = \chi_{y_1 \leq 0}(y)$. Setze $\omega := \psi^*(\rho) \in A_c^{n-1}(U_{\xi})$.

Zur Berechnung der rechten Seite (*) beachte $i^*(\omega) = i^*(\psi^*(\rho)) = \lambda_\xi^*(\rho)$ für die Inklusion $i : V_\xi \hookrightarrow U_\xi$ und den Pullback von ω auf $V_\xi = U_\xi \cap \{y_1 = 0\}$ wegen $\lambda_\xi = \psi \circ i$. Es verbleibt daher für $\omega = \psi^*(\rho) \in A_c^{n-1}(U_\xi)$ und die Inklusion $i : V_\xi \hookrightarrow U_\xi$ zu zeigen (**)

$$\int_{U_\xi} \chi_{y_1 \leq 0} \cdot d\omega = \int_{V_\xi} i^*(\omega) \quad \text{für} \quad V_\xi = U_\xi \cap \{y_1 = 0\},$$

Legt man $U_\xi \cap (\mathbb{R}_{\leq 0} \times \mathbb{R}^{n-1})$ in einen genügend grossen Quader $Q \subseteq \mathbb{R}^n$, dessen eine Wand durch die Hyperebene $y_1 = 0$ gegeben ist, und setzt die Form ω durch Null auf den Quader Q fort [möglich, da ω kompaktem Träger in $U_\xi \cap (\mathbb{R}_{\leq 0} \times \mathbb{R}^{n-1})$ besitzt], dann folgt diese letzte verbliebene Aussage (**) unmittelbar aus dem Satz von Stokes für Quader (Satz 4.34). $\square$

9.4 Standardintegral auf der Kugeloberfläche

Der Euklidische **Stern-Operator** $*$: $A^i(\mathbb{R}^n) \rightarrow A^{n-i}(\mathbb{R}^n)$ ist definiert durch

$$*(\sum_I f_I(x) dx_I) = \sum_I f_I(x) *dx_I$$

vermöge $*dx_I = \varepsilon \cdot dx_{I^c}$, wobei I^c die Komplementärmenge von I in $\{1, \dots, n\}$ ist und ε ein durch $dx_I \wedge *dx_I = \omega_n = dx_1 \wedge \dots \wedge dx_n$ bestimmtes Vorzeichen. Wie in Lemma 5.29 gezeigt wurde, gilt für $\eta \in A^i(\mathbb{R}^n)$ und Substitutionen M aus $SO(n, \mathbb{R})$

$$M^*(\eta) = *(M^*(\eta)),$$

d.h. der Stern-Operator vertauscht mit dem Pullback M^* . Daraus folgt

Lemma 9.4. Die Formen $\rho_0 = \frac{1}{2}r^2$ sowie $\rho_1 = d\rho_0$ und $\sigma_{n-1} = *\rho_1$ sind drehinvariant, d.h. invariant unter Pullbacks von Abbildungen $M \in SO(n, \mathbb{R})$, d.h. es gilt $M^*(\sigma_i) = \sigma_i$.

Beweis. Berücksichtigt man $M^*(r^2) = r^2 \iff {}^TMM = id$, ist die Aussage klar für ρ_0 . Da die Pullbacks M^* mit der Cartan Ableitung vertauschen, folgt die Aussage für σ_1 und wegen Lemma 5.29 damit auch für σ_{n-1} . $\square$

Im Fall der Euklidischen Einheitssphäre

$$S^{n-1} = \{x \in \mathbb{R}^n \mid \|x\| = 1\}$$

würden wir aber gerne etwas wie die Gesamtfläche definieren. Wie soll man das machen? Hierauf gibt es eine sehr befriedigende Antwort: Betrachte die drehinvariante Differentialform $\sigma_{n-1}(x) \in A^{n-1}(\mathbb{R}^n)$, die sogenannte **Oberflächenform**

$$\sigma_{n-1}(x) = *d(\frac{1}{2}r^2) = \sum_{i=1}^n x_i *dx_i.$$

Hierbei entsteht $*dx_i$ aus $\omega_n = dx_1 \wedge \dots \wedge dx_n$ bis auf ein Vorzeichen $(-1)^{i-1}$ durch Weglassen des Terms dx_i

$$*dx_i = \partial_i \lrcorner \omega_n = (-1)^{i-1} dx_1 \wedge \dots \wedge \widehat{dx_i} \wedge \dots \wedge dx_n.$$

9 Integration auf Mannigfaltigkeiten

Es gilt daher $dx_i \wedge *dx_i = \omega_n$. Zum Beispiel gilt $\sigma_1 = xdy - ydx$ im Fall $n = 2$ und es gilt $\sigma_2 = xdy \wedge dz - ydx \wedge dz + zdx \wedge dy$ im Fall $n = 3$. Optisch schöner

$$\sigma_2 = xdy \wedge dz + ydz \wedge dx + zdx \wedge dy .$$

Es gilt offensichtlich

$$\boxed{d\sigma_{n-1} = n \cdot \omega_n} .$$

Die Kugeloberfläche (das Wort Oberfläche ist natürlich ein Sprachmissbrauch in Dimensionen $n \neq 3$) ist dann definiert durch

$$\text{vol}(S) := \int_S \sigma_{n-1} .$$

Mit Hilfe der Kugeloberflächenform σ_{n-1} definieren wir jetzt allgemeiner für beliebige stetige Funktionen $f(x)$ auf der Sphäre $S = S^{n-1}$ ein abstraktes Integral

$$I : C(S^{n-1}) \rightarrow \mathbb{R} ,$$

das sogenannte **Standard Integral**. Wir erklären dieses Integral zuerst für Funktionen $f(x)$ aus dem Unterraum $A \subseteq C(S^{n-1})$,

$$I(f) := \int_{S^{n-1}} f(x) \cdot \sigma_{n-1}(x) \quad , \quad f(x) \cdot \sigma_{n-1}(x) \in A^{n-1}(\mathbb{R}^n) ,$$

welcher durch Einschränkungen von Funktionen in $f \in C^\infty(\mathbb{R}^n)$ erklärt ist (es würde sogar genügen den Raum der Einschränkungen aller Polynome im $\mathbb{R}^n$ auf S^{n-1} zu betrachten). Dieser Unterraum A ist eine Algebra und liegt nach Lemma 5.32 und dem Satz von Stone-Weierstraß dicht in $C(S^{n-1})$ bezüglich der Supremums-Norm. Auf Grund von Lemma 9.6 und Korollar 9.7 (weiter unten) kann man I (wie in Abschnitt 3.6) eindeutig zu einem abstrakten Integral auf $C(S^{n-1})$ fortsetzen.

Lemma 9.5. *Das Standardintegral $I(f)$ ist ein abstraktes Integral auf dem Verband A und ist invariant unter Drehungen¹ aus der Gruppe $SO(n, \mathbb{R})$, d.h.*

$$\boxed{I(f) = I(f \circ M) \quad , \quad M \in SO(n, \mathbb{R})} .$$

Beweis. Die Abbildung $I : A \rightarrow \mathbb{R}$ ist per Definition $\mathbb{R}$ -linear; die Drehinvarianz folgt aus der Drehinvarianz der Oberflächenform σ_{n-1} . Es verbleibt daher nur noch für $f(x) \geq 0$ zu zeigen $I(f) \geq 0$.

Mittels einer Partition der Eins kann man sich auf den Fall beschränken, daß der Träger von $f(x)$ in einer Karte enthalten ist. Durch eine Drehung kann man dann annehmen, der Träger von $f(x)$ sei enthalten in der rechten ‘Hemisphäre’ $S_+^{n-1} := \{x \in S^{n-1} \mid x_1 > 0\}$. Eine Kartenabbildung

$$\lambda : \{x = (x_2, \dots, x_n) \in \mathbb{R}^{n-1} \mid \rho^2 := \sum_{i=2}^n x_i^2 < 1\} \longrightarrow S_+^{n-1}$$

¹Eine $\mathbb{R}$ -lineare Abbildung $M : \mathbb{R}^n \rightarrow \mathbb{R}^n$ ist in der Gruppe $SO(\mathbb{R}^n)$ und erhält damit die Sphäre S^{n-1} , wenn $M = {}^T M^{-1}$ sowie $\det(M) = 1$ gilt. Siehe auch Abschnitt 5.13.

für die rechte Hemisphäre S_+^{n-1} ist für $\rho^2 := \|x\|_{\mathbb{R}^{n-1}}^2 = \sum_{i=2}^n x_i^2$ gegeben durch die lokale Parametrisierung

$$\lambda(x_2, \dots, x_n) = (+\sqrt{1-\rho^2}, x_2, \dots, x_n).$$

Die Behauptung $I(f) \geq 0$ für $f \geq 0$ folgt daher aus dem nächsten Lemma. $\square$

Lemma 9.6. Für Funktionen f mit kompaktem Träger in der rechten Hemisphäre S_+^{n-1} gilt

$$\int_{S_+^{n-1}} f(x) \cdot \sigma_{n-1}(x) = \int_{\|x\| < 1} \frac{f(\lambda(x_2, \dots, x_n))}{+\sqrt{1-\rho^2}} dx_2 \cdots dx_n.$$

Beweis. Bis auf $dx_2 \wedge \cdots \wedge dx_n$ ist der Pullback $\lambda^*(\sigma_{n-1})$ gleich

$$\sqrt{1-\rho^2} - \sum_{i=2}^n x_i \frac{\partial}{\partial x_i} \sqrt{1-\rho^2} = (1-\rho^2)^{-1/2} \cdot (1-\rho^2 + \sum_{i=2}^n x_i^2) = (1-\rho^2)^{-1/2}.$$

Korollar 9.7. Das Standard Integral I auf $C(S^{n-1})$ ist ein abstraktes Integral.

Beweis. Für monotone Folgen $f_n \nearrow f$ stetiger Funktionen $f_n \in C(S^{n-1})$ mit stetiger Grenzfunktion $f \in C(S^{n-1})$ ist zu zeigen $I(f_n) \nearrow I(f)$. Durch eine Partition der Eins kann man obdA annehmen, daß alle f, f_n kompakten Träger in der rechten Hemisphäre S_+^{n-1} besitzen. In diesem Fall folgt die Aussage sofort aus dem letzten Lemma und dem Satz von Beppo Levi angewendet auf das reelle Integral $\int_{\|x\| < 1} \frac{f(\lambda(x_2, \dots, x_n))}{+\sqrt{1-\rho^2}} dx_2 \cdots dx_n$. $\square$

Lemma 9.8. $\int_{\|x\| < r} f(x) \omega_n = \int_0^r \left(\int_S f(t\xi) \sigma_{n-1}(\xi) \right) t^{n-1} dt.$

Beweis. ObdA hat f Träger in der rechten Hemisphäre. Die Abbildung $\psi : \{\|\xi\| < 1\} \times (0, r) \rightarrow B_r(0)$ sei definiert durch $\psi(\xi_2, \dots, \xi_n; t) = (t\sqrt{1-\rho^2}, t\xi_2, \dots, t\xi_n)$. Sie hat dann die Eigenschaft $\det(D\psi)(\xi) = \frac{t^{n-1}}{\sqrt{1-\rho^2}}$ (Laplace Entwicklungssatz). Die Aussage folgt daher aus der Substitutionsformel Satz 4.27, sowie dem Satz von Fubini und Lemma 9.6. $\square$

Das letzte Lemma motiviert die Definition des Standardintegrals I_R auf der Sphäre $S(R)$ vom Radius R durch die Formel

$$I_R(f) := R^{n-1} \cdot \int_S f(R\xi) \sigma_{n-1}(\xi).$$

Lemma 9.8 schreibt sich dann suggestiver in der Form

$$\int_{\|x\| < r} f(x) \omega_n = \int_0^r I_t(f) dt.$$

9.5 Greensche Formel

Der Satz von Stokes für die abgeschlossenen Kugeln vom Radius R resp. r zeigt dann durch Subtraktion das folgende Resultat für die **Kugelschale** $X = \{x \in \mathbb{R}^n \mid r \leq \|x\| \leq R\}$:

9 Integration auf Mannigfaltigkeiten

Für $\eta \in A^{n-1}(U)$ und U offen in $\mathbb{R}^n$ mit $X \subset U$ gilt

$$\int_X d\eta = \int_{S^{n-1}(R)} \eta - \int_{S^{n-1}(r)} \eta.$$

Hierbei bezeichne $S^{n-1}(R)$ resp. $S^{n-1}(r)$ die Sphären vom Radius R resp. r . Wir schreiben die rechte Seite dieser Formel symbolisch wieder als $\int_{\partial X}$, wobei der Rand ∂X jetzt aus den beiden Sphären (mit unterschiedlicher Orientierung, nämlich $+$ für den äusseren Teil und $-$ für den inneren Teil) besteht

$$\int_X d\eta = \int_{\partial X} \eta.$$

Greensche Formel. Der Laplace Operator $\Delta = \sum_{i=1}^n \partial_i^2$ operiert auf Funktionen $f, g \in C^\infty(U)$. Für das Skalarprodukt $\langle f, g \rangle_{L^2(X)} = \int_X \bar{f}(x)g(x)dx$ im Hilbertraum $L^2(X)$ gilt

$$\boxed{\langle f, \Delta(g) \rangle_{L^2(X)} - \langle \Delta(f), g \rangle_{L^2(X)} = \int_X (f\Delta(g) - g\Delta(f))\omega_n}$$

(beachte $f = \bar{f}$ und $g = \bar{g}$) und man hat

Lemma 9.9. (Greensche Formel). Für Kugelschalen X gilt

$$\boxed{\int_X (f\Delta(g) - g\Delta(f))\omega_n = \int_{\partial X} f \sum_{i=1}^n \partial_i(g) * dx_i - \int_{\partial X} g \sum_{i=1}^n \partial_i(f) * dx_i}.$$

Beweis. Benutze $d(\sum_i f \partial_i(g) * dx_i) = \sum_i df \wedge \partial_i(g) * dx_i + \sum_i f \cdot d(\partial_i(g) * dx_i)$ wegen $d(*dx_i) = 0$, und damit $d(\sum_i f \partial_i(g) * dx_i) = \sum_{i=1}^n (\partial_i f)(\partial_i g) \omega_n + f \Delta(g) \omega_n$ wegen $df \wedge *dx_i = \sum_j (\partial_j f) dx_j \wedge *dx_i = (\partial_i f) \omega_n$. Vertauscht man f und g und bildet die Differenz, folgt daraus sofort das Lemma mit Hilfe des Satzes von Stokes. $\square$

10 Harmonische Analysis

10.1 Der Hilbertraum $L^2(S)$

Sei $n \geq 2$ und sei $S = S^{n-1} = \{x \in \mathbb{R}^n \mid \|x\| = 1\}$ die **Einheitssphäre** im $\mathbb{R}^n$. Das in Abschnitt 9.4 für stetige Funktionen $f \in C(S)$ erklärte Standard Integral $I(f) = I_S(f)$ ist nach Korollar 9.7 ein abstraktes Integral. Daher kann man den Hilbertraum $L^2(S)$ definieren mit dem Skalarprodukt $\langle f, g \rangle = I_S(\bar{f} \cdot g)$, d.h.

$$\langle f, g \rangle = \int_S \bar{f}(x) g(x) \sigma_{n-1}(x)$$

gebildet zur Differentialform $\sigma_{n-1}(x) = \sum_{i=1}^n x_i * dx_i$ (siehe Abschnitt 9.4).

Der $\mathbb{C}$ -Vektorraum aller Polynome auf $\mathbb{R}^n$ ist eine Algebra. Da bereits lineare Funktionen Punkte trennen, ist diese Algebra punktetrennend auf $\mathbb{R}^n \setminus \{0\}$. Nach Stone-Weierstraß ist daher der Raum A der Einschränkungen aller Polynome auf S ein dichter Unterraum des Raumes $C(S)$ der stetigen Funktionen auf S , denn S ist kompakt. Nach Lemma 5.22 wird A von den Einschränkungen der harmonischen Polynome auf $\mathbb{R}^n$ aufgespannt. Es folgt

Lemma 10.1. *Der Vektorraum aller harmonischen Polynome liegt dicht in $L^2(S)$.*

In Lemma 10.6.4 zeigen wir, daß homogene harmonische Polynome verschiedenen Grades orthogonal zueinander sind bezüglich des Skalarproduktes von $L^2(S)$. Wir beschränken uns daher für den Moment auf homogene *harmonische* Polynome in $\mathcal{H}_l(\mathbb{R}^n)$ festes Grad l . Das Skalarprodukt $\langle P, Q \rangle$ von $L^2(S)$ liefert eine positiv definite symmetrische $\mathbb{R}$ -Bilinearform auf $\mathcal{H}_l(\mathbb{R}^n)$. Sei $\{P_{l,k}(x) \in \mathcal{H}_l(\mathbb{R}^n) \mid 1 \leq k \leq \dim \mathcal{H}_l(\mathbb{R}^n)\}$ eine **reelle ON-Basis** $P_{l,k}(x)$ dieses Skalarproduktes auf $\mathcal{H}_l(\mathbb{R}^n)$. Verschwindet ein Polynom in $\mathcal{H}_l(\mathbb{R}^n)$ auf der Sphäre S , dann ist es wegen der Homogenität das Nullpolynom. Die Einschränkung $\mathcal{H}_l(\mathbb{R}^n) \rightarrow C(S)$ ist also eine injektive Abbildung. Wir definieren dann für $x, \xi \in \mathbb{R}^n$

$$Z_l(x, \xi) := \sum_{k=1}^{\dim(\mathcal{H}_l(\mathbb{R}^n))} P_{l,k}(x) P_{l,k}(\xi) .$$

Für festes ξ ist die Funktion $Z_l(x, \xi)$ ein harmonisches Polynom der Variable x vom Grad l , und umgekehrt gilt dies auch. Nach Definition gilt

$$Z_l(x, \xi) = Z_l(\xi, x) \quad , \quad Z_l(\lambda x, \lambda \xi) = \lambda^{2l} Z_l(x, \xi) .$$

$Z_l(x, \xi)$ ist **vollkommen bestimmt** durch die folgende **Reproduktionsformel**

$$\langle Z_l(x, \xi), P(x) \rangle = P(\xi).$$

[Diese folgt unmittelbar aus $\langle \sum_k a_k P_{l,k}(x), P_{l,j}(x) \rangle = a_j$.] Aus der Drehinvarianz des Standard Integrals auf S folgt $\langle Z_l(Mx, M\xi), P(x) \rangle = \langle Z_l(x, M\xi), P(M^{-1}x) \rangle$ für $M \in SO(n, \mathbb{R})$. Aber $\langle Z_l(x, M\xi), P(M^{-1}x) \rangle = P(M^{-1}M\xi) = P(\xi)$, also $\langle Z_l(Mx, M\xi), P(x) \rangle = P(\xi)$. Aus obiger Eindeutigkeitsaussage folgt daher

$$Z_l(Mx, M\xi) = Z_l(x, \xi) \quad , \quad M \in SO(n, \mathbb{R}) .$$

Insbesondere ist $Z_l(x, \xi)$ als Funktion von $x \in S$ (bei festem $\xi \in S$) damit invariant unter dem Stabilisator $SO(n-1, \mathbb{R})$ des Punktes ξ . Im Spezialfall $\xi = (1, 0, \dots, 0)$ hängt daher $Z_l(x, \xi)$, $x \in \mathbb{R}^n$ nur von $x_1 = x \cdot \xi$ und $r = \|x\|$ ab, ist daher proportional zu dem **zonalen harmonischen Polynom** $P_{l,0}(x)$ vom Grad l (definiert in Abschnitt 5.10). D.h. man erhält die **Additionsformel**

$$Z_l(x, \xi) = \text{const}(l) \cdot P_{l,0}(x) \quad , \quad \xi = (1, 0, \dots, 0) .$$

Lemma 10.2. Es gilt $|Z_l(x, \xi)| \leq Z_l(x, x) = \frac{\dim(\mathcal{H}_l(\mathbb{R}^n))}{\text{vol}(S)}$ für alle $x, \xi \in S$.

Beweis. Wegen $Z_l(x, x) = \sum_k P_{l,k}(x)^2$ und der $SO(n, \mathbb{R})$ -Invarianz ist $Z_l(x, x) = c$ eine konstante Funktion auf S . Also $\text{vol}(S)c = \int_S Z_l(x, x) \sigma_{n-1} = \sum_k \langle P_{l,k}, P_{l,k} \rangle = \dim_{\mathbb{R}}(\mathcal{H}_l(\mathbb{R}^n))$. Es gilt $|Z_l(x, \xi)| = |\sum_k P_{l,k}(x) P_{l,k}(\xi)| \leq |Z_l(x, x)|^{1/2} |Z_l(\xi, \xi)|^{1/2} = Z_l(x, x)$. Dies benutzt die Schwarz Ungleichung und die Konstanz von $Z_l(x, x) \geq 0$ auf S . $\square$

10.2 Poisson Kern

Sei U offen in $\mathbb{R}^n$ und $f : U \rightarrow \mathbb{R}$ eine harmonische Funktion. Insbesondere gilt $f \in C^\infty(U)$. Ist $x_0 = 0$ in U , dann sind die Taylor Koeffizienten

$$c_l(x) = \frac{1}{l!} \left(\frac{d}{dt} \right)^l f(tx)_{t=0}$$

harmonische Polynome in $\mathcal{H}_l(\mathbb{R}^n)$; siehe Abschnitt 5.12.

Wir wenden dies auf den **Poisson Kern** $P(x, \xi)$ an im Bereich $\|x\| < \|\xi\|$, d.h. auf

$$P(x, \xi) = \frac{1}{\text{vol}(S)} \frac{\|\xi\|^2 - \|x\|^2}{\|\xi - x\|^n} \quad , \quad x \neq \xi .$$

Beachte $\|x\|^2 - \|\xi\|^2 = \|x - \xi\|^2 + 2(x - \xi, \xi)$. Also $-\text{vol}(S)P(x, \xi) = \|x - \xi\|^{-\kappa} + \frac{2(x - \xi, \xi)}{\|x - \xi\|^{\kappa+2}} = P_0^*(x - \xi) + P_1^*(x - \xi)$. Dies ist eine Summe von Kelvin Transformaten der harmonischen Polynome $P_0(x) = 1$ und $P_1(x) = 2(x, \xi)$ mit $\kappa = n - 2$. Also ist $P(x, \xi)$ eine harmonische Funktion in $x - \xi$ und damit harmonisch in x , und wegen $P(x, \xi) = -P(\xi, x)$ damit auch harmonisch in ξ .

Offensichtlich gilt $P(Mx, M\xi) = P(x, \xi)$ für alle $M \in SO(n, \mathbb{R})$. Die Taylor Koeffizienten $c_{l,\xi}(x) := T_l(P(x, \xi))(x)$ sind dann harmonische Polynome mit derselben Drehsymmetrie.

Somit sind $c_{l,\xi}(x)$ und $Z_l(x, \xi)$ **zonale harmonische Polynome** in x vom Grad l und daher proportional zu $P_{l,0}(x)$ nach Abschnitt 5.10. [OBdA ist ξ gleich $\xi_0 = (1, \dots, 0)$ wegen der Drehsymmetrie. Dann hängt $\text{vol}(S)P(x, \xi_0) = \frac{1-r^2}{(1-2x_1+r^2)^{n/2}}$ nur ab von der ersten Koordinate $x_1 = (x, \xi_0)/2$ und dem Radius $r = \|x\|$. Dasselbe gilt daher auch für die Taylor Koeffizienten $c_l(x) = c_l(x, \xi_0)$]. Daraus folgern wir weiter unten

Satz 10.3. Es gilt $c_l(x, \xi) = Z_l(x, \frac{\xi}{\|\xi\|^2})\|\xi\|^{-\kappa}$ für alle $l \in \mathbb{N}$, und für alle $x, \|x\| < \|\xi\|$ gilt

$$P(x, \xi) = \sum_{l=0}^{\infty} Z_l(x, \frac{\xi}{\|\xi\|^2})\|\xi\|^{-\kappa}.$$

Die Potenzreihe auf der rechten Seite konvergiert auf jeder kompakten Teilmenge von der Kugel $\{x \in \mathbb{R}^n \mid \|x\| < \|\xi\|\}$ absolut und gleichmässig.

Bemerkung. Für $x, \xi \in S$ sind alle Koeffizienten $Z_l(x, \xi)$ symmetrisch in x und ξ . Dennoch ist die Summe $P(x, \xi)$ antisymmetrisch in x und ξ . Zur Illustration dieses scheinbaren Widerspruchs betrachte dazu im eindimensionalen Fall $(\xi - x)^{-1} = \sum_{l=0}^{\infty} z_l(x, \frac{\xi}{\xi^2})\xi^{-1}$ für $z_l(x, \xi) = x^l \xi^l$ im Bereich $|x| < |\xi|$ (im wesentlichen die geometrische Reihe).

Beweis. Durch Reskalierung $(x, \xi) \mapsto (tx, t\xi)$ ist wegen $P(tx, t\xi) = t^{-\kappa}P(x, \xi)$ und der analogen Reskalierungseigenschaft der rechten Seite oBdA $\xi \in S$ und $\|x\| < 1$. Auf Grund der Drehinvarianz beider Seiten ist dann oBdA $\xi = \xi_0 = (1, 0, \dots, 0)$. Wir zeigen nun, daß auf der rechten Seite $\sum_{l=0}^{\infty} Z_l(x, \xi_0)$ für $\|x\| < 1$ absolut konvergiert, und in diesem Bereich damit eine drehinvariante harmonische Funktion $f(x)$ darstellt. Dazu benützen wir die Abschätzung $|Z_l(x, \xi_0)| \leq t^l Z_l(\xi_0, \xi_0)$, welche aus Lemma 10.2 folgt, und können für die absolute Konvergenz oBdA annehmen $x = t\xi_0$. Aus $\text{vol}(S)Z_l(t\xi_0, \xi_0) = \dim(\mathcal{H}_l(\mathbb{R}^n)) \cdot t^l$ (Lemma 10.2) und $\kappa \cdot \dim(\mathcal{H}_l) = (2l + \kappa) \binom{n+l-3}{l}$ oBdA für $\kappa = n - 2 \geq 1$ (Satz 5.20), sowie

$$\sum_l \frac{2l + \kappa}{\kappa} \cdot \binom{n+l-3}{l} t^l = \frac{1}{(1-t)^\kappa} + \frac{2t}{\kappa} \frac{d}{dt} \frac{1}{(1-t)^\kappa} = \frac{1-t^2}{(1-t)^n} = \frac{\|\xi_0\|^2 - \|t\xi_0\|^2}{\|\xi_0 - t\xi_0\|^n},$$

folgt $f(t\xi_0) = P(t\xi_0, \xi_0)$. Wir haben dabei benutzt $\frac{1}{t^l} \left(\frac{d}{dt}\right)^l (1-t)^{-\kappa} \big|_{t=0} = \binom{\kappa+l-1}{l}$ für $\kappa > 0$. Im Fall $n = 2$ reduziert sich alles auf $-1 + 2/(1-t) = (1-t^2)/(1-t)^2$. Dies zeigt die absolute und gleichmässige Konvergenz.

Wegen der Proportionalität der zonalen harmonischen Polynome $c_l(x, \xi) = \text{const}(l) \cdot Z_l(x, \xi)$ genügt wegen $Z_l(\xi_0, \xi_0) > 0$ für $\text{const}(l) = 1$ die gezeigte Gleichheit in den Punkten $x = t \cdot \xi_0$. Dies zeigt, daß alle Taylor Koeffizienten der harmonischen C^∞ -Funktion $f(x) - P(x, \xi_0)$ im Mittelpunkt $x_0 = 0$ der Kugel $U = \{x \mid \|x\| < 1\}$ verschwinden. Aus Lemma 10.8 im nächsten Abschnitt folgt daher $f(x) = P(x, \xi_0)$ auf U . $\square$

Eine geringfügige Modifikation des Beweises zeigt analog

Satz 10.4. Für $n \geq 3$ und $\kappa = n - 2$ sei $Q(x, \xi)$ das **Coulomb oder Newton Potential**

$$Q(x, \xi) = \frac{1}{\kappa \cdot \text{vol}(S)} \frac{1}{\|x - \xi\|^\kappa}$$

im $\mathbb{R}^n$. Für alle x mit $\|x\| < \|\xi\|$ gilt dann

$$Q(x, \xi) = \sum_{l=0}^{\infty} \frac{1}{2l + \kappa} \cdot Z_l(x, \frac{\xi}{\|\xi\|^2}) \|\xi\|^{-\kappa}.$$

Die Reihe konvergiert auf jedem Kompaktum der Kugel $\{x \mid \|x\| < \|\xi\|\}$ absolut und gleichmäßig.

10.3 Orthogonalität

In diesem Abschnitt sei $n \geq 2$. Für reelle Zahlen $0 < \rho < R$ und $r \in [\rho, R]$ sei $X = X[r, R]$ die abgeschlossene **Kugelschale** im $\mathbb{R}^n$ mit den Radien r und R um $x_0 = 0$. Sei $V \subseteq \mathbb{R}^n$ eine offene Teilmenge, welche X für alle $r \in [\rho, R]$ enthält. Für harmonische Funktionen $f(x), g(x) \in C^\infty(V)$ verschwindet die linke Seite der Greenschen Formel wegen $\Delta(f) = \Delta(g) = 0$. Die Greensche Formel (Lemma 9.9) liefert damit folgende **Vertauschungsformel**:

$$\int_{\partial X} f \sum_{i=1}^n \partial_i(g) * dx_i = \int_{\partial X} g \sum_{i=1}^n \partial_i(f) * dx_i.$$

Damit beweisen wir das nächste Lemma. Beachte, ∂X ist die Vereinigung der beiden Sphären $S(R), S(r)$ vom Radius R und r mit unterschiedlicher Orientierung

$$\partial X = S(R) - S(r).$$

Die Vertauschungsformel zeigt daher, daß $\int_{S(r)} f \sum_i \partial_i(g) * dx_i - \int_{S(r)} g \sum_i \partial_i(f) * dx_i$ unabhängig vom Radius $r \in [\rho, R]$ ist.

Definition. Für harmonisches $f(x) \in C^\infty(V)$ und $P(x) \in \mathcal{H}_l(\mathbb{R}^n)$ und $r \in [\rho, R]$ setzen wir

$$\langle f, P \rangle_r := r^{-n-2l} \int_{S(r)} f(x) P(x) \sigma_{n-1}(x).$$

Lemma 10.5. Sei $P(x)$ ein harmonisches Polynom auf $\mathbb{R}^n$ und $f(x)$ eine harmonische Funktion auf einer offenen Menge $V \subseteq \mathbb{R}^n$, welche $\bigcup_{r \in [\rho, R]} S(r)$ enthält. Dann existieren Konstanten $\alpha = \alpha(P, f)$ und $\beta = \beta(P, f)$, welche nicht von $r \in [\rho, R]$ abhängen, so daß für $\kappa = n - 2$ gilt

$$(\kappa + 2l) \cdot \langle f, P \rangle_r = \alpha + \beta \cdot r^{-\kappa-2l} \quad \text{bzw. für } (n, l) = (2, 0) \quad \langle f, P \rangle_r = \alpha + \beta \cdot \log(r).$$

Beweis. Die **Kelvin Transformierte** $g(x) = P^*(x) = P(x) \|x\|^{-\kappa-2l}$ des harmonischen Polynoms $P(x)$ hat eine *Singularität* im Punkt $x_0 = 0$ [für $(n, l) = (2, 0)$ setzt man analog $g(x) = \log(\|x\|)$], ist aber harmonisch auf der offenen Menge $V \setminus \{x_0\}$, welche $X = X(r, R)$

enthält. Wir wenden nun für $g = P^*$ und f die oben angegebene **Vertauschungsformel** an. Wegen

$$\partial_i g(x) = -(\kappa + 2l) \frac{P(x)x_i}{\|x\|^{n+2l}} + \frac{\partial_i P(x)}{\|x\|^{\kappa+2l}}$$

gilt

$$f(x) \sum_i \partial_i(g) * dx_i = -(\kappa + 2l) \frac{f(x)P(x)\sigma_{n-1}(x)}{\|x\|^{n+2l}} + \frac{f(x) \sum_i \partial_i(P) * dx_i}{\|x\|^{\kappa+2l}}.$$

Die Nenner sind Potenzen von r , also konstant auf $S(r)$. Nach der Vertauschungsformel ist daher $-\alpha := \int_{S(r)} f \sum_i \partial_i(g) * dx_i - \int_{S(r)} g \sum_i \partial_i(f) * dx_i$ unabhängig von $r \in [\rho, R]$. Konkret ist

$$-\alpha = -(\kappa + 2l) \langle f, P \rangle_r + \frac{1}{r^{\kappa+2l}} \int_{S(r)} f(x) \sum_i \partial_i(P) * dx_i - \frac{1}{r^{\kappa+2l}} \int_{S(r)} P(x) \sum_i \partial_i(f) * dx_i.$$

Analog zeigt die Vertauschungsformel angewandt für $g = P$ und f die r -Unabhängigkeit von

$$\beta := \int_{S(r)} f(x) \sum_i \partial_i(P) * dx_i - \int_{S(r)} P(x) \sum_i \partial_i(f) * dx_i.$$

Beides zusammen liefert unsere Behauptung. $\square$

Für $(n, l) = (2, 0)$ überlassen wir es dem Leser $\langle f, 1 \rangle_r = \frac{1}{r^2} \int_{S(r)} f(x) \sigma_1(x) = \alpha + \beta \cdot \log(r)$ zu zeigen mit einem analogen Argument.

Lemma 10.6. Sei U eine offene Kugel vom Radius $> R$ um $x_0 = 0$. Wir nehmen an $1 < R$. Dann gilt für harmonisches $f(x) \in C^\infty(U)$ und harmonisches $P(x) \in \mathcal{H}_l(\mathbb{R}^n)$

1. Nach Definition gilt $\langle f, P \rangle_1 = \langle f, P \rangle$ für das Skalarprodukt $\langle f, P \rangle = \langle f, P \rangle_{L^2(S)}$.
2. $\langle f, P \rangle_r$ hängt nicht ab von der Wahl von r .
3. Verschwinden die Taylor Koeffizienten $T_\nu(f)(x)$ für $\nu \leq l$, dann gilt $\lim_{r \rightarrow 0} \langle f, P \rangle_r = 0$.
4. Orthogonalität: Es gilt $\langle \mathcal{H}_m(\mathbb{R}^n), \mathcal{H}_l(\mathbb{R}^n) \rangle = 0$ für $m \neq l$.
5. Es gilt $\langle f, P \rangle = \lim_{r \rightarrow 0} \langle f, P \rangle_r = \langle T_l(f), P \rangle$.

Beweis. Behauptung 2 folgt aus Lemma 10.5, da $f(x)$ nach Annahme stetig im Punkt $x_0 = 0$ ist und deshalb für $r \rightarrow 0$ das Integral $\langle f, P \rangle_r$ beschränkt bleibt. Daraus folgt $\beta = 0$ und $(\kappa + 2l) \cdot \langle f, P \rangle_r = \alpha$ in Lemma 10.5. Beachte $\kappa + 2l > 0$ ausser im Fall $(n, l) = (2, 0)$; dieser geht aber analog. Für Behauptung 3 betrachten wir die Taylor Koeffizienten $T_\nu(f)$ von $f(x)$ im Punkt $x_0 = 0$ (siehe Abschnitt 5.12). Aus $T_\nu(f) = 0$ für alle $\nu \leq l$ und Lemma 5.26 folgt $f(x) = \|x\|^l \cdot H(x)$ für eine stetige Funktion H auf U mit $H(0) = 0$. Wegen $f(rx)P(rx)\sigma_{n-1}(rx) = r^{2l+n} \cdot H(rx)P(x)\sigma_{n-1}(x)$ für $x \in S = S(1)$ folgt Behauptung 3 aus $\lim_{r \rightarrow 0} H(rx) = 0$. Behauptung 4 folgt aus den Behauptungen 1,2,3, denn wegen $\langle f, g \rangle = \overline{\langle g, f \rangle}$ ist obdA $m > l$. Für die harmonische Funktion $h(x) = f(x) - \sum_{\nu=0}^l T_\nu(f)(x)$ folgt $\lim_{r \rightarrow 0} \langle h, P \rangle_r = 0$ aus Behauptung 3 und Abschnitt 5.12. Wegen Behauptung 1 und 2 ist daher $\langle h, P \rangle = 0$ und damit $\langle f, P \rangle = \sum_{m=0}^l \langle T_m(f), P \rangle$, wegen Behauptung 4 also $\langle f, P \rangle = \langle T_l(f), P \rangle$. Dies zeigt Behauptung 5. $\square$

Lemma 10.6.4 und Lemma 10.1 zusammen ergeben

Satz 10.7. Die Funktionen $P_{l,k}(x)$ bilden eine Hilbertraum-Basis von $L^2(S)$.

Aus Lemma 10.6.1-3 folgt daher

Lemma 10.8. Verschwinden alle Taylor Koeffizienten einer harmonischen Funktion $f \in C^\infty(U)$ im Mittelpunkt x_0 einer offenen Kugel U , dann verschwindet f auf U .

Irreduzibilität. Es existiert kein¹ $\mathbb{R}$ -Untervektorraum $0 \neq V \subsetneq \mathcal{H}_l(\mathbb{R}^n)$ mit der Eigenschaft: $P(x) \in V \Rightarrow P(Mx) \in V$ für alle $M \in SO(V)$.

10.4 Harmonische Funktionen sind analytisch

Sei $U \subset \mathbb{R}^n$ eine offene Menge und

$$f : U \rightarrow \mathbb{R}$$

eine harmonische Funktion auf U . Sei die abgeschlossene Kugel vom Radius R um x_0 in U enthalten

$$\{x \mid \|x - x_0\| \leq R\} \subset U.$$

Für die Sphäre $S(R, x_0)$ mit Mittelpunkt x und Radius R gilt dann

Satz 10.9 (Poisson Formel). Unter obigen Voraussetzungen gilt für alle x in der Kugel B definiert durch $\|x\| < R$

$$f(x + x_0) = \frac{1}{R^2 \text{vol}(S)} \int_{S(R,0)} f(\xi + x_0) \frac{\|\xi\|^2 - \|x\|^2}{\|\xi - x\|^n} \sigma_{n-1}(\xi) d\xi.$$

Für $x = 0$ liefert dies folgende **Mittelpunktsformel**

Satz 10.10. Sei $n \geq 2$ und U offen im $\mathbb{R}^n$ und $f \in C^\infty(U)$ eine harmonische Funktion. Für jede abgeschlossene Kugel vom Radius R in U mit Mittelpunkt x_0 gilt

$$f(x_0) = \frac{1}{\text{vol}(S(R))} \cdot \int_{S(R,0)} f(\xi + x_0) \sigma_{n-1}(\xi) d\xi.$$

Beweis. Zum Beweis der Poisson Formel ist obdA $x_0 = 0$ und $R = 1$. Das Integral auf der rechten Seite der Poisson Formel definiert eine Funktion $g(x)$ auf B . Wegen Lemma 10.8 genügt es, daß alle (höheren) Ableitungen von f und g in x_0 übereinstimmen. Zum Beweis entwickeln wir g auf B mittels Satz 10.4 in eine konvergente Potenzreihe². Vertauschen von Summation

¹Definiere $Z_l^V(x, \xi)$ mit einer ON-Basis von V analog zu $Z_l(x, \xi)$. Dann gilt analog eine Reproduktionsformel auf V und damit $Z_l^V(Mx, M\xi) = Z_l^V(x, \xi)$ für $M \in SO(V)$. Daher sind $Z_l^V(x, \xi)$ und $Z(x, \xi)$ proportional [setze $\xi = (1, 0, \dots, 0)$] und die Reproduktionsformeln für $Z_l^V(x, \xi)$, $Z(x, \xi)$ zeigen $V = \mathcal{H}_l(\mathbb{R}^n)$ oder $V = 0$.

²Die Konvergenz ist gleichmäßig für festes $\xi \in S$ und alle x aus einer kompakten Teilkugel von B . Wegen der $SO(n, \mathbb{R})$ -Invarianz von $P(x, \xi)$ ist sie daher auch gleichmäßig für alle solchen x und alle $\xi \in S$.

und Integration (Korollar 6.14) liefert

$$g(x) = \int_S f(\xi) P(x, \xi) \sigma_{n-1}(\xi) = \sum_{l=0}^{\infty} \int_S f(\xi) Z_l(x, \xi) \sigma_{n-1}(\xi) .$$

Die Summanden $\int_S f(\xi) Z_l(x, \xi) \sigma_{n-1}(\xi) = \langle Z_l(\xi, x), f(\xi) \rangle$ sind homogen vom Grad l in x , somit gleich den Taylor Koeffizienten $T_l(g)(x)$. Also ist $\langle Z_l(\xi, x), f(\xi) \rangle = \langle Z_l(\xi, x), T_l(f)(\xi) \rangle$ nach Lemma 10.6.5, und wegen der Reproduktionsformel in Abschnitt 10.1 dann gleich $T_l(f)(x)$. Also gilt $T_l(f)(x) = T_l(g)(x)$ für alle l . $\square$

Das benutzte Argument zeigt unter gleichen Voraussetzungen

Satz 10.11. Die harmonische Funktion $f(x)$ lässt sich um x_0 (hier obdA $x_0 = 0$) in eine konvergente Potenzreihe vom Konvergenzradius $\geq R$ entwickeln

$$f(x) = \sum_{l=0}^{\infty} P_l(x) .$$

Die Funktionen $P_l(x) \in \mathcal{H}_l(\mathbb{R}^n)$ sind harmonische Polynome auf $\mathbb{R}^n$ und homogen vom Grad l , und sind wie folgt durch f bestimmt

$$P_l(x) = \frac{1}{R^2} \int_S f(\xi) Z_l(x, \xi) \sigma_{n-1}(\xi) .$$

Bemerkung. Der Konvergenzradius der Potenzreihenentwicklung ist daher mindestens (!) so groß wie der Radius R jeder Vollkugel um x_0 , die vollkommen im Definitionsbereich U von f enthalten ist.

Aus der Mittelpunktsformel kann man ohne große Mühe das sogenannte **Maximumsprinzip** folgern: Nimmt eine harmonische Funktion $f : U \rightarrow \mathbb{R}$ auf einer offenen Menge $U \subset \mathbb{R}^n$ ihr Maximum (Minimum) in einem Punkt $x_0 \in U$ an, dann ist f konstant auf jeder offenen Kugel in U mit Mittelpunkt x_0 .

10.5 Entwicklung auf Kugelschalen

Sei $X = X[\rho, R]$ eine abgeschlossene Kugelschale und $V \subseteq \mathbb{R}^n$ eine offene Menge, die X enthält. ObdA sei $0 < \rho < R$ und

$$f : V \rightarrow \mathbb{R} \quad \text{harmonisch} .$$

Satz 10.12. Auf V harmonische Funktionen $f(x)$ lassen sich auf Kugelschalen $X[\rho, R] \subset V$ in absolut und gleichmässig konvergente Reihen entwickeln der Gestalt (siehe Abschnitt 5.11)

$$f(y) = \sum_{l=0}^{\infty} \sum_{k=1}^{\dim(\mathcal{H}_l(\mathbb{R}^n))} (a_{lk} \cdot P_{l,k}(y) + b_{lk} \cdot P_{l,k}^*(y)) .$$

Im Fall $(n, l) = (2, 0)$ ist hierbei $P_{0,1}^*(x)$ formal durch $\log(r)$ zu ersetzen.

Beweis. Sei $n > 2$; der Fall $n = 2$ geht analog. Wir entwickeln f in eine Potenzreihe auf $X[\rho, R]$. Die Einschränkung von f auf jede Sphäre $S(r)$ für $r \in [\rho, R]$ definiert eine stetige Funktion auf $S(r)$, und ist somit in $L^2(S(r))$ enthalten. Für harmonische Polynome $P \in \mathcal{H}_l(\mathbb{R}^n)$ gilt daher durch Transformation auf die Einheitssphäre S , mittels $y = r\xi$ und der Notation $f_r(\xi) := f(r\xi)$,

$$\int_S f_r(\xi) P(\xi) \sigma_{n-1}(\xi) = r^{-n-l} \int_{S(r)} f(y) P(y) \sigma_{n-1}(y) = ar^l + br^{-l-\kappa}$$

für Konstanten a, b , welche nur von f, P abhängen aber nicht von r . Letzteres folgt aus Lemma 10.5, wobei wir der Einfachheit halber $(n, l) \neq (2, 0)$ angenommen haben! Anwenden von $r \frac{d}{dr}$ liefert³ auf Grund des Vertauschungssatzes 4.32

$$\int_S (Ef)_r(\xi) P(\xi) \sigma_{n-1}(\xi) = l \cdot ar^l - (l + \kappa) \cdot br^{-l-\kappa},$$

und damit $\frac{2l+\kappa}{l+\kappa} \cdot ar^l = \langle f_r + \frac{1}{l+\kappa} (Ef)_r, P \rangle$, beziehungsweise wegen $(r \frac{d}{dr})^m f_r = (E^m f)_r$

$$\frac{(2l + \kappa)l^m}{l + \kappa} \cdot ar^l = \langle (E^m f)_r + \frac{1}{l + \kappa} (E^{m+1} f)_r, P \rangle.$$

Das Skalarprodukt rechts kann durch $c_1 \cdot \max_{\xi \in S} |P(\xi)|$ abgeschätzt werden. Die Konstante c_1 hängt dabei nur von f und m ab und nicht von l . Für $P = P_{l,k}$, $a = a_{lk}$ und $l > 0$ liefert dies $|a_{lk}r^l| \leq C_1 \cdot l^{-m} \max_{\xi \in S} |P_{l,k}(\xi)|$ für eine Konstante C_1 , welche nur von f und m abhängt. Dies schätzt $\sum_{k=1}^{\dim(\mathcal{H}_l(\mathbb{R}^n))} |a_{lk} \cdot P_{l,k}(r\xi)|$ durch $C_1 l^{-m} \cdot \sum_{k=1}^{\dim(\mathcal{H}_l(\mathbb{R}^n))} \max_{\xi \in S} (P_{l,k}(\xi)^2)$ ab. Aus Lemma 10.2 folgt $0 \leq P_{l,k}(\xi)^2 \leq Z_l(\xi, \xi) \leq \frac{\dim(\mathcal{H}_l(\mathbb{R}^n))}{\text{vol}(S)}$. Wegen $\frac{\dim(\mathcal{H}_l(\mathbb{R}^n))}{\text{vol}(S)} \leq C_2 \cdot l^n$ (Satz 5.20) gilt also

$$\sum_{k=1}^{\dim(\mathcal{H}_l(\mathbb{R}^n))} |a_{lk} \cdot P_{l,k}(r\xi)| \leq \text{vol}(S) C_1 C_2^2 \cdot l^{2n-m} \leq \text{const.} \cdot l^{-2},$$

da wir obdA $m = 2n + 2$ wählen können. Ähnlich kann man $\sum_k b_{lk} P_{l,k}^*(r\xi)$ abschätzen. Gilt $x \in S$ und $r \in [\rho, R]$, folgt aus $\sum_{l=1}^{\infty} l^{-2} < \infty$ daher für $y = r \cdot x \in X[\rho, R]$ die absolute und gleichmässige Konvergenz der Reihe auf der rechten Seite der noch zu zeigenden Identität

$$f_r(x) = \sum_{l=0}^{\infty} \sum_{k=1}^{\dim(\mathcal{H}_l(\mathbb{R}^n))} (a_{lk}r^l + b_{lk}r^{-l-\kappa}) \cdot P_{l,k}(x).$$

Damit ist die rechte Seite der Formel in Satz 10.12 wohldefiniert und harmonisch auf ihrem Definitionsbereich $X[\rho, R]$. Weiterhin haben die linke und die rechte Seite der Formel in Satz 10.12 die selben Skalarprodukte mit allen Funktionen der Hilbertraum-Basis $P_{l,k}(x)$ (Satz 10.7) von $L^2(S)$. Deshalb sind beide Seiten f.ü. gleich als Funktion auf S (Korollar 8.9) bei festem r . Damit sind beide Seiten aus Stetigkeitsgründen (Satz 2.24 und Lemma 8.11) gleich auf ganz S . Aus $P_{k,l}(y) = r^l P_{k,l}(x)$ resp. $P_{k,l}^*(y) = r^{-l-\kappa} P_{k,l}(x)$ resp. $f(y) = f_r(x)$ folgt daher die Behauptung (im Fall $n > 2$). $\square$

³Für das Eulerfeld $E = \sum_{i=1}^n x_i \partial_i$ gilt $r^2 \Delta(Ef) = E(r^2 \Delta(f)) = 0$. Deshalb ist auch Ef harmonisch auf V , sowie dann alle $E^m f$ für $m \in \mathbb{N}$. Aus der Kettenregel folgt $(Ef)(r\xi) = r \frac{d}{dr} f(r\xi)$ für reelles $r > 0$ und $\xi \in \mathbb{R}^n$.

Satz 10.13 (Hebbarkeitssatz). Sei $B = \{x \in \mathbb{R}^n \mid \|x - x_0\| < R\}$ und $f : B \setminus \{x_0\} \rightarrow \mathbb{R}$ eine harmonische Funktion. Ist f beschränkt auf $\{x \neq x_0 \in \mathbb{R}^n \mid \|x - x_0\| \leq r\}$ für ein r mit $0 < r < R$, dann lässt sich f zu einer harmonischen Funktion auf ganz B fortsetzen.

Aus Lemma 10.5 folgt wie im Beweis von Lemma 10.6.2 aus der Beschränktheit von f das Verschwinden der Koeffizienten $b = b_{kl}$. Damit folgt das Resultat leicht aus Satz 10.12.

10.6 Die Potential Gleichung

Für eine gegebene C^∞ -Funktion $\rho : \mathbb{R}^n \rightarrow \mathbb{R}$ (hier der Einfachheit halber $n \geq 3$) suchen wir Lösungen $u : \mathbb{R}^n \rightarrow \mathbb{R}$ der **Potential Gleichung** oder auch **Laplace Gleichung**

$$\boxed{-\Delta u(x) = \rho(x)}.$$

Hat man eine Lösung $u(x)$ gefunden, dann ist jede andere Lösung von der Gestalt $u(x) + f(x)$ für eine harmonische Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$. Unter geeigneten Bedingungen an das Abklingen der Lösungen im Unendlichen (etwa $\leq \text{const} \cdot r^{-\varepsilon}$ für ein $\varepsilon > 0$) ist die Lösung der Poisson Gleichung eindeutig, d.h. es gilt $f(x) = 0$. Dazu benutzt man

Satz 10.14. Beschränkte harmonische Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}$ sind konstant.

Beweis. (Im Prinzip ist das Satz 10.13). Wegen Satz 4.8 genügt zu zeigen $df(x_0) = 0$ für alle $x_0 \in \mathbb{R}^n$. Zum Beweis von $df(x_0) = 0$ sei obdA $x_0 = 0$. Satz 10.9 und Satz 4.32 zeigen dann für den Poissonkern $P(x, \xi)$

$$\frac{\partial}{\partial x_i} f(0) = \frac{1}{R^2} \int_{S(R)} f(\xi) \left(\frac{\partial}{\partial x_i} P(x, \xi) \Big|_{x=0} \right) \sigma_{n-1}(\xi) \, d\xi.$$

Aus $|\text{vol}(S) \partial_i P(x, \xi)|_{x=0}| = n |\xi_i| \|\xi\|^{-n} \leq n R^{1-n}$ folgt $|\text{vol}(S) \partial_i f(0)| \leq R^{-2} \cdot \max(|f|) \cdot n R^{1-n} \text{vol}(S(R)) = \frac{n \text{vol}(S)}{R}$. Also folgt $\partial_i f(0) = 0$ im Limes $R \rightarrow \infty$. $\square$

Das **Coulomb** oder **Newton Potential** $Q(x, y)$ zum (fixierten) Punkt y

$$\boxed{Q(x, y) = \frac{1}{\kappa \cdot \text{vol}(S^{n-1}) \|x - y\|^\kappa}}, \quad \kappa = n - 2 > 0$$

ist lokal eine integrierbare Funktion und definiert daher eine Distribution auf $\mathbb{R}^n$. Aus Korollar 7.6 und Lemma 7.7 folgt für $Q(x) = Q(x, 0)$

Satz 10.15 (Existenz). Für $\rho(x) \in C_c^\infty(\mathbb{R}^n)$ und $\kappa = n - 2 > 0$ ist

$$\boxed{u(x) = (Q * \rho)(x) = \int_{\mathbb{R}^n} \frac{\rho(y)}{(n-2) \text{vol}(S^{n-1}) \|x - y\|^\kappa} \omega_n(y) \, dy}$$

in $C^\infty(\mathbb{R}^n)$ für die Fundamentallösung $Q(x) = \frac{1}{\kappa \cdot \text{vol}(S^{n-1}) \|x\|^\kappa}$, und $u(x)$ definiert eine Lösung der Potential Gleichung $-\Delta u = \rho$.

Mit der so erhaltenen Formel lässt sich die gefundene Lösung der Poisson Gleichung deuten als Verschmierung des **Coulomb** oder **Newton Potentials** $Q(x, y)$ im Punkt y .

Bemerkung. In der **Elektrostatik** reduziert sich das Vektorpotential A auf dem $\mathbb{R}^4$ (siehe Abschnitt 5.16) auf die 1-Form $A = -u(x, y, z)dt$ und der Vierer-Strom j reduziert sich auf die 1-Form $j = -\rho(x, y, z)dt$. Die Maxwellgleichungen vereinfachen sich daher zu der Gleichung

$$-\Delta u(x, y, z) = \rho(x, y, z)$$

für die statische Ladungsverteilung $\rho = \rho(x, y, z)$. Das Lösungspotential u liefert Satz 10.15. Durch das Postulat, daß im Unendlichen die Lösung $u(x, y, z)$ abklingt oder zumindestens beschränkt ist, wird die Lösung u (bis auf eine Konstante) eindeutig.

Bemerkung. Sei X eine Kugelschale X um einen Punkt x_0 und sei $\text{obd}A \ x_0 = 0$. Ist die Ladungsdichte $\rho(x, y, z)$ Null auf X (eine physikalisch typische Situation wenn alle Ladungen im Inneren konzentriert sind), dann ist $u(x, y, z)$ auf X eine harmonische Funktion und kann dort nach Satz 10.12 in eine Potenzreihe entwickelt werden. Nimmt man sogar an $X = \{x \in \mathbb{R}^n \mid \|x\| \geq r\}$, d.h. ist also $R = +\infty$, und ist $u(x, y, z)$ ausserdem beschränkt auf X , dann folgt aus Satz 10.13 die harmonische Fortsetzbarkeit der Kelvin Transformaten $u^*(x, y, z)$ von $\{x \in \mathbb{R}^n \mid 0 < \|x\| < \frac{1}{r}\}$ auf die offene Kugel $\{x \in \mathbb{R}^n \mid \|x\| < \frac{1}{r}\}$. Damit lässt sich $u(x, y, z)$ im Bereich $\|x\| \geq r + \varepsilon$ (für jedes $\varepsilon > 0$) in eine absolut und gleichmässig konvergente Reihe der Gestalt

$$u(x, y, z) = \sum_{l=0}^{\infty} \sum_{k=-l}^{k=+l} b_{kl} \cdot P_{k,l}^*(x, y, z)$$

entwickeln für gewisse Koeffizienten b_{kl} . Aus der Kenntnis von u auf X , d.h. aus dieser Entwicklung, kann man die Ladungsdichte ρ ausserhalb von X aber nicht rekonstruieren! Im Gegenteil: Für den Beobachter in X erscheint es auf Grund dieser Formel eher so, als sei alle Ladung in infinitesimaler Nähe⁴ des Ursprungs $x_0 = 0$ konzentriert!

Bemerkung 10.16. Im Gegensatz zur Potentialgleichung besitzt die d'Alembertgleichung

$$\square A = 0$$

viele beschränkte Lösungen, etwa $A = f(\sum_{\nu=1}^4 a_{\nu} x_{\nu})$ für $a = (a_1, a_2, a_3, a_4)$ mit $q_L(a) = 0$ und beliebige beschränkte Funktionen $f \in C^2(\mathbb{R})$. Die einzige Lösung im Schwartzraum $\mathcal{S}(\mathbb{R}^4)$ ist dagegen $A = 0$, denn die Fourier Transformierte $\mathcal{F}(A) \in \mathcal{S}(\mathbb{R}^4)$ hat dann ihren Träger im Lichtkegel (dies folgt aus Lemma 8.15). Da der Lichtkegel eine Nullmenge im $\mathbb{R}^4$ ist, folgt daraus $\|A\|^2 = \|\mathcal{F}(A)\|^2 = 0$ und damit $A = 0$.

⁴Physiker nennen dies die **Multipolentwicklung** des Potentials in x_0 . Die Terme $P_{k,l}^*(x, y, z)$ können als Potentiale von 2^l infinitesimal in der Nähe von x_0 konzentrierten Ladungen gedeutet werden. Im Beispiel $l = 1$ betrachten wir zwei Elementarladungen entgegen gesetzten Vorzeichens im Punkt $\frac{1}{2}\xi$ und $-\frac{1}{2}\xi$ (**Dipol**). Im Limes $t \rightarrow 0$ für $\xi = t \cdot \xi_0$ gilt

$$\frac{1}{\|\xi\|} \left(\frac{1}{\|x + \frac{1}{2}\xi\|} - \frac{1}{\|x - \frac{1}{2}\xi\|} \right) \longrightarrow \frac{(\xi_0/\|\xi_0\|, x)}{\|x\|^3} = P^*(x)$$

für das lineare harmonische Polynom $P(x) = (\xi_0/\|\xi_0\|, x)$.

11 Ausgewählte Themen II

11.1 Kugelvolumina

Sei $E^n \subseteq \mathbb{R}^n$ die Einheitskugel, d.h. $x \in E^n \Leftrightarrow \|x\| \leq 1$.

Polarkoordinaten. Für die Polarkoordinaten Abbildung $(0, 1) \times \prod_{i=1}^{n-2} [0, \pi) \times [0, 2\pi) \rightarrow E^n$

$$(x_1, \dots, x_n) = r \cdot \left(\cos(\alpha_1), \cos(\alpha_2) \sin(\alpha_1), \dots, \cos(\alpha_{n-1}) \prod_{i=1}^{n-2} \sin(\alpha_i), \prod_{i=1}^{n-1} \sin(\alpha_i) \right).$$

hat die Jacobideterminante die Gestalt $r^{n-1} \cdot g(\alpha_1, \dots, \alpha_{n-1})$ für eine Funktion $g(\alpha_1, \dots, \alpha_{n-1})$ in den **Eulerwinkeln** $\alpha_1, \dots, \alpha_{n-1}$. Setze $\text{vol}(S^{n-1}) := \int_0^\pi \cdots \int_0^{2\pi} g(\alpha_1, \dots, \alpha_{n-1}) d\alpha_1 \cdots d\alpha_{n-1}$.

Durch Berechnung in Polarkoordinaten (Satz 4.27) oder mit Hilfe von Lemma 9.8 folgt

$$\boxed{\text{vol}(\{x \in \mathbb{R}^n \mid \|x\| \leq R\}) = \text{vol}(E^n) \cdot R^n}$$

für das Volumen $\text{vol}(E^n)$ der Einheitskugel, sowie

Lemma 11.1.

$$\boxed{\text{vol}(S^{n-1}) = \int_{S^{n-1}} \sigma_{n-1} = n \cdot \text{vol}(E^n)}.$$

Um das Volumen $\text{vol}(S^{n-1})$ der Einheitssphäre und damit $\text{vol}(E^n)$ zu bestimmen berechnen wir folgendes Integral einmal mit Hilfe von Lemma 8.15.4

$$I = \int_{\mathbb{R}^n} \exp(-r^2) dx_1 \dots dx_n = \prod_{i=1}^n \int_{\mathbb{R}} \exp(-x^2) dx = \pi^{\frac{n}{2}}$$

und einmal in Polarkoordinaten (oder mit Hilfe von Lemma 9.8):

$$I = \text{vol}(S^{n-1}) \int_{r>0} r^{n-1} \exp(-r^2) dr = \frac{\text{vol}(S_{n-1})}{2} \int_{r>0} r^{\frac{n}{2}} \exp(-r) \frac{dr}{r}.$$

Für die **Gammafunktion** $\Gamma(s) := \int_{r>0} r^s \exp(-r) \frac{dr}{r}$ liefert partielle Integration

$$\boxed{\Gamma(s+1) = s\Gamma(s)} \quad , \quad (s > 0).$$

Der Vergleich berechnet $\text{vol}(S^{n-1})$ in Termen der Γ -Funktion

Lemma 11.2.

$$\boxed{\text{vol}(S^{n-1}) = 2 \frac{\pi^{\frac{n}{2}}}{\Gamma(\frac{n}{2})}}.$$

Aus $\text{vol}(S^0) = \text{vol}(E^1) = 2$ und $\text{vol}(S^1) = 2\text{vol}(E^2) = 2\pi$ folgt daher $\Gamma(1) = 1$ und $\Gamma(\frac{1}{2}) = \pi^{1/2}$. Mittels der Rekursionsgleichung der Gammafunktion kann man damit dann alle weiteren Werte $\Gamma(\frac{n}{2})$ leicht bestimmen. Es ergibt sich

$$\text{vol}(S^{n-1}) = 2, 2\pi, 4\pi, 2\pi^2, \dots$$

in den Fällen $n = 1, 2, 3, 4, \dots$.

Bemerkung. Die Eigenschaft

$$d\sigma_{n-1} = n \cdot \omega_n$$

und die Drehinvarianz $M^*(\sigma_{n-1}) = \sigma_{n-1}$ unter Drehungen M in der orthogonalen Gruppe $SO(n, \mathbb{R})$ (in allen Teilebenen in den i, j -Koordinatenrichtungen für $1 \leq i < j \leq n$) bestimmen¹ die Kugelflächen Form $\sigma = \sigma_{n-1}$ in $A^{n-1}(\mathbb{R}^n)$ eindeutig.

In der Fußnote wird unter anderem gezeigt, daß die der Kugelflächenform σ_{n-1} zugeordnete **Potentialform**

$$\rho = \frac{\sigma_{n-1}}{r^n} \in A^{n-1}(\mathbb{R}^n \setminus \{0\}),$$

welche einen Singularität im Nullpunkt besitzt, eine geschlossene rotationssymmetrische Differentialform auf $U = \mathbb{R}^n \setminus \{0\}$ ist:

$$d\rho = 0.$$

Bei Integration über die Sphäre S^{n-1} liefert die Potentialform ρ dasselbe Integral $nc(n) \neq 0$ wie die Oberflächenform σ_{n-1} , da $r = 1$ ist auf der Einheitskugel. Im Gebiet $U = \mathbb{R}^n \setminus \{0\}$ ist deshalb das Poincaré Lemma nicht erfüllt im Grad $i = n-1$ (das verallgemeinert die Bemerkung auf Seite 73 im Fall $n = 2$)

Lemma 11.3. $\rho \notin d(A^{n-2}(U)).$

Beweis. Angenommen $\rho = d\eta$, dann folgt $nc(n) = \int_{S^{n-1}} \rho = \int_{S^{n-1}} d\eta$ für ein $\eta \in A^{n-2}(U)$. Eine stärkere Version (!) des Satzes von Stokes (Übungsaufgabe; hier ohne Beweis) würde dann folgenden Widerspruch liefern $0 \neq n \cdot c(n) = \int_{S^{n-1}} d\eta = \int_{\partial S^{n-1}} \eta = 0$, letzteres wegen $\partial S^{n-1} = \emptyset$. $\square$

¹Ist $\sigma = \sum f_i(x) * dx_i$ invariant unter allen Drehungen $M = M(\alpha)$ in Richtung der x_1, x_2 -Ebene mit Winkeln α , d.h. gilt $M^*(\sigma) = \sigma$, dann folgt $f_1(x) = \cos(\alpha)f_1(\cos(\alpha)x_1 + \sin(\alpha)x_2, -\sin(\alpha)x_1 + \cos(\alpha)x_2, \dots) - \sin(\alpha)f_2(\cos(\alpha)x_1 + \sin(\alpha)x_2, -\sin(\alpha)x_1 + \cos(\alpha)x_2, \dots)$ für alle α . Wenn man diese Relation nach α ableitet und $\alpha = 0$ setzt, folgt daraus $x_1 f_2(x) = x_2 f_1(x)$. Macht man dies analog für andere Koordinaten-Ebenen, so folgt $f_i(x) = x_i \cdot f(x)$ für eine Funktion $f(x)$, welche auf Sphären konstant ist und daher nur vom Radius r abhängt. Also $\sigma = f(r) \cdot \sigma_{n-1}$. Fordert man noch $d\sigma = n \cdot \omega_n$, so folgt aus $d\sigma = df \wedge \sigma_{n-1} + n f \omega_n$ die inhomogene lineare Differentialgleichung $\sum_{i=1}^n x_i \partial_i f(x) + n f(x) = n$, oder wegen $\partial_i f(r) = \frac{x_i}{r} \frac{d}{dr} f(r)$ dann $r \frac{d}{dr} f(r) + n f = n$. Also $f = 1 + g$ für eine Lösung $g(r) = c \cdot r^{-n}$ der homogenen Gleichung $r \frac{d}{dr} g(r) + n g(r) = 0$. Nur für $c = 0$ ist diese Lösung differenzierbar im Punkt Null.

11.2 Überdeckungskompaktheit

Satz 11.4 (Heine-Borel). Für einen folgenkompakten metrischen Raum (X, d) gilt: Jede Überdeckung $X = \bigcup_{i \in I} U_i$ durch offene Teilmengen $U_i \subseteq X$ besitzt eine Überdeckung durch endlich viele der Mengen U_i . (Die Umkehrung gilt auch).

Beweis. 1) Für jedes $\varepsilon > 0$ besitzt ein folgenkompakter Raum (X, d) eine Überdeckung durch endlich viele offene Kugeln $K_\varepsilon(x_i)$ vom Radius ε . [Anderenfalls gäbe es eine Folge $x_1, x_2, \dots$ in X mit $x_n \notin \bigcup_{i=1}^{n-1} K_\varepsilon(x_i)$. Diese würde wegen $d(x_n, x_m) > \varepsilon$ keine konvergente Teilfolge besitzen. Ein Widerspruch zur Folgenkompaktheit von (X, d)]. Für $\varepsilon = 1/n$ liefert diese Beobachtung endlich viele Kugelmittelpunkte und die Vereinigung all dieser Punkte für $n \in \mathbb{N}$ ist eine abzählbare Teilmenge M von X .

2) M liegt dicht in X . [Sei $x \in X$ beliebig und $U = K_\varepsilon(x)$ eine offene Kugel um x vom Radius $\varepsilon > 0$. Sei $1/n < \varepsilon$, dann existiert ein $\xi \in M$ mit $x \in K_{1/2n}(\xi)$. Aus der Dreiecksungleichung folgt dann $x \in K_{1/2n}(\xi) \subset U$].

3) Sei $X = \bigcup_{i \in I} U_i$ eine beliebige Überdeckung von X durch offene Mengen $U_i \subset X$. Für jedes $x \in X$ gibt es dann ein $i \in I$ mit $x \in U_i$ sowie dann eine Kugel $U = U_x = K_\varepsilon(x)$ in U_i , da U_i offen in X ist. Offensichtlich ist dann auch $X = \bigcup_{x \in X} U_x$ eine Überdeckung von X durch offene Teilmengen. Wegen Schritt 2) gilt dann $X = \bigcup_{n, \xi} K_{1/2n}(\xi)$ für eine gewisse Teilmenge J der $(n, \xi) \in \mathbb{N} \times M$. J ist daher abzählbar! Es gilt also $X = \bigcup_{j \in J} V_j$ für $V_j = K_{1/2n}(\xi)$ derart daß nach Schritt 2) jedes $V_j, j \in J$ in einer der offenen Mengen $U_i, i \in I$ enthalten ist. Lässt sich X durch endlich viele der V_j überdecken, dann erst recht durch endlich viele der U_i . Dies reduziert den Beweis auf den Fall, wo die Indexmenge I abzählbar ist.

4) Wegen des letzten Schrittes sei also $X = \bigcup_{i=1}^{\infty} U_i$ eine abzählbare Vereinigung. Man kann die U_i durch die aufsteigende Folge der offenen Mengen $\bigcup_{j \leq i} U_j$ ersetzen und daher oBdA annehmen $U_1 \subseteq U_2 \subseteq \dots$. Dann bilden die abgeschlossenen Komplemente $A_i = X \setminus U_i$ der U_i eine absteigende Folge $A_1 \supseteq A_2 \supseteq \dots$. Es genügt jetzt $A_n = \emptyset$ für ein geeignetes n , denn dann gilt $X = U_n$. Gäbe es ein solches n nicht, gibt es eine Folge $x_n \in X$ mit $x_n \in A_n$. Da X folgenkompakt ist, kann durch Übergang zu einer Teilfolge oBdA angenommen werden $x_n \rightarrow x$ sei konvergent in (X, d) . Es gilt $x \in U_m$ für ein m . Da U_m offen ist, liegen dann fast alle Folgenglieder der Teilfolge x_n in U_m , also $x_n \notin A_m$ für alle $n \geq n_0$. OBdA ist hierbei $n_0 \geq m$. Dies ist ein Widerspruch wegen $x_n \in A_n \subset A_m$ für alle $n \geq n_0$. $\square$

Metrische Räume können skurrile Eigenschaften besitzen²³

²Eine Menge X mit $d(x, y) = 0$ für $x = y$ und $d(x, y) = 1$ für $x \neq y$ definiert einen metrischen Raum (X, d) . Eine Teilmenge $A \subset X$ ist kompakt in (X, d) genau dann, wenn A endlich ist. Jede Kugel in (X, d) ist kompakt oder gleich X . Ist X nicht abzählbar, dann ist X keine abzählbare Vereinigung von kompakten Teilmengen und besitzt keine abzählbare dichte Teilmenge.

³Sei $X \subset \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 \leq 1\}$ die Teilmenge aller Punkte $P = (x, y)$ für die im Fall $y \neq 0$ die Zahl x/y in $\mathbb{Q}$ ist. Versieht man X mit der Wegemetrik d (Hierbei sei $d(P, Q)$ die Länge des kürzesten Weges zwischen P und Q in X), dann ist keine Kugel in X um $P = (0, 0)$ kompakt. Dennoch ist X eine abzählbare Vereinigung von kompakten Teilmengen und besitzt eine abzählbare dichte Teilmenge. Jede Funktion $f \in C_c(X)$ verschwindet im Punkt $x = 0$.

11.3 Residuensatz

Sei $\varepsilon > 0$ und

$$f : U = \{z \in \mathbb{C} \mid 0 < \|z - z_0\| < \varepsilon\} \longrightarrow \mathbb{C}$$

eine auf U holomorphe Funktion.

Satz 11.5 (Laurent Entwicklung). $f(z)$ ist dann auf jedem kompakten Kreisring in U in eine absolut und gleichmässig konvergente Potenzreihe entwickelbar

$$f(z) = \sum_{l=0}^{\infty} a_l \cdot (z - z_0)^l + \sum_{l=1}^{\infty} b_l \cdot (z - z_0)^{-l}.$$

Man nennt den Koeffizient b_1 das **Residuum** $\text{Res}_{z_0}(f)$ der Funktion f bei z_0 .

Beweis. ObdA $z_0 = 0$. Nach Lemma 5.8 sind Real- und Imaginärteil von $f(z)$ harmonisch. Die harmonischen Polynome in $\mathcal{H}_l(\mathbb{R}^2)$ sind $\text{Re}(z^l)$ und $\text{Im}(z^l)$ und somit folgt aus Satz 10.12 die Existenz einer Entwicklung $f(z) = a_0 + b_0 \cdot \log(|z|) + \sum_{l=1}^{\infty} (a_l z^l + \tilde{a}_l \bar{z}^l) + \sum_{l=1}^{\infty} (b_l z^{-l} + \tilde{b}_l \bar{z}^{-l})$. Nach Annahme gilt $\bar{\partial}_z f = 0$ für $\bar{\partial}_z := \partial_x + i\partial_y$ (siehe Abschnitt 5.2). Gliedweises Ableiten mit $\bar{\partial}_z$ gibt $\frac{b_0}{z} + \sum_{l=0}^{\infty} l \tilde{a}_l \bar{z}^{l-1} - \sum_{l=1}^{\infty} l \tilde{b}_l \bar{z}^{-l-1} = 0$, d.h. $b_0 = \tilde{b}_l = \tilde{a}_l = 0$, $l \neq 0$. $\square$

Bemerkung. Der letzte Satz liefert für $\gamma : [0, 2\pi] \rightarrow \mathbb{C}$ mit $\gamma(t) = z_0 + r \cdot \exp(it)$ und $0 < r < \varepsilon$ die Formel

$$\int_{\gamma} f(z) dz = 2\pi i \cdot \text{Res}_{z_0}(f).$$

Satz 6.12 zeigt nämlich $\int_{\gamma} f(z) dz = \sum_{l=0}^{\infty} a_l \cdot \int_{\gamma} (z - z_0)^l dz + \sum_{l=1}^{\infty} b_l \cdot \int_{\gamma} (z - z_0)^{-l} dz$. Beachte

$$\int_{\gamma} (z - z_0)^{-1} dz = 2\pi i \quad \text{sowie} \quad \int_{\gamma} (z - z_0)^n dz = 0, \quad (n \neq -1)$$

wegen $(z - z_0)^n dz = d \frac{(z - z_0)^{n+1}}{n+1}$ für $n \neq -1$ (für beides siehe Abschnitt 5.2).

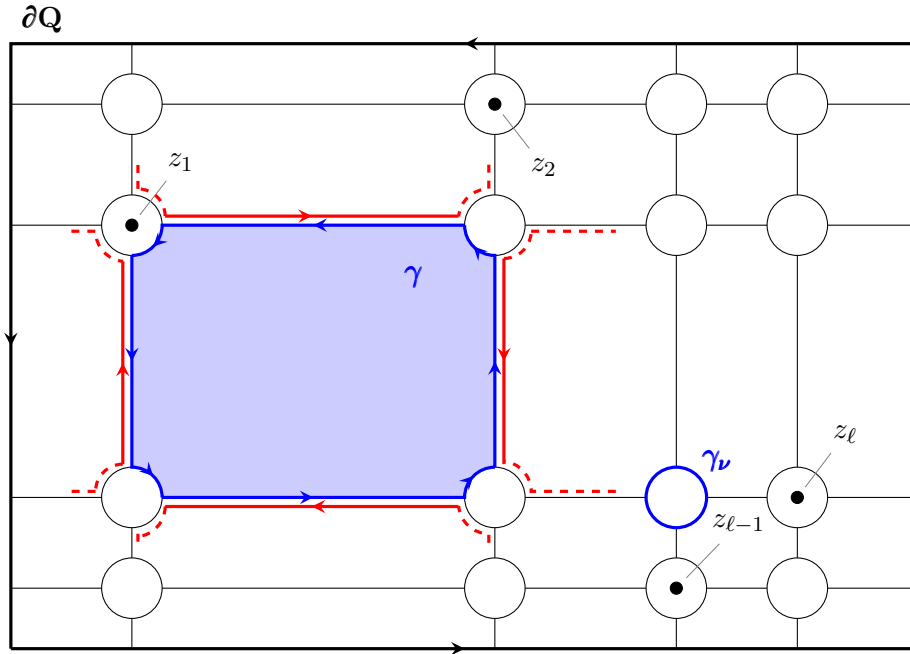
Lemma 11.6. Ist $f : U \rightarrow \mathbb{C}$ eine holomorphe Funktion und ist U eine sternförmige offene Teilmenge von $\mathbb{C}$, dann gilt $\int_{\gamma} f(z) dz = 0$ für jeden geschlossenen Weg γ in U .

Beweis. Nach Lemma 5.4.1 ist die Differentialform $\omega = f(z) dz$ für holomorphes $f(z)$ geschlossen: $d\omega = 0$. Ist der Definitionsbereich U von f offen und sternförmig, folgt aus dem Poincaré Lemma die Existenz eines Potentials ϕ mit $\omega = d\phi$. Ist $\gamma : [a, b] \rightarrow U$ ein Weg in U , gilt also $\int_{\gamma} \omega = \phi(\gamma(b)) - \phi(\gamma(a))$. Ist der Weg γ in U geschlossen, folgt $\int_{\gamma} f(z) dz = 0$. $\square$

Sei nun $Q = [a_1, a_2] \times [b_1, b_2]$ ein beschränkter Quader im $\mathbb{R}^2$. Wir entfernen aus Q offene Kreise vom Radius $\varepsilon > 0$ um die Eckpunkte und erhalten eine abgeschlossene Menge X . Nach Konstruktion gilt $X \subset Q$. Sei U eine offene Teilmenge von $\mathbb{R}^2$ mit $X \subset U$ und sei $f : U \rightarrow \mathbb{C}$ eine holomorphe Funktion. Wurde $\varepsilon > 0$ klein genug gewählt, enthält U eine sternförmige offene Teilmenge $V \subseteq U$ welche ganz X enthält. Wegen Lemma 11.6 folgt daraus $\int_{\partial X} f(z) dz = 0$. Dies zeigt

Satz 11.7 (Residuensatz). Sei $Q \subset \mathbb{R}^2$ ein Quader und $U \subset \mathbb{R}^2$ eine offene Menge, welche Q enthält. Seien $z_1, \dots, z_l$ endlich viele Punkte in $Q \setminus \partial Q$ und sei $f : U \setminus \{z_1, \dots, z_l\} \rightarrow \mathbb{C}$ holomorph. Dann gilt

$$\frac{1}{2\pi i} \int_{\partial Q} f(z) dz = \sum_{\nu=1}^l \text{Res}_{z_\nu}(f).$$



Beweis. Wir zerlegen den Quader Q in endlich viele Teilquader so, daß alle Punkte z_ν für $\nu = 1, \dots, l$ zu Eckpunkten werden. Wir wählen dann $\varepsilon > 0$ klein genug so daß das Innere aller Teilquader nach Herausnahme kleiner Kreisscheiben vom Radius ε um die Eckpunkte sternförmig wird. Seien $\gamma_\nu : [0, 2\pi] \rightarrow U$ kleine Kreiströge in U definiert durch $\gamma_\nu(t) = z_\nu + \varepsilon \cdot \exp(it)$. Dann gilt $\frac{1}{2\pi i} \int_{\partial Q} f(z) dz = \sum_{\nu=1}^l \frac{1}{2\pi i} \int_{\gamma_\nu} f(z) dz$, denn die Differenz beider Seiten schreibt sich als eine endliche Summe von Integralen $\int_\gamma f(z) dz$ über Wege, die in einem sternförmigen offenen Teil von U enthalten sind und daher nach Lemma 11.6 verschwinden. Ist ε klein genug, dann liegt γ_ν in U und f ist holomorph auf $\{z \mid 0 < \|z - z_\nu\| < 2\varepsilon\}$. Daraus folgt $\frac{1}{2\pi i} \int_{\gamma_\nu} f(z) dz = \text{Res}_{z_\nu}(f)$, wie oben gezeigt wurde. Dies zeigt die Behauptung. $\square$

11.4 Wärmeleitungskern

Die Funktion $f_t(x) = t^{-n/2} e^{-\pi \frac{\|x\|^2}{t}}$ ist definiert auf $\mathbb{R}_{>0} \times \mathbb{R}^n$ und stellt dort eine C^∞ -Funktion dar. Man zeigt leicht, daß sie dort eine Lösung der **Wärmeleitungsgleichung**

$$4\pi \partial_t f_t(x) = \Delta f_t(x)$$

ist. Wir studieren im Folgenden das Verhalten der Funktion im rechtsseitigen Limes $t \rightarrow 0^+$ (im Spezialfall $n = 1$). Im physikalischen Kontext bedeutet dies in der Regel: Entweder die Zeit t

geht gegen Null oder die Temperatur oder Energie $T = \frac{1}{t}$ geht gegen unendlich. Man nennt die Funktion $f_t(x - y)$ den **Wärmeleitungskern**.

Lemma 11.8. Für $t > 0$ und $f_t(x) = t^{-1/2} e^{-\pi \frac{x^2}{t}}$ gilt

$$\int_{\mathbb{R}} f_t(x) dx = 1.$$

Beweis. Man reduziert dies mittels Variablensubstitution auf den Spezialfall $t = 1$, der in Lemma 8.15 bewiesen wurde. $\square$

Für $0 < t \leq \delta$ kann das Integral $\int_{|x| \geq \delta} f_t(x) dx$ durch $\int_{|x| \geq 1} \frac{\delta}{t^{1/2}} \exp(-\pi \frac{\delta^2}{t} x^2) dx$, oder damit $\int_1^\infty \frac{\delta}{t^{1/2}} \exp(-\pi \frac{\delta^2}{t} y) dy = \frac{t^{1/2}}{\pi \delta} \cdot \exp(-\pi \frac{\delta^2}{t}) \leq c(\delta) t^{1/2}$ für $c(\delta) = \frac{\exp(-\pi \delta)}{\pi \delta}$ abgeschätzt werden.

Lemma 11.9. Für eine beschränkte stetige Funktion $g(x)$ auf $\mathbb{R}$ gilt⁴

$$g(0) = \lim_{t \rightarrow 0^+} \int_{\mathbb{R}} g(x) f_t(x) dx.$$

Beweis. Indem man $g(x)$ durch $g(x) - g(0)$ ersetzt kann obdA $g(0) = 0$ angenommen werden [benutze $\int_{\mathbb{R}} f_t(x) dx = 1$]. Nach Annahme gilt $|g(x)| \leq C$ für eine Konstante C . Wegen der Stetigkeit von g gibt es für $\varepsilon > 0$ ein $\delta > 0$ mit $|g(x) - g(0)| < \varepsilon$ für $|x| < \delta$. Es gilt $|\int_{|x| \leq \delta} g(x) f_t(x) dx| \leq \varepsilon \int_{|x| \leq \delta} f_t(x) dx$ wegen $f_t(x) \geq 0$. Weiterhin ist $|\int_{|x| \geq \delta} g(x) f_t(x) dx| \leq C \int_{|x| \geq \delta} f_t(x) dx$. Für $t \leq \delta$ ist dies wie oben gezeigt $\leq C c(\delta) \cdot t^{1/2}$ und geht gegen Null im Limes $t \rightarrow 0$. Es folgt

$$-\varepsilon \cdot \lim_{t \rightarrow 0^+} \int_{\mathbb{R}} f_t(x) dx \leq \lim_{t \rightarrow 0^+} \int_{\mathbb{R}} g(x) f_t(x) dx \leq +\varepsilon \cdot \lim_{t \rightarrow 0^+} \int_{\mathbb{R}} f_t(x) dx.$$

Aus $\int_{\mathbb{R}} f_t(x) dx = 1$ (Lemma 11.8) folgt dann sofort die Behauptung. $\square$

11.5 Spinordarstellung

Der K -Vektorraum $\mathcal{S}_n$ der Polynome in den antikommutierenden Variablen $\theta_1, \dots, \theta_n$ mit Koeffizienten in K besitzt folgende ausgezeichnete K -lineare Endomorphismen für $\nu = 1, \dots, n$

$$\theta_\nu: \mathcal{S}_n \longrightarrow \mathcal{S}_n \quad (\text{Linksmultiplikation mit } \theta_\nu) \quad , \quad \frac{\partial}{\partial \theta_\nu}: \mathcal{S}_n \longrightarrow \mathcal{S}_n \quad (\text{Ableiten nach } \theta_\nu)$$

wobei $\partial_\nu = \frac{\partial}{\partial \theta_\nu}$ durch $\partial_\nu(\theta_\nu \cdot g(\theta)) = g(\theta)$ sowie $\partial_\nu \theta^I = 0$, falls $\nu \notin I$, definiert ist. Der von diesen Endomorphismen aufgespannte Untervektorraum V von $\text{End}_K(\mathcal{S}_n)$ hat $\dim_K(V) = 2n$. Beachte $\dim_K(\mathcal{S}_n) = 2^n$ und $\dim_K(\text{End}_K(\mathcal{S}_n)) = (2^n)^2 = 2^{2n}$.

⁴Insbesondere gilt hier im Fall $g(0) \neq 0$ also $g(0) = \lim_{t \rightarrow 0^+} \int_{\mathbb{R}} g(x) f_t(x) dx \neq \int_{\mathbb{R}} \lim_{t \rightarrow 0^+} g(x) f_t(x) dx = 0$. Dies zeigt, wie essentiell die Bedingungen in Satz 6.13 sind, denn $\lim_{t \rightarrow 0^+} f_t(x) = 0$ nach Abschnitt 5.17.

Wie man leicht auf Monomen $f(\theta) = \theta^I$ nachprüft, gilt in $\text{End}_K(\mathcal{S}_n)$, in Analogie zu den Heisenberg Kommutatorrelationen $\frac{\partial}{\partial x_\mu} \circ x_\nu - x_\nu \circ \frac{\partial}{\partial x_\mu} = \delta_{\nu\mu} \cdot \text{id}$, die Formel

$$\boxed{\frac{\partial}{\partial \theta_\mu} \circ \theta_\nu + \theta_\nu \circ \frac{\partial}{\partial \theta_\mu} = \delta_{\nu\mu} \cdot \text{id}}.$$

Neben der Fourier Transformation $\mathcal{F} \in \text{End}_K(\mathcal{S}_n)$ hat man folgende Automorphismen φ von $\mathcal{S}_n$; erstens die *Koordinatenwechsel*⁵ $f(\theta) \mapsto \det(U)^{-1/2} \cdot f(U \cdot \theta)$ für invertierbare komplexe Matrizen $U \in \text{Gl}(n, K)$, zweitens Multiplikationen $f(\theta) \mapsto f_A(\theta) \cdot f(\theta)$ mit **Gaußfunktionen** $f_A(\theta)$ für *schiefsymmetrische*⁶ Matrizen A

$$f_A(\theta) = e^{-\sum_{\nu < \mu} A_{\nu\mu} \theta_\nu \theta_\mu} \quad , \quad A_{\nu\mu} = -A_{\mu\nu} \in K.$$

Es gilt $\partial_\nu f_A(\theta) = (\sum_{\mu < \nu} A_{\nu\mu} \theta_\mu - \sum_{\mu < \nu} A_{\nu\mu} \theta_\mu) \cdot f_A(\theta)$ sowie $f_A(\theta) f(\theta) = f(\theta) f_A(\theta)$ für alle $f \in \mathcal{S}_n$. Beachte auch $f_A(\theta) f_{A'}(\theta) = f_{A+A'}(\theta)$.

Alle genannten Automorphismen $\varphi \in \text{Gl}_K(\mathcal{S}_n)$ haben die Eigenschaft^{7,8,9}

$$\boxed{\varphi^{-1} \circ V \circ \varphi \subseteq V}$$

Für die von diesen Automorphismen φ in $\text{Gl}_K(\mathcal{S}_n)$ erzeugte Untergruppe $G \subset \text{Gl}_K(\mathcal{S}_n)$ existiert daher ein Gruppenhomomorphismus

$$\iota : G \longrightarrow \text{Gl}_K(V) \quad , \quad \iota(\varphi)v := \varphi^{-1} \circ v \circ \varphi \text{ für } v \in V.$$

Der Kern $\text{Kern}(\iota)$ besteht aus skalaren Vielfachen der Identität wegen

Lemma 11.10. *Sei $\varphi \in \text{Gl}_K(\mathcal{S}_n)$. Gilt $\varphi \circ v = v \circ \varphi$ für alle $v \in V$, so ist $\varphi = \lambda \cdot \text{id}$ ein skalares Vielfaches der Identität.*

Beweis. Für $P := \varphi(1) \in \mathcal{S}_n$ gilt $v(P) = \varphi(v(1)) = 0$ für $v = \partial_\nu \in V, \forall \nu$. Also $P(\theta) = \lambda$ für $\lambda \in K$ (benutze Induktion nach dem Grad). Erneut durch Induktion nach $r = |I|$ folgt dann $\varphi(\theta^I) = \lambda \cdot \theta^I$ für alle I . Benutze dazu $\varphi(\theta_\nu \theta^I) = \theta_\nu \varphi(\theta^I)$ für $\theta_\nu \in V$. $\square$

In geeigneten Koordinaten von V sind die $2n \times 2n$ -Matrizen $\iota(\varphi)$ gegeben durch

$$M(U) := \begin{pmatrix} U & 0 \\ 0 & {}^t U^{-1} \end{pmatrix} \quad , \quad M(A) := \begin{pmatrix} E & A \\ 0 & E \end{pmatrix} \quad , \quad S = \begin{pmatrix} 0 & E \\ E & 0 \end{pmatrix}$$

⁵d.h. wir ersetzen die Variablen θ_ν durch $\sum_{\mu=1}^n U_{\mu\nu} \theta_\mu$ und multiplizieren das entstehende Polynom mit der Konstante $\det(U)^{-1/2}$ (die Wurzel existiere in K^* ; diese ist nur eindeutig bis auf ein Vorzeichen!)

⁶ $\sum_{\nu < \mu} A_{\nu\mu} \theta_\nu \theta_\mu = \frac{1}{2} \sum_{\nu, \mu} A_{\nu\mu} \theta_\nu \theta_\mu$ definiert eine symmetrische Bilinearform $(\eta, \theta)_A = \sum_{\nu, \mu} A_{\nu\mu} \eta_\nu \theta_\mu$, d.h. es gilt $(\eta, \theta)_A = (\theta, \eta)_A$.

⁷ $f_{-A}(\theta) \theta_\nu f_A(\theta) f(\theta) = f(\theta)$ und $f_{-A}(\theta) \partial_\nu f_A(\theta) f(\theta) = (\partial_\nu + (\sum_{\mu < \nu} A_{\nu\mu} \theta_\mu - \sum_{\mu < \nu} A_{\nu\mu} \theta_\mu)) f(\theta)$.

⁸ $\mathcal{F}^{-1}(\partial_\nu \mathcal{F}(f)) = \theta_\nu f$ sowie damit $\mathcal{F}^{-1}(\theta_\nu \mathcal{F}(f)) = \partial_\nu f$ zeigt man vollkommen analog zu Lemma 8.15. Zeige zuerst $\partial_\nu \mathcal{F}(f) = \mathcal{F}(\theta_\nu f)$ mit Hilfe von $\frac{\partial}{\partial \eta_\nu} \int P(\theta, \eta) d\theta = \int \frac{\partial}{\partial \eta_\nu} P(\theta, \eta) d\theta$. Analog $\mathcal{F}_L^{-1} \circ \theta_\nu \circ \mathcal{F}_L = \lambda_\nu^{-1} \partial_\nu$.

⁹Der Fall von Koordinatenwechseln ist klar.

und liegen daher in der orthogonalen Gruppe $O(S, K)$ der symmetrischen $2n \times 2n$ -Matrix S . Die Matrix S liegt in $SO(S, K)$ genau dann, wenn n gerade ist; die anderen liegen in $SO(S, K)$. Wir überlassen es dem Leser folgendes nachzuweisen¹⁰

Lemma 11.11. Für $K = \mathbb{C}$ ist die von den Matrizen $M(U), M(A), SM(A)S^{-1}$ erzeugte Untergruppe $H \subset Gl_{\mathbb{C}}(V)$ die Gruppe $H = SO(S, \mathbb{C})$. Also ist $G/Kern(\iota) = SO(S, \mathbb{C})$ oder $O(S, \mathbb{C})$ je nachdem ob $\dim_{\mathbb{C}}(V)$ durch 4 teilbar (d.h. n gerade) ist oder nicht.

Lemma 11.12. $Kern(\iota) = \pm id_{S_n}$.

Beweis. $[f, h] := \int f^*(\theta) \cdot h(\theta) d\theta$ ist eine G -invariante K -Bilinearform auf S_n . Benutze dazu $f_A^*(\theta) = f_{-A}(\theta)$, die ‘Leibnizformel’ $\int h(TU \cdot \theta) d\theta = \det(U) \cdot \int h(\theta) d\theta$ (Lemma 5.27) und¹¹

$$[\mathcal{F}(f), \mathcal{F}(h)] = [f, h] \quad , \quad f, h \in S_n .$$

Für $g = \lambda \cdot id \in G$ gilt deshalb $[\lambda \cdot f, \lambda \cdot h] = [f, h]$, also $\lambda^2 = 1$. $\square$

Das Urbild von $SO(S, \mathbb{C})$ in G definiert wegen Lemma 11.12 eine zweiblättrige Überlagerung der Gruppe $SO(S, \mathbb{C})$, die sogenannte **Spingruppe** $Spin(S, \mathbb{C}) \subseteq G$. Die kanonische Einbettung von G in $Gl_{\mathbb{C}}(S_n)$ definiert eine Darstellung

$$Spin(S, \mathbb{C}) \hookrightarrow Gl_{\mathbb{C}}(S_n) .$$

Die Operation von $Spin(S, \mathbb{C})$ erhält die Teilräume $S_n^{\pm}$ der geraden resp. ungeraden Polynome in S_n . Unsere Darstellung zerfällt daher in zwei Teildarstellungen $S_n^{\pm}$ der Dimension $\dim_{\mathbb{C}}(S_n^{\pm}) = 2^{n-1}$, die beiden **Spindarstellungen** der Gruppe $Spin(S, \mathbb{C})$

$$Spin(S, \mathbb{C}) \longrightarrow Gl_{\mathbb{C}}(S_n^{\pm}) .$$

Physiker Notation: $(S_n)_L$ und $(S_n)_R$ für ‘links’ und ‘rechts’ anstatt S_n^+ und S_n^- .

11.6 Oszillatordarstellung

Auf dem Schwartzraum $\mathcal{S}(\mathbb{R}^N)$ hat man $\mathbb{C}$ -lineare Endomorphismen für $\nu = 1, \dots, N$

$$2\pi i x_{\nu} : \mathcal{S}(\mathbb{R}^N) \longrightarrow \mathcal{S}(\mathbb{R}^N) \quad (\text{Multiplikation mit } 2\pi i x_{\nu}) \quad , \quad \frac{\partial}{\partial x_{\nu}} : \mathcal{S}(\mathbb{R}^N) \longrightarrow \mathcal{S}(\mathbb{R}^N) .$$

¹⁰Hinweis: Der Hyperbelfall $n = 1$ ist instruktiv, denn hier ist

$$O(S, K) = \left\{ \begin{pmatrix} a & 0 \\ 0 & a^{-1} \end{pmatrix} \right\} \cup \left\{ \begin{pmatrix} 0 & a \\ a^{-1} & 0 \end{pmatrix} \right\} .$$

Für $n \geq 2$ und $M \in SO(S, K)$ setzt man $v = M(e_1), w = M(e_{n+1})$ und zeigt, daß wegen $q_S(v) = q_S(w) = 0$ und ${}^T v S w = 1$ eine Substitution $g \in H$ existiert mit $g(v) = e_1$ und $g(w) = e_{n+1}$. Dann läßt gM den Teilraum $W = K \cdot e_1 + K \cdot e_{n+1}$ identisch fest. Man kann sich auf das Orthokomplement $W^{\perp} = \{v \in V \mid {}^T v S w = 0 \forall w \in W\}$ beschränken und schließt dann per Induktion.

¹¹ S_n zerfällt in orthogonale Summanden $W = Kf \oplus K\mathcal{F}(f)$, $f(\theta) = \theta^I$ mit $[f, f] = [\mathcal{F}(f), \mathcal{F}(f)] = 0$ wegen Lemma 5.28. Somit genügt $[\mathcal{F}(f), \mathcal{F}^2(f)] = [f, \mathcal{F}(f)]$ wegen $[h, g] = \varepsilon_n[g, h]$ [beachte $\int h^* g d\theta = \varepsilon_n \int (h^* g)^* d\theta = \varepsilon_n \int g^* f h d\theta$]. Setze nun $h := \mathcal{F}(f)$ und $g := \mathcal{F}^2(f) = \varepsilon_n f$.

Der davon aufgespannte Untervektorraum V in $\text{End}_{\mathbb{C}}(\mathcal{S}(\mathbb{R}^N))$ hat Dimension $\dim_{\mathbb{C}}(V) = 2N$. Neben der Fourier Transformation $\mathcal{F} \in \text{End}_{\mathbb{C}}(\mathcal{S}(\mathbb{R}^N))$ betrachten wir folgende Automorphismen φ von $\mathcal{S}(\mathbb{R}^N)$: Die *Koordinatenwechsel* $f(x) \mapsto \det(U)^{+1/2} \cdot f(TU \cdot \theta)$ für invertierbare reelle Matrizen $U \in \text{Gl}(n, \mathbb{R})$, sowie Multiplikationen $f(x) \mapsto f_S(x) \cdot f(x)$ mit Gaußfunktionen $f_S(x)$ für *symmetrische* reelle Matrizen S

$$f_S(x) = e^{-\pi i \sum_{\nu, \mu} S_{\nu\mu} x_{\nu} x_{\mu}} \quad , \quad S_{\nu\mu} = S_{\mu\nu} \in \mathbb{R} .$$

Es gilt $\partial_{\nu} f_S(x) = (\sum_{\mu} S_{\nu\mu} x_{\mu}) \cdot f_S(x)$. Die genannten Automorphismen $\varphi \in \text{Gl}_{\mathbb{C}}(\mathcal{S}(\mathbb{R}^N))$ haben wegen Lemma 8.15, der Substitutionsregel und wegen $f_{-S}(x) x_{\nu} f_S(x) f(x) = f(x)$ und $f_{-S}(x) \partial_{\nu} f_S(x) f(x) = (\partial_{\nu} + \sum_{\mu=1}^N S_{i\mu} \cdot 2\pi i x_{\mu}) f(x)$ die Eigenschaft

$$\boxed{\varphi^{-1} \circ V \circ \varphi \subseteq V} .$$

Für die von diesen Automorphismen φ in $\text{Gl}_{\mathbb{C}}(\mathcal{S}(\mathbb{R}^N))$ erzeugte Untergruppe $G \subset \text{Gl}_{\mathbb{C}}(\mathcal{S}(\mathbb{R}^N))$ existiert daher ein Gruppenhomomorphismus

$$\iota : G \longrightarrow \text{Gl}_{\mathbb{C}}(V) \quad , \quad \iota(\varphi)v := \varphi^{-1} \circ v \circ \varphi \text{ für } v \in V .$$

$\text{Kern}(\iota)$ besteht aus skalaren Matrizen. [$f(x) := \varphi(\exp(-\pi r^2))$ ist gleich $\lambda \exp(-\pi r^2)$, da $f(x)$ von den $\partial_{x_{\nu}} - 2\pi x_{\nu}$ annulliert wird. Man schließt dann $\varphi = \lambda \cdot \text{id}$ analog zu Lemma 11.10 auf der dichten Teilmenge $\bigoplus_{l=0}^{\infty} \mathcal{P}_l(\mathbb{R}^N) \cdot \exp(-\pi r^2)$ von $\mathcal{S}(\mathbb{R}^N)$]. In geeigneten Koordinaten von V sind die Matrizen $\iota(\varphi)$ gegeben durch

$$M(U) := \begin{pmatrix} U & 0 \\ 0 & {}^t U^{-1} \end{pmatrix} \quad , \quad M(A) := \begin{pmatrix} E & S \\ 0 & E \end{pmatrix} \quad , \quad J = \begin{pmatrix} 0 & E \\ -E & 0 \end{pmatrix} .$$

Diese liegen in der **symplektischen Gruppe** $Sp(2N, \mathbb{R}) \subset \text{Gl}_{\mathbb{C}}(V)$ aller reellen $2N \times 2N$ Matrizen M mit der Eigenschaft ${}^t M J M = J$. Ähnlich wie Lemma 11.11 zeigt man dann

Lemma 11.13. *Die von den Matrizen $M(U), M(S), J$ erzeugte Untergruppe in $\text{Gl}_{\mathbb{C}}(V)$ ist die symplektische Gruppe $Sp(2N, \mathbb{R})$. Insbesondere ist $G/\text{Kern}(\iota) = Sp(2N, \mathbb{R})$.*

Normiert¹² man die reelle Fourier Transformation $\mathcal{F}$ mit einer geeigneten komplexen Zahl γ vom Betrag 1, gilt $\text{Kern}(\iota) = \pm \text{id}_{\mathcal{S}(\mathbb{R}^N)}$ (der Nachweis dafür ist viel diffiziler als im Fall von Lemma 11.12). Damit wird dann G eine zweiblättrige Überlagerung der symplektischen Gruppen, die sogenannte **metaplektische Gruppe** $Mp(2N, \mathbb{R})$. Die kanonische Einbettung der Gruppe $G = Mp(2N, \mathbb{R})$ in die Gruppe $\text{Gl}_{\mathbb{C}}(\mathcal{S}(\mathbb{R}^N))$ wird realisiert durch eine Einbettung in die unitäre Gruppe von $L^2(\mathbb{R})$

$$\boxed{Mp(2N, \mathbb{R}) \longrightarrow U(L^2(\mathbb{R}^N))} ,$$

wie man durch direkte (!) Inspektion auf den Erzeugern $M(U), M(S), \gamma \cdot \mathcal{F}$ sieht. Dies definiert die **metaplektische Darstellung** oder auch **Oszillatordarstellung**.

¹²In $\text{Gl}_{\mathbb{C}}(V)$ gilt $J^2 = M(U)$ für $U = -E$. Der U -Koordinatenwechsel liefert $\det(U)^{+1/2} f(-x) = i^N f(-x)$. Ersetze $\mathcal{F}$ durch $\gamma \cdot \mathcal{F}$ für $\gamma \in \mathbb{C}^*$ mit $\gamma^2 = i^N$ wegen $\mathcal{F}^2(f)(x) = f(-x)$.

11.7 Spinor-Matrizen

Sei $S \in Gl(K, k)$ eine symmetrische Matrix mit Koeffizienten $S_{\nu\mu}$ im Körper K . Wir nennen Endomorphismen $\gamma_1, \dots, \gamma_k$ eines K -Vektorraums W Spinoren für S , wenn gilt

$$\gamma_\nu \circ \gamma_\mu + \gamma_\mu \circ \gamma_\nu = 2S_{\nu\mu} \cdot id_W.$$

Beispiel ($k = 2n$). Sei S die Matrix von Abschnitt 11.5. Dann sind $\gamma_i = \theta_i, \gamma_{i+n} = \frac{\partial}{\partial \theta_i}$ ($i = 1, \dots, n$) Spinoren für $\frac{1}{2}S$ und $W = \mathcal{S}_n \cong \mathbb{C}^{2^n}$. Die 2^{2n} angeordneten Monome γ^J , $J \subseteq \{1, \dots, 2n\}$ sind linear unabhängig in $M_{2^n, 2^n}(\mathbb{C})$, wie man leicht sieht. Sie bilden daher eine Basis wegen $\dim_{\mathbb{C}}(M_{2^n, 2^n}(\mathbb{C})) = (2^n)^2 = 2^{2n}$.

Basiswechsel. Sind $\gamma_1, \dots, \gamma_k$ Spinoren für S in $End_K(W)$ und ist $U \in Gl(k, K)$, dann sind $\tilde{\gamma}_\mu := \sum_{\nu=1}^k U_{\nu\mu} \gamma_\nu$ Spinoren¹³ für die symmetrische Matrix $\tilde{S} = {}^T U S U$ in $End_K(W)$.

Die Abbildung $O(\tilde{S}, K) \ni M \mapsto U^{-1} M U \in O(S, K)$ definiert einen Isomorphismus¹⁴

$$O(\tilde{S}, K) \cong O(S, K).$$

Die Abbildung $\gamma_\nu \mapsto \tilde{\gamma}_\nu$ setzt sich zu einem Ringisomorphismus ρ des Ringes $End_K(K^{2^n})$ fort und damit zu einem Gruppenautomorphismus $\rho: Gl(2^n, K) \rightarrow Gl(2^n, K)$.

Der Fall $K = \mathbb{C}$. Sei dann $Spin(\tilde{S}, \mathbb{C}) = \rho(Spin(S, \mathbb{C}))$ für die Spinorgruppe $Spin(S, \mathbb{C})$ definiert in Abschnitt 11.5. Beachte: Es gilt $g^{-1} \tilde{V} g \subset \tilde{V}$ für den K -Aufspann $\tilde{V}$ der Spinormatrizen $\tilde{\gamma}$. Die Abbildung $g \mapsto (\tilde{v} \mapsto g^{-1} \tilde{v} g)$ definiert einen Gruppenhomomorphismus (2-blättrige Überlagerung)

$$Spin(\tilde{S}, \mathbb{C}) \longrightarrow SO(\tilde{S}, \mathbb{C}).$$

Die Einbettung $\tilde{G} \hookrightarrow Gl(2^n, \mathbb{C})$ definiert die Spindarstellungen.

Da zwei symmetrische invertierbare Matrizen $\tilde{S}, S$ lassen sich über dem Körper $K = \mathbb{C}$ durch einen Basiswechsel ineinander überführen, existiert $U \in Gl(2n, \mathbb{C})$ mit der Eigenschaft $U^T S U = \tilde{S}$. Daraus folgt

Lemma 11.14. Für jedes $S = {}^T S \in Gl(2n, \mathbb{C})$ existieren Spinoren $\gamma_1, \dots, \gamma_{2n}$ in $End(\mathbb{C}^{2^n})$.

Beispiel. Im Fall $2n = 4$ definieren die Matrizen¹⁵

$$\gamma_\nu = \begin{pmatrix} 0 & \sigma_\nu \\ -\sigma_\nu & 0 \end{pmatrix} \quad \text{für } \nu = 1, 2, 3 \quad \text{sowie} \quad \gamma_4 = \begin{pmatrix} 0 & \sigma_4 \\ \sigma_4 & 0 \end{pmatrix}$$

Spinoren für die Lorentzform $\tilde{S} = \text{diag}(-1, -1, -1, 1)$. Hierbei sind $\sigma_1, \sigma_2, \sigma_3$ die **Pauli-Matrizen**. Die beiden Spinordarstellungen $\mathcal{S}^\pm$ haben die Dimension $2^{n-1} = 2$.

¹³Für $\lambda_1, \dots, \lambda_k \in K$ gilt $\sum_n \sum_m \lambda_n \lambda_m (\tilde{\gamma}_n \tilde{\gamma}_m + \tilde{\gamma}_m \tilde{\gamma}_n) = \sum_{\nu, \mu, n, m} \lambda_n \lambda_m U_{\nu n} U_{\mu m} (\gamma_\nu \gamma_\mu + \gamma_\mu \gamma_\nu) = 2 \sum_{\nu, \mu, n, m} \lambda_n \lambda_m U_{\nu n} U_{\mu m} S_{\nu\mu} = 2 \sum_n \sum_m \lambda_n \lambda_m \tilde{S}_{nm}$. Setze jetzt $\lambda_i = 0$ für $i \neq \nu, \mu$.

¹⁴ $M \in O(\tilde{S}, K) \iff {}^T M \tilde{S} M = \tilde{S} \iff {}^T M^T U S U M = U^T \tilde{S} U \iff U M U^{-1} \in O(S, K)$.

¹⁵ $\theta_1 + \frac{\partial}{\partial \theta_1}, \dots, \theta_n + \frac{\partial}{\partial \theta_n}, \theta_1 - \frac{\partial}{\partial \theta_1}, \dots, \theta_n - \frac{\partial}{\partial \theta_n}$ definieren S -Spinoren für die Matrix $S = \text{diag}(1, \dots, 1, -1, \dots, -1)$. Damit konstruiert man sofort ganz explizit $\tilde{S}$ -Spinoren für beliebige Diagonalmatrizen S .

Die Pauli-Matrizen sind wie folgt erklärt: Der **Minkowski Raum** $(\mathbb{R}^4, q_L)$ ist der $\mathbb{R}^4$ versehen mit der quadratischen Lorentz Form

$$q_L(x_1, x_2, x_3, x_4) = -x_1^2 - x_2^2 - x_3^2 + x_4^2.$$

Die zugehörige orthogonale Gruppe $O(3, 1)$ ist die **Lorentzgruppe**. Der Minkowski Raum kann mit dem $\mathbb{R}$ -Vektorraum $Herm$ der komplex hermiteschen 2×2 Matrizen identifiziert werden

$$\mathbb{R}^4 \ni x = (x_1, x_2, x_3, x_4) \mapsto \sum_{i=1}^4 x_i \sigma_i = \begin{pmatrix} x_4 + x_1 & x_2 + ix_3 \\ x_2 - ix_3 & x_4 - x_1 \end{pmatrix} = \not{x} \in Herm$$

mit $x_4 = \frac{1}{2} Tr(\not{x})$. Matrixen $M \in Sl(2, \mathbb{C})$ operieren auf $Herm$ vermöge $H \mapsto M^\dagger H M$ durch $\mathbb{R}$ -lineare Transformationen $\psi(M)$. Wegen $det(M^\dagger H M) = det(H)$ und $q_L(x) = det(\not{x})$ definiert dies einen Homomorphismus $M \mapsto \psi(M)$

$$\boxed{\psi : Sl(2, \mathbb{C}) \longrightarrow O(3, 1)(\mathbb{R})}.$$

Man kann zeigen $Kern(\psi) = \pm id_{2,2}$, und für $\varepsilon = diag(1, 1, 1, -1)$ gilt

$$\boxed{O(3, 1)(\mathbb{R}) = Bild(\psi) \cup -Bild(\psi) \cup \varepsilon \cdot Bild(\psi) \cup -\varepsilon \cdot Bild(\psi)}.$$

11.8 Heisenberggruppe

Nach Satz 8.18 ist die Fourier Transformation unitär, definiert also einen Automorphismus des Zustandsraumes $L^2(\mathbb{R})$. Lemma 8.15 2) und 3) lässt sich so deuten, daß die Fourier Transformation die physikalischen Operatoren von Impuls $\frac{1}{2\pi i} X = \frac{1}{2\pi i} \frac{d}{dx}$ und Ort $\frac{1}{2\pi i} Y = x$ vertauscht; also als Transformation von der (kohärenten) **Ortsraumdarstellung** in die (kohärente) **Impulsraumdarstellung**. Das Wirkungsquantum $\hbar$ wurde hierzu der Einfachheit halber zu 1 normiert. Man hat die unitären Transformationen

$$U_t(f): f(x) \mapsto f(x+t) \quad , \quad V_s(f): f(x) \mapsto e^{2\pi i x s} \cdot f(x)$$

$$W_r(f): f(x) \mapsto e^{2\pi i r} \cdot f(x)$$

des Hilbertraumes $L^2(\mathbb{R})$ in sich, welche man (hier nur symbolisch¹⁶) auch in der Form $U_t = \exp(2\pi i t \tilde{X})$ und $V_s = \exp(2\pi i s \tilde{Y})$ schreibt, denn es gelten die Funktionalgleichungen $U_t \circ U_{t'} = U_{t+t'}$ und $V_s \circ V_{s'} = V_{s+s'}$ mit

$$\frac{d}{dt} U_t(f)|_{t=0}(x) = \frac{d}{dt} f(x+t)|_{t=0} = \partial_x f(x) = (Xf)(x)$$

$$\frac{d}{ds} V_s(f)|_{s=0}(x) = 2\pi i x \cdot f(x) = (Yf)(x) ,$$

¹⁶Man würde gerne schreiben $U_t : f(x) \mapsto e^{\partial_x t} f(x)$, denn $f(x+t) = e^{\partial_x t} f(x)$ gilt für analytische Funktionen f und kleine t ; siehe Lemma 4.39. Für beliebige Funktionen $f \in L^2(X)$ ist der Ausdruck $e^{\partial_x t} f(x)$ aber sinnlos!

$$\frac{d}{dr} W_r(f)|_{r=0}(x) = 2\pi i \cdot f(x) = 2\pi i \cdot (id_{L^2(\mathbb{R})} f)(x).$$

Die Operatoren U_t, V_s und W_r (für $r, s, t \in \mathbb{R}$) erzeugen eine Gruppe unitärer Operatoren, die sogenannte **Heisenberggruppe**. Es gilt

$$U_t \circ V_s = W_{st} \circ V_s \circ U_t = e^{2\pi i st} \cdot V_s \circ U_t.$$

Bis auf einen **Phasenfaktor** $W_{st} = \exp(2\pi i st)$ vertauschen also U_t und V_s . Die Operatoren W_s induzieren die identische Abbildung auf dem Zustandsraum aller Geraden $\mathbb{C} \cdot v$ im Hilbertraum. In der Tat bildet W_s die Gerade $\mathbb{C} \cdot v$ auf $\mathbb{C} \cdot \exp(2\pi i r) \cdot v = \mathbb{C} \cdot v$ ab.

Die Heisenberggruppe ist isomorph zur Matrixgruppe aller Matrizen der Gestalt

$$\left\{ \begin{pmatrix} 1 & t & r \\ 0 & 1 & s \\ 0 & 0 & 1 \end{pmatrix} \right\}$$

beschrieben werden, mit

$$u_t = \begin{pmatrix} 1 & t & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad v_s = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & s \\ 0 & 0 & 1 \end{pmatrix} \quad w_r = \begin{pmatrix} 1 & 0 & r \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Diese Matrixgruppe operiert auf $L^2(\mathbb{R})$ durch die Zuordnung $u_t \mapsto U_t, v_s \mapsto V_s, w_r \mapsto W_r$, denn man prüft leicht nach

$$u_t \circ v_s = w_{st} \circ v_s \circ u_t.$$

11.9 Vertauschungslemma

Für einen nicht degenerierten kompakten Quader $Q \subset \mathbb{R}^n$ sei $f_n(x)$ eine Folge nach $\frac{\partial}{\partial x_\nu}$ differenzierbarer Funktionen mit stetigen Ableitungen $g_n(x) := \frac{\partial f_n(x)}{\partial x_\nu} \in C(Q)$. Dann ist $f_n(x) \in C(Q)$.

Lemma 11.15. *Konvergieren $f_n(x) \rightarrow f(x)$ und $\frac{\partial f_n(x)}{\partial x_\nu} \rightarrow g(x)$ gleichmäßig (d.h. in $C(Q)$), dann ist die Grenzfunktion $f(x)$ nach $\frac{\partial}{\partial x_\nu}$ differenzierbar und es gilt $\frac{\partial f(x)}{\partial x_\nu} = g(x)$.*

Beweis. Die Folge $G_n(x) = \int_{a_\nu}^{b_\nu} g_n(x) dx_\nu$ ist eine Cauchyfolge in $C(Q)$ (Boxungleichung) und konvergiert daher in $C(Q)$ gegen eine stetige Grenzfunktion $G(x)$ (Satz 2.24). Nach Satz 4.14 gilt $\frac{\partial G_n(x)}{\partial x_\nu} = g_n(x)$, also $\frac{\partial}{\partial x_\nu}(G_n(x) - f_n(x)) = 0$. Somit ist $G_n(x) - f_n(x)$ eine stetige Funktion $F_n(x)$, welche nicht von x_ν abhängt. Im Limes $n \rightarrow \infty$ hängt auch $G(x) - f(x) = \lim_n F_n(x)$ nicht von x_ν ab. Es folgt $\frac{\partial}{\partial x_\nu}(G(x) - f(x)) = 0$, und damit $\frac{\partial f(x)}{\partial x_\nu} = \frac{\partial G(x)}{\partial x_\nu} = g(x)$, letzteres nach Satz 4.32. $\square$

Betrachte $f : [c, d] \rightarrow \mathbb{R}^N$ für $a : [c, d] \times \mathbb{R}^N \rightarrow \mathbb{R}^N$. Wir nehmen an $a = (a_1, \dots, a_n)$ mit $a_\nu \in C_c^\infty([c, d] \times \mathbb{R}^N)$. Dann besitzt die Differentialgleichung $f'(t) = a(t, f(t))$ eine eindeutige

Lösung $\varphi_t(x)$ für gegebenes $f(t_0) = x \in \mathbb{R}^n$ nach ?? . Angenommen, die Lösung $f(t, x)$ ist unendlich oft partiell differenzierbar in (t, x) . Dann folgt aus der Kettenregel und der Symmetrie der Hessematrix

$$\frac{d}{dt}f = a(t, f) \quad , \quad \frac{d}{dt}(\partial_\nu f) = \sum_\mu \frac{\partial a(t, f)}{\partial y_\mu} \frac{\partial f_\mu}{\partial x_\nu}$$

Diese Differentialgleichung für $F = (f, \partial_1 f, \dots, \partial_n f)$ bestimmt F eindeutig durch die Anfangswerte (WELCHE). Benutzt man die Picard Iteration für f und F , folgt aus dem letzten Lemma die partielle Differenzierbarkeit von $f(t, x)$ (sukzessive in allen Ordnungen).