

Vorlesungsskript zur Analysis

Emil Wiedemann
Universität Ulm, 2018/19

Vorbemerkungen

Dies ist ein Skript zur Vorlesung Analysis I und II im akademischen Jahr 2018/19 an der Universität Ulm. Diese Vorlesung vermittelt nicht nur eine Einführung in die Analysis als Teildisziplin der Mathematik, sondern auch in die Mathematik überhaupt. Insbesondere das erste Kapitel vermittelt Grundlagen der Mathematik, die nicht für die Analysis spezifisch sind.

Die Analysis befaßt sich, grob gesprochen, mit Grenzwerten von Folgen und Funktionen und daran anknüpfend mit Ableitungen und Integralen von Funktionen. Sie ist im Vergleich zu den beiden anderen großen klassischen Teilgebieten der Mathematik (Algebra und Geometrie) relativ jung: Ihr Anfang läßt sich auf die gleichzeitige Entwicklung der Differential- und Integralrechnung durch NEWTON und LEIBNIZ in der zweiten Hälfte des 17. Jahrhunderts terminieren. Gleichwohl gehen Grundideen der Integralrechnung bereits auf die griechische Antike zurück. Die mathematisch rigorose Fundierung der zentralen Konzepte der Analysis in der Form, wie wir sie im Laufe dieser Vorlesung kennenlernen werden, stammt aus dem 19. Jahrhundert und ist mit den Namen CAUCHY, WEIERSTRASS, RIEMANN und vielen anderen verbunden. Im 20. Jahrhundert dienten die Werkzeuge und Konzepte der Analysis maßgeblich zur Entwicklung neuerer Zweige der (angewandten) Mathematik wie der Stochastik, Numerik oder Finanzmathematik. Bei alledem sollte man die enge Wechselwirkung zwischen verschiedenen Bereichen der Mathematik betonen: So kann man etwa mit analytischen Methoden an Probleme der Zahlentheorie oder der Geometrie herangehen, andererseits zeigen sich z.B. in der Funktionalanalysis die Vorzüge der Verwendung von Begriffen der Linearen Algebra innerhalb der Analysis. Eine strikte und eindeutige Abgrenzung der mathematischen Teilgebiete voneinander ist daher weder sinnvoll noch möglich.

Als Begleitlektüre empfehle ich die Analysis-Lehrbücher von FORSTER und von AMANN-ESCHER, die für weite Teile dieser Vorlesung als Grundlage dienen. Für die ‚mathematische Allgemeinbildung‘ und zur Erweiterung Ihres Horizonts lohnt ein Blick in die weiteren im Literaturverzeichnis angegebenen Bücher. Darüber hinaus empfehle ich die Lektüre des Eintrags ‚Philosophy of Mathematics‘ in der auch sonst sehr lesenswerten *Stanford Encyclopedia of Philosophy* (<https://plato.stanford.edu/>). Das vorliegende Manuskript ist allerdings ‚self-contained‘, also ohne Zuhilfenahme weiterer Literatur lesbar, und wird für den Prüfungserfolg ausreichend sein.

Beachten Sie bitte die folgenden Warnhinweise:

- Das Skript wird Ihnen als Lernhilfe zur Verfügung gestellt, gibt aber nicht notwendigerweise die gesamten prüfungsrelevanten Inhalte wieder. *Prüfungsrelevant ist genau der in der Vorlesung besprochene Stoff.*
- Es ist davon auszugehen, daß sich (zahlreiche) Fehler in das Skript einschleichen werden. Seien Sie also stets kritisch und senden Sie mir Hinweise auf vermutete Fehler an meine E-Mail-Adresse emil.wiedemann@uni-ulm.de. Besonders fleißigen Einsendenden winken attraktive Prämien.

Ich bedanke mich herzlich bei Dr. Manfred Sauter und Frederic Weber für ihre hervorragende Leitung des Übungsbetriebs und zahlreiche Hinweise, die zur Verbesserung dieses Manuskripts geführt haben.

Inhaltsverzeichnis

Vorbemerkungen	2
Kapitel 1. Grundlagen der Mathematik	5
1.1. Logik	5
1.2. Mengenlehre	11
1.3. Rationale Zahlen	20
Kapitel 2. Folgen und Vollständigkeit	29
2.1. Konvergenz	29
2.2. Reelle Zahlen	32
2.3. Vollständigkeit	35
2.4. Teilfolgen und Häufungspunkte	37
Kapitel 3. Stetige Funktionen	40
3.1. Charakterisierungen von Stetigkeit	40
3.2. Eigenschaften stetiger Funktionen	44
Kapitel 4. Differentiation und Integration	49
4.1. Ableitungen	49
4.2. Monotonie und Konvexität	57
4.3. Das Integral stetiger Funktionen	61
4.4. Der Hauptsatz der Differential- und Integralrechnung	66
4.5. Uneigentliche Integrale	69
4.6. Der Satz von TAYLOR	70
Kapitel 5. Potenzreihen	73
5.1. Komplexe Zahlen	73
5.2. Reihen	76
5.3. Gleichmäßige Konvergenz	82
5.4. Potenzreihen	86
5.5. Spezielle Funktionen	89
Kapitel 6. Topologische Grundlagen	97
6.1. Topologische Räume	97
6.2. Metrische Räume	103
6.3. Die Topologie des $\mathbb{R}^n$	109
Kapitel 7. Differentialrechnung in mehreren Variablen	111
7.1. Partielle Differentiation	111
7.2. Totale Differentiation	114
7.3. Der Satz von TAYLOR und lokale Extrema	119
7.4. Implizite Funktionen und lokale Invertierbarkeit	127
Kapitel 8. Integralrechnung in $\mathbb{R}^n$	132
8.1. Mehrfachintegrale	132
8.2. Die Transformationsformel	138

8.3. Integration stetiger Funktionen auf kompakten Mengen	146
8.4. Der GAUSSsche Integralsatz	155
Literaturverzeichnis	159

KAPITEL 1

Grundlagen der Mathematik

Dieses einführende Kapitel hat noch nichts spezifisch mit Analysis zu tun, sondern ist grundlegend für die Mathematik als Ganze. Wir behandeln Aussagen- und Prädikatenlogik, Mengenlehre, Relationen und Abbildungen, und schließlich (natürliche, ganze, rationale, reelle und komplexe) Zahlen. Bei den logischen und mengentheoretischen Grundlagen gehen wir ‚naiv‘ vor, d.h. wir geben lediglich unpräzise Erklärungen der Begriffe ‚Aussage‘ und ‚Menge‘. Für eine axiomatische Einführung dieser Begriffe verweisen wir auf die Literatur, etwa [7] und [6].

1.1. Logik

Logik untersucht die Struktur von Aussagen und die Regeln, aufgrund derer Aussagen aus anderen Aussagen hergeleitet werden können. Dabei betrachtet sie Aussagen sowohl in natürlichen als auch in formalen Sprachen. Während die Logik in der Antike als reines ‚Organon‘ (Werkzeug) für die Wissenschaften gesehen wurde, hat sie sich seit dem 19. Jahrhundert zu einer eigenständigen wissenschaftlichen Disziplin entwickelt, die heute in der Philosophie, der Mathematik und der theoretischen Informatik gleichermaßen beheimatet ist.

1.1.1. Aussagenlogik.

1.1.1.1. *Aussagen und Junktoren.* Eine *Aussage* ist ein Satz (in natürlicher oder formaler Sprache), dem sinnvoll ein eindeutiger *Wahrheitswert* zugeordnet werden kann, der also entweder wahr oder falsch ist. Beispiele für Aussagen sind etwa

- *Heute scheint in Ulm die Sonne* (wahr oder falsch, je nachdem, ob in Ulm heute die Sonne scheint);
- $2+2=4$ (wahr);
- $2+2=5$ (falsch);
- *Jede gerade Zahl größer als 2 ist Summe zweier Primzahlen* (Goldbachsche Vermutung; seit 1742 weder bewiesen noch widerlegt).

Das letzte Beispiel zeigt, daß eine Äußerung selbst dann eine Aussage sein kann, wenn ihr Wahrheitswert unbekannt ist. Ambiguitäten wie im ersten Beispiel (wo der Wahrheitswert vom Zeitpunkt der Aussage abhängt) treten bei mathematischen Aussagen üblicherweise nicht auf.

Keine Aussagen sind hingegen Fragen, Exklamationen, oder syntaktisch unzulässige Zeichenfolgen, wie z.B.

- *Wer hat die Kokosnuß geklaut?*
- *Zefix noch amol!*
- $x+ = \int$.

Ferner unterscheidet man zwischen Aussagen und bloßen *Termen* wie x^2 , ‚Mahatma Gandhi‘ oder $\sum_{n=1}^{\infty} \frac{1}{n^2}$. Aussagen sind aus Termen zusammengesetzt, aber ein Term allein ist noch keine Aussage.

Man kann neue Aussagen gewinnen, indem man sie aus bereits bekannten zusammensetzt. Beispiele solcher zusammengesetzter Aussagen sind

- *In Ulm scheint heute die Sonne und in München regnet es;*

- *Kräht der Hahn auf dem Mist, so ändert sich das Wetter oder es bleibt wie es ist;*
- *Eine reelle Zahl ist genau dann nichtnegativ, wenn sie das Quadrat einer reellen Zahl ist.*

Um dies zu formalisieren, bezeichnen wir für den Rest dieses Unterkapitels Aussagen mit großen lateinischen Buchstaben $A, B, C, \dots$ und führen die folgenden *Junktoren* ein:

- $\neg$ Negation (Verneinung)¹;
- $\wedge$ Konjunktion (,und‘);
- $\vee$ Disjunktion (,oder‘);
- $\Rightarrow$ Implikation (,wenn, dann‘);
- $\Leftrightarrow$ Äquivalenz (,genau dann, wenn‘).

Die klassische Aussagenlogik, die heute von fast allen Mathematiker*innen akzeptiert wird, folgt dem *Extensionalitätsprinzip*: Der Wahrheitswert einer zusammengesetzten Aussage hängt ausschließlich von den Wahrheitswerten der einzelnen Teilaussagen ab². Dies mag selbstverständlich erscheinen, aber wir werden gleich im Zusammenhang mit der Implikation sehen, daß natürliche Sprachen kaum als extensional betrachtet werden können.

Das Extensionalitätsprinzip ermöglicht uns, die Bedeutung der oben aufgeführten Junktoren festzulegen, indem wir die Wahrheitswerte der jeweiligen Junktion für alle möglichen Wahrheitswerte der Teilaussagen angeben. Man macht dies oft mithilfe von *Wahrheitstafeln*:

A	$\neg A$	A	B	$A \wedge B$	A	B	$A \vee B$
w	f	w	w	w	w	w	w
w	f	w	f	f	w	f	w
f	w	f	w	f	f	w	w
		f	f	f	f	f	f

A	B	$A \Rightarrow B$	A	B	$A \Leftrightarrow B$
w	w	w	w	w	w
w	f	f	w	f	f
f	w	w	f	w	f
f	f	w	f	f	w

Diese Konventionen stimmen im Falle der Negation und der Konjunktion gewiß mit dem umgangssprachlichen Gebrauch überein: Die Negation einer Aussage ist wahr, wenn die Aussage falsch ist, und umgekehrt; ‚ A und B ‘ ist nur dann wahr, wenn sowohl A als auch B wahr sind. Bei der Disjunktion ist zu beachten, daß es sich um ein *einschließendes Oder* handelt: ‚ A oder B ‘ gilt auch dann als wahr, wenn *beide* Teilaussagen wahr sind. Betrachte hierzu ein Beispiel aus der Umgangssprache:

Du entschuldigst dich sofort bei mir oder ich erzähle es meiner Mama!

Im Alltag würde die beschuldigte Person diese Drohung zurecht dahingehend auffassen, daß sie durch die geforderte Entschuldigung das Verpetztwerden umgehen kann. Aussagenlogisch wäre es allerdings auch zulässig, die Person trotz sofortiger Entschuldigung zu verpetzen.

Eine noch größere Diskrepanz zwischen formalem und alltäglichem Sprachgebrauch besteht bei der Implikation. Unproblematisch sind noch solche Sätze, in denen das Antezedens A die kausale Ursache für das Konsequens B ist:

Wenn es regnet, wird die Straße naß.

¹Die Negation ist strenggenommen kein Junktor, da sie nicht zwei Aussagen miteinander verknüpft, sondern nur auf eine einzelne Aussage wirkt.

²Einen Überblick über nicht-extensionale (intensionale) Logiken gibt WILHOLT [13]. Dort finden Sie auch eine ausgezeichnete, ausführlichere Darstellung der klassischen Aussagen- und Prädikatenlogik.

Für Verwirrung sorgen dagegen (wahre) Implikationen, bei denen Antezedens und Konsequens in keinerlei Sinnzusammenhang stehen:

- *Wenn Hans im April Geburtstag hat, hat das Ulmer Münster den höchsten Kirchturm der Welt;*
- *Wenn Hans im April Geburtstag hat, hat der Kölner Dom den höchsten Kirchturm der Welt;*
- *Wenn Hans im September Geburtstag hat, hat das Ulmer Münster den höchsten Kirchturm der Welt.*

Angenommen, Hans habe tatsächlich im September Geburtstag, so handelt es sich um allesamt wahre Aussagen (allein die Aussage *Wenn Hans im September Geburtstag hat, hat der Kölner Dom den höchsten Kirchturm der Welt* wäre falsch). Wir beobachten daran zweierlei: Erstens kann eine Implikation auch dann wahr sein, wenn das Konsequens falsch ist – sofern das Antezedens ebenfalls falsch ist (lat. *ex falso quodlibet*, ‚aus Falschem folgt Beliebiges‘). Zweitens sind die Implikationen wahr, obwohl Hansens Geburtstag mit der Frage nach dem höchsten Kirchturm inhaltlich überhaupt nichts zu tun hat. Dies kann aufgrund des Extensionalitätsprinzips auch gar nicht anders sein: Demzufolge hängt ja der Wahrheitswert einer Implikation nur von den Wahrheitswerten, nicht aber vom ‚Sinn‘ (der *Intension*) der beiden Teilaussagen ab.

Noch deutlicher wird die Differenz zwischen Implikation und Kausalität an folgendem Beispiel: Angenommen, ich habe schon länger nichts mehr von meiner Cousine gehört; insbesondere bin ich über den Ausgang ihrer rezenten Abiturprüfung nicht informiert. Nun schreibt sie mir, sie habe sich für ein Studium an der Uni Ulm eingeschrieben. Daraus folgere ich (zu recht), daß sie die Abiturprüfung bestanden hat³, erkenne also die folgende Implikation als wahr:

Wenn die Cousine nun studiert, dann hat sie zuvor ihre Abiturprüfung bestanden.

Dabei ist das Studium hier keineswegs die *Ursache* für das bestandene Abitur, sondern eher im Gegenteil. Bei der Implikation geht es also um zulässige logische Schlüsse, nicht um Kausalbeziehungen. In der Mathematik spielt der Begriff der Kausalität ohnehin keine Rolle – obwohl Mathematiker*innen in ihren Argumentationen regelmäßig das Wort ‚weil‘ verwenden.

Noch ein Hinweis zur Terminologie: Im Falle der Implikation $A \Rightarrow B$ sagt man auch ‚ A impliziert B ‘, ‚ B folgt aus A ‘, ‚ A ist hinreichend für B ‘, ‚ B ist notwendig für A ‘, ‚ A ist stärker als B ‘, oder ‚ B ist schwächer als A ‘.

1.1.1.2. *Analyse zusammengesetzter Aussagen und Tautologien.* Wahrheitstabellen erlauben nicht nur die Definition der gängigen Junktoren, sondern auch die Analyse komplizierterer zusammengesetzter Aussagen, die man aus den mit $A, B, \dots$ bezeichneten ‚atomaren Aussagen‘ durch wiederholte Verknüpfung mittels der Junktoren gewinnt. Dazu legen wir zunächst ‚Vorfahrtsregeln‘ für Junktoren fest: Wir vereinbaren, daß $\neg$, $\wedge$ und $\vee$ stärker binden als $\Rightarrow$ und $\Leftrightarrow$ (vgl. ‚Punkt-vor-Strich‘ in der Arithmetik). So meinen wir etwa mit $A \Rightarrow B \wedge C$ die Aussage $A \Rightarrow (B \wedge C)$ und nicht etwa die Aussage $(A \Rightarrow B) \wedge C$.

Als Beispiel wollen wir wissen, wie der Wahrheitswert der Aussage $A \wedge B \Leftrightarrow B \vee C$ von den Wahrheitswerten der Teilaussagen A, B, C abhängt. Darüber gibt folgende Wahrheitstafel Aufschluß:

³Wir vernachlässigen hier mögliche Zugänge zu einem Universitätsstudium ohne Abitur.

A	B	C	$A \wedge B$	$B \vee C$	$A \wedge B \Leftrightarrow B \vee C$
w	w	w	w	w	w
w	w	f	w	w	w
w	f	w	f	w	f
w	f	f	f	f	w
f	w	w	f	w	f
f	w	f	f	w	f
f	f	w	f	w	f
f	f	f	f	f	w

Betrachten wir als ein weiteres Beispiel die Aussage $A \wedge (A \Rightarrow B) \Rightarrow B$:

A	B	$A \Rightarrow B$	$A \wedge (A \Rightarrow B)$	$A \wedge (A \Rightarrow B) \Rightarrow B$
w	w	w	w	w
w	f	f	f	w
f	w	w	f	w
f	f	w	f	w

Hier fällt auf, daß die Aussage $A \wedge (A \Rightarrow B) \Rightarrow B$ *unabhängig von den Wahrheitswerten von A und B* stets wahr ist. Das ist nicht besonders überraschend, denn diese Aussage bedeutet: Ist A wahr und folgt B aus A , so ist auch B wahr.

Man nennt zusammengesetzte Aussagen, die unabhängig vom Wahrheitswert ihrer atomaren Teilaussagen wahr sind, *logisch wahr* oder *tautologisch*. Der Begriff ‚logisch wahr‘ bezieht sich darauf, daß solche Aussagen allein kraft ihrer logischen Struktur wahr sind und nicht nur aufgrund kontingenter Gegebenheiten. Hat etwa Hans im September Geburtstag, so ist die Aussage *Hans hat im September oder im April Geburtstag* wahr, aber nicht tautologisch; die Aussage *Hans hat im September oder in einem anderen Monat Geburtstag* ist dagegen eine Tautologie, weil sie wahr ist, egal wann Hans Geburtstag hat.

Wir geben weitere Beispiele für bekannte Tautologien, zusammen mit ihren traditionellen Bezeichnungen⁴:

$A \wedge (A \Rightarrow B) \Rightarrow B$	(modus ponens/direkter Beweis)
$(A \Rightarrow B) \wedge \neg B \Rightarrow \neg A$	(modus tollens/reductio ad absurdum/Widerspruchsbeweis)
$A \vee \neg A$	(tertium non datur/Prinzip des ausgeschlossenen Dritten)
$(A \Rightarrow B) \wedge (\neg A \Rightarrow B) \Rightarrow B$	(klassisches Dilemma/Fallunterscheidung)
$\neg \neg A \Leftrightarrow A$	doppelte Negation
$(A \Rightarrow B) \Leftrightarrow (\neg B \Rightarrow \neg A)$	Kontraposition
$(A \Rightarrow B) \wedge (B \Rightarrow A) \Leftrightarrow (A \Leftrightarrow B)$	

1.1.1.3. *Logische Schlußregeln als mathematische Beweisstrategien.* Solche Tautologien geben insbesondere mögliche Strategien für mathematische Beweise vor. Wir geben je ein Beispiel für *reductio ad absurdum* und für die Fallunterscheidung:

BEISPIEL 1.1 (Widerspruchsbeweis). Wir zeigen, daß $\sqrt{2}$ irrational ist.⁵ Sei dazu A die Aussage $\sqrt{2}$ ist rational. Wir wollen die Annahme A zum Widerspruch führen, d.h.

⁴Der Nachweis der Tautologie kann jeweils mittels Wahrheitstafeln erbracht werden, was den Lesenden zur Übung empfohlen sei. – WILHOLT [13, S. 37] weist darauf hin, daß die Bezeichnungen der Schlußregeln (z.B. ‚Prinzip vom ausgeschlossenen Dritten‘) sich auf die korrespondierenden metasprachlichen Schlüsse und nicht, wie hier insinuiert, auf die objektsprachlichen Formeln beziehen. Wir ignorieren solche Einwände, da Objekt- und Metasprache in der mathematischen Praxis ohnehin kaum jemals scharf unterschieden werden.

⁵Wir werden die Begriffe ‚rationale Zahl‘ und ‚Quadratwurzel‘ noch eingehend besprechen. Für das Beispielmateriale sei aber zunächst noch auf Ihre Schulkenntnisse verwiesen.

aus ihr eine falsche Aussage folgern. Mit der Regel *modus tollens* folgt dann $\neg A$, also die gewünschte Irrationalität von $\sqrt{2}$.

Wenn also $\sqrt{2}$ rational wäre, so ließe es sich durch einen vollständig gekürzten Bruch $\frac{p}{q}$ darstellen, wobei p und q ganze Zahlen ohne gemeinsamen Teiler sind. Nach Definition der Quadratwurzel wäre dann

$$2 = (\sqrt{2})^2 = \frac{p^2}{q^2}$$

bzw. $p^2 = 2q^2$. Die Zahl p^2 ist also durch 2 teilbar, und damit ist auch p selbst durch 2 teilbar. Wir können also schreiben $p = 2r$ für eine ganze Zahl r , und es folgt $4r^2 = 2q^2$ bzw. $q^2 = 2r^2$. Also ist auch q durch 2 teilbar, und wir folgern die Aussage B : *p und q haben einen gemeinsamen Teiler*. Doch zu Beginn des Beweises hatten wir p und q teilerfremd gewählt, sodaß also $\neg B$ wahr ist. Wir haben damit einerseits $A \Rightarrow B$ und andererseits $\neg B$ gezeigt, sodaß $\neg A$ (die Irrationalität von $\sqrt{2}$) folgt. $\square$

BEISPIEL 1.2 (Fallunterscheidung). Sei n eine natürliche Zahl, die nicht durch 3 teilbar ist. Wir zeigen: n^2 läßt bei Division durch 3 den Rest 1.

Hier wählen wir als zu beweisende Aussage B : *n^2 läßt bei Division durch 3 den Rest 1*; die Aussage A laute: *n läßt bei Division durch 3 Rest 1*. Nach dem Prinzip der Fallunterscheidung ist B gezeigt, wenn wir es sowohl aus A als auch aus $\neg A$ herleiten können.

Fall A: Falls A wahr ist, gibt es eine natürliche Zahl k mit $n = 3k + 1$. Dann aber ist

$$n^2 = (3k + 1)^2 = 9k^2 + 6k + 1 = 3(3k^2 + 2k) + 1,$$

was bei Division durch 3 Rest 1 ergibt. Somit ist $A \Rightarrow B$ gezeigt.

Fall $\neg A$: Falls $\neg A$, so läßt n bei Division durch 3 den Rest 0 oder 2. Rest 0 ist nach Voraussetzung ausgeschlossen (da n nicht durch 3 teilbar ist). Also existiert eine natürliche Zahl m mit $n = 3m + 2$, und damit

$$n^2 = (3m + 2)^2 = 9m^2 + 12m + 4 = 3(3m^2 + 4m + 1) + 1,$$

also läßt n^2 bei Division durch 3 den Rest 1, und wir haben $\neg A \Rightarrow B$ gezeigt. Aus beiden Fällen zusammen folgt B . $\square$

Eine Warnung scheint hier angebracht: Logische Schlußregeln bilden immer nur das ‚Gerüst‘ eines mathematischen Beweises. Um dieses Gerüst mit zielführendem Inhalt zu füllen, ist fast immer ein gewisses Maß an Originalität, Intuition und/oder Erfahrung erforderlich (letztere erwirbt man nur mit harter Arbeit). Es gibt kein ‚Kochrezept‘ für mathematische Beweise.

1.1.2. Prädikatenlogik. Viele mathematische Sätze treffen nicht nur Aussagen über einzelne Objekte (so wie *$\sqrt{2}$ ist irrational*), sondern beanspruchen allgemeine Gültigkeit für alle Objekte aus einer bestimmten Klasse (z.B. *Jede natürliche Zahl kann eindeutig in Primfaktoren zerlegt werden*). Andere Sätze wiederum behaupten die Existenz eines bestimmten Objekts (*Es existiert eine positive reelle Zahl x mit $x^2 = 2$*). Die Formalisierung solcher Aussagen erfordert die Erweiterung der Aussagenlogik zur *Prädikatenlogik*, die zusätzlich zu den aussagenlogischen Junktoren noch über die *Quantoren* $\forall$ (‚für alle‘) und $\exists$ (‚es existiert‘) verfügt.

1.1.2.1. Prädikate und Quantoren. Betrachten wir die Formel $x^2 = 2$: Solange nicht festgelegt ist, welchen Wert x annimmt, kann dieser Formel kein Wahrheitswert zugeordnet werden (für manche x ist die Gleichung erfüllt, für andere nicht). Sie ist deshalb keine Aussage im Sinne von Abschnitt 1.1.1. Man bezeichnet einen solchen Satz, der von einer oder mehreren Variablen abhängt, als *Aussageform* oder *Prädikat*. Ein weiteres Beispiel für ein Prädikat wäre *x hat im September Geburtstag*. Mehrere Prädikate können wie gehabt durch die aussagenlogischen Junktoren verbunden werden: Für jede Wahl der Variablen

ergibt sich dann der Wahrheitswert der zusammengesetzten Aussageform aus den Wahrheitswerten der Teilaussageformen.

Eine Möglichkeit, aus einem Prädikat eine Aussage zu machen, besteht in der *Quantifizierung* über die Variablen: Die beiden Sätze *Es existiert ein x mit $x^2 = 2$* sowie *Für alle x gilt $x^2 = 2$* sind Aussagen, zumindest wenn vereinbart wurde, aus welcher Zahlenmenge x gewählt werden darf. Reden wir über reelle Zahlen, so ist die erste Aussage wahr, die zweite falsch. In der Prädikatenlogik schreibt man die beiden Aussagen wie folgt:

$$\exists x : x^2 = 2 \quad \text{bzw.} \quad \forall x : x^2 = 2.$$

Prädikate können mehrere Argumente haben: So können wir ein zweistelliges Prädikat $P(x, y)$ definieren⁶ durch $x \leq y$, und die folgende (wahre) Aussage formulieren:

$$\forall x \forall y : P(0, x) \wedge P(0, y) \wedge P(x, y) \Rightarrow P(x^2, y^2),$$

das bedeutet: Für alle x und y mit $x \geq 0$, $y \geq 0$ und $x \leq y$ gilt $x^2 \leq y^2$.

Wenn nicht aus dem Kontext klar ist, über welche Variablenwerte quantifiziert wird, sollte man dies explizit angeben, z.B. wie folgt:

$$\exists x \in \mathbb{R} : x^2 = 2 \quad \text{bzw.} \quad \exists x \in \mathbb{Q} : x^2 = 2,$$

wobei $\mathbb{R}$ und $\mathbb{Q}$ die Mengen der reellen bzw. rationalen Zahlen bezeichnen; man beachte unbedingt, daß die erste Aussage wahr, die zweite aber falsch ist!

Eine weitere Bemerkung betrifft die *gebundene Umbenennung*: In der Aussage $\forall x : P(x)$ nimmt x keinen bestimmten Wert an, sondern ‚läuft‘ durch die Menge der zulässigen Werte. Deshalb ist es gleichgültig, wie die Laufvariable heißt, und man kann sie beliebig umbenennen; $\forall x : P(x)$ ist also gleichbedeutend mit $\forall y : P(y)$ oder mit $\forall \alpha : P(\alpha)$ etc. Das Gleiche gilt natürlich für Existenzquantoren. Ähnliche Formen von Laufvariablen, die beliebig umetikettiert werden dürfen, werden uns auch bald in Summen und später in Integralen begegnen.

1.1.2.2. Negation und Quantorenvertauschung. Die Negation von *Alle Menschen sind sterblich* lautet natürlich nicht *Alle Menschen sind unsterblich*, sondern *Nicht alle Menschen sind sterblich* oder, äquivalent dazu, *Es gibt unsterbliche Menschen*. Man kann sich dies vor Augen führen, indem man zunächst formuliert *Es trifft nicht zu, daß alle Menschen unsterblich sind*. Allgemein gilt

$$\neg \forall x : P(x) \Leftrightarrow \exists x : \neg P(x).$$

Auch im Falle mehrerer Quantoren trifft folgende Regel zu: Ein Negationszeichen kann vom Prädikat vor die Quantoren gezogen werden (oder umgekehrt), sofern jeder Quantor umgedreht wird (d.h. aus einem All- wird ein Existenzquantor und umgekehrt). Betrachten wir als Beispiel die Definition der Stetigkeit einer Funktion, wie wir sie später in diesem Semester ausführlich behandeln werden. In Prädikatenlogik lautet diese Definition: Eine Funktion f heißt stetig im Punkt x , wenn

$$\forall \epsilon > 0 \quad \exists \delta > 0 \quad \forall y : |x - y| < \delta \Rightarrow |f(x) - f(y)| < \epsilon.$$

Was bedeutet es dann, daß eine Funktion im Punkt x *nicht* stetig ist? Nach der eben beschriebenen Regel erhalten wir für die Negation:

$$\exists \epsilon > 0 \quad \forall \delta > 0 \quad \exists y : |x - y| < \delta \not\Rightarrow |f(x) - f(y)| < \epsilon.$$

Hier haben wir $\not\Rightarrow$ für ‚impliziert nicht‘ geschrieben. (Man überzeuge sich zur Übung, daß $A \not\Rightarrow B$ logisch äquivalent zu $A \wedge \neg B$ ist.)

⁶Genauer gesagt wird der zweistellige Prädikatbuchstabe P durch die *Relation* (s. nächstes Unterkapitel) ‚ $\leq$ ‘ semantisch interpretiert; wir verzichten hier und im Folgenden auf eine strikte Unterscheidung zwischen syntaktischen und semantischen Begriffen, da dies in der mathematischen Praxis in der Regel keinen Mehrwert bietet und das aktuelle Kapitel unnötig aufblähen würde.

Eine berüchtigte Fallgrube ist die Vertauschung von Quantoren. Es gilt: All- und Existenzquantoren dürfen untereinander vertauscht werden, die Vertauschung eines Allquantors mit einem Existenzquantor ist hingegen unzulässig. Es ist also

$$\forall x \forall y : P(x, y) \Leftrightarrow \forall y \forall x : P(x, y),$$

$$\exists x \exists y : P(x, y) \Leftrightarrow \exists y \exists x : P(x, y),$$

aber im Allgemeinen

$$\forall x \exists y : P(x, y) \not\Leftrightarrow \exists y \forall x : P(x, y).$$

Letzteres bedarf der Erklärung. Sei x dafür aus der Menge der Töpfe und y aus der Menge der Deckel, und das Prädikat $P(x, y)$ werde interpretiert als *Auf den Topf x paßt der Deckel y* . Dann bedeutet $\forall x \exists y : P(x, y)$: *Für jeden Topf gibt es einen Deckel, der auf ihn paßt* oder einfacher: *Auf jeden Topf paßt ein Deckel*. Die Aussage $\exists y \forall x : P(x, y)$ dagegen bedeutet: *Es gibt einen Deckel, der auf jeden Topf paßt*. Dies sind offenkundig zwei völlig verschiedene Aussagen: Im ersten Falle darf jeder Topf einen unterschiedlichen Deckel haben, im zweiten Falle müßte *ein und derselbe Deckel* alle Töpfe zudecken! Einen solchen ‚Universaldeckel‘ gibt es in Wirklichkeit natürlich nicht, die erste Aussage dagegen ist viel plausibler, wenn nicht sogar wahr.

Halten wir fest: $\exists y \forall x : P(x, y)$ ist eine viel stärkere Behauptung als $\forall x \exists y : P(x, y)$, weil im ersten Fall das y nicht von x abhängen darf, im zweiten Fall aber schon. Wir werden auf diesen Unterschied später im Zusammenhang mit gleichmäßiger Stetigkeit von Funktionen und auch mit gleichmäßiger Konvergenz von Funktionenfolgen zurückkommen.

1.1.2.3. *Syllogismen*. Zu den bekanntesten logischen Figuren zählen seit ARISTOTELES die *Syllogismen*, deren 24 Unterarten im Mittelalter jeder Student, gleich welcher Fakultät, im ersten Studienjahr auswendig lernen mußte. Dies wird von Ihnen nicht verlangt werden. Gleichwohl ist es instruktiv zu sehen, wie man Syllogismen umstandslos in Prädikatenlogik schreiben kann.

BEISPIEL 1.3 (*Modus Barbara*). Folgende Aussage ist logisch wahr⁷:

$$\forall x : (P(x) \Rightarrow Q(x)) \wedge \forall y : (Q(y) \Rightarrow R(y)) \Rightarrow \forall z : (P(z) \Rightarrow R(z)).$$

Zum Beispiel: Alle Griechen sind Menschen. Alle Menschen sind sterblich. Also sind alle Griechen sterblich.

Hier haben wir definiert: $P(x)$ ‚ x ist ein Grieche‘, $Q(x)$ ‚ x ist ein Mensch‘, $R(x)$ ‚ x ist sterblich‘.

BEISPIEL 1.4 (*Modus Darii*). Auch die folgende Aussage ist logisch wahr:

$$\forall x : (P(x) \Rightarrow Q(x)) \wedge \exists y : (P(y) \wedge R(y)) \Rightarrow \exists z : (Q(z) \wedge R(z)).$$

Zur Übung denke man sich ein sinnvolles Beispiel aus.

1.2. Mengenlehre

Logische Zusammenhänge bringen lediglich die Struktur, die äußere Form eines mathematischen Arguments zum Ausdruck. Wovon aber ‚handeln‘ mathematische Aussagen, was also sind die Objekte, über die Mathematiker*innen sprechen? Ungeachtet kontroverser philosophischer Debatten um diese Fragen könnten wir naiv antworten: In der Mathematik geht es um Zahlen, Mengen, Funktionen, algebraische Strukturen (z.B. Körper oder Vektorräume), geometrische Figuren usw. Diese Konzepte sind scheinbar sehr divers, aber seit Ende des 19. Jahrhunderts wurde klar, daß alle diese mathematischen Objekte ultimativ

⁷Da in dieser Aussage alle Variablen gebunden vorkommen, können sie beliebig umbenannt werden. Insbesondere könnten wir genauso gut schreiben

$$\forall x : (P(x) \Rightarrow Q(x)) \wedge \forall x : (Q(x) \Rightarrow R(x)) \Rightarrow \forall x : (P(x) \Rightarrow R(x)).$$

als Mengen beschrieben werden können. In diesem Abschnitt werden wir z.B. sehen, wie der Funktionsbegriff auf den Mengenbegriff zurückgeführt werden kann.

1.2.1. Mengen und Mengenoperationen.

1.2.1.1. *Grundlegendes.* CANTOR, der als Begründer der Mengenlehre gilt, erläuterte den Begriff der Menge 1895 wie folgt [4]:

Unter einer ‚Menge‘ verstehen wir jede Zusammenfassung M von bestimmten wohlunterschiedenen Objecten m unsrer Anschauung oder unseres Denkens (welche die ‚Elemente‘ von M genannt werden) zu einem Ganzen.

Beispiele für Mengen sind $\{1, 2, 3\}$, $\{\pi\}$, die *leere Menge* (bezeichnet mit $\{\}$ oder $\emptyset$), die kein Element enthält, oder die Menge $\{\emptyset\}$, deren einziges Element die leere Menge ist. Mengen können auch unendlich viele Elemente enthalten, z.B. die Menge $\mathbb{R}$ der reellen Zahlen oder die Menge der gleichschenkligen Dreiecke in der Ebene. Elemente von Mengen müssen nicht zwingend mathematische Objekte sein: Auch {rot, grün, blau} oder {Kemal, Claudia, Giovanni, Kim} sind Mengen. Die Reihenfolge, in der die Elemente angegeben sind, spielt keine Rolle; ebenso werden Mehrfachnennungen ignoriert: $\{3, 1, 2, 3, 2, 3, 3\}$ ist gleichbedeutend mit $\{1, 2, 3\}$. Ist x ein Element von M , so schreibt man $x \in M$ (oder manchmal $M \ni x$). Ist x kein Element von M , so schreibt man $x \notin M$. Man beachte, daß $x \notin M$ nur eine Kurzschreibweise für $\neg x \in M$ ist.

Anstatt eine Menge durch Aufzählung aller Elemente anzugeben (was ohnehin nur für endliche Mengen funktioniert), kann man sie auch durch Prädikate definieren:

$$\mathbb{R}_0^+ = \{x \in \mathbb{R} : x \geq 0\}$$

ist die Menge der nichtnegativen reellen Zahlen. Ist P ein einstelliges Prädikat, so scheint es also, daß man die Menge $\{x : P(x)\}$ bilden könne. RUSSELL machte aber in einem Brief an FREGE 1902 eine folgenschwere Beobachtung: Betrachte das Prädikat $\neg(x \in x)$, dessen Argument x selbst eine Menge ist; Es sagt also aus, daß die Menge x sich nicht selbst als Element enthält. (Die Frage, ob es Mengen mit dieser wundersamen Eigenschaft überhaupt gibt, muß uns hier zum Glück gar nicht beschäftigen.) Definiere die Menge

$$M := \{x : \neg(x \in x)\},$$

also die Menge aller Mengen, die sich nicht selbst enthalten. Gilt dann $M \in M$? Falls ja, ist $M \notin \{x : \neg(x \in x)\} = M$; falls nein, ist $M \in \{x : \neg(x \in x)\} = M$. Es kann also weder $M \in M$ noch $M \notin M$ gelten, was wiederum dem Prinzip *tertium non datur* (Prinzip des ausgeschlossenen Dritten) widerspricht. Diese RUSSELLsche Antinomie kann nur umgangen werden, indem man die Mengenbildung $\{x : P(x)\}$ *nicht* für beliebige Prädikate zuläßt. Diese Gedanken führten um die Wende zum 20. Jahrhundert zu tiefgreifenden Entwicklungen in der Mengenlehre, unter anderem zu RUSSELLs Typentheorien und vor allem zu ZERMELOS axiomatischer Mengenlehre. In letzterer wird die Mengenbildung durch Prädikate in der Weise eingeschränkt, daß nur Vorschriften der Form

$$\{x \in G : P(x)\}$$

erlaubt sind, wobei G irgendeine Menge und P ein Prädikat ist.

Genau wie Aussagen unterliegen auch Mengen einem Extensionalitätsprinzip:

DEFINITION 1.5. Zwei Mengen sind *gleich*, wenn sie dieselben Elemente haben.

Auch diese Version des Extensionalitätsprinzips erscheint selbstverständlich. Betrachte aber die Mengen

$$M = \{x \in \mathbb{R} : x \geq 0\}, \quad N = \{x \in \mathbb{R} : \exists y \in \mathbb{R} : x = y^2\}.$$

Wir behandeln bald systematisch die reellen Zahlen, aber unter Rückgriff auf Schulwissen ist bekannt, daß jede nichtnegative reelle Zahl ein Quadrat ist (nämlich ihrer Quadratwurzel), und daß umgekehrt jedes Quadrat einer reellen Zahl nichtnegativ ist. Nach dem Extensionalitätsprinzip ist also $M = N$, obwohl die Vorschrift (die Intension), nach der die beiden Mengen gebildet wurden, völlig verschieden sind: Für die Definition von M wird die *Anordnung* der reellen Zahlen (also die Relation $\geq$) verwendet, für N eine *algebraische Operation* (nämlich das Quadrieren bzw. die Multiplikation).

DEFINITION 1.6 (Teilmengen).

- i) Eine Menge M heißt *Teilmenge* einer Menge N , falls

$$\forall x : (x \in M \Rightarrow x \in N).$$

In diesem Falle schreibt man $M \subset N$.

- ii) M heißt *echte Teilmenge* von N , falls $M \subset N$ und

$$\exists x : (x \in N \wedge x \notin M).$$

Man schreibt dann $M \subsetneq N$.

Man beachte, daß die leere Menge Teilmenge jeder Menge ist, und daß jede Menge Teilmenge ihrer selbst ist.

PROPOSITION 1.7. ⁸ *Es ist $M = N$ genau dann, wenn $M \subset N$ und $N \subset M$.*

BEWEIS. Nach Extensionalitätsprinzip (Definition 1.5) ist $M = N$ genau dann, wenn

$$\forall x : (x \in M \Leftrightarrow x \in N). \quad (1.1)$$

Mit der Tautologie $(A \Rightarrow B) \wedge (B \Rightarrow A) \Leftrightarrow (A \Leftrightarrow B)$ (Nachweis durch Wahrheitstafel!), angewendet auf die Aussagen $x \in M$ und $x \in N$, ist (1.1) äquivalent zu

$$\forall x : ((x \in M \Rightarrow x \in N) \wedge (x \in N \Rightarrow x \in M)),$$

und dies ist nach Definition 1.6 wiederum äquivalent dazu, daß $M \subset N$ und $N \subset M$. Dies war zu beweisen. $\square$

1.2.1.2. *Mengenoperationen.* Ebenso wie man aus bekannten Aussagen neue Aussagen durch diverse Verknüpfungen gewinnen kann, erhält man aus Mengen mithilfe bestimmter Operationen neue Mengen. Wir definieren:

$$M \cup N := \{x : x \in M \vee x \in N\} \quad \text{Vereinigung,}$$

$$M \cap N := \{x : x \in M \wedge x \in N\} \quad \text{Durchschnitt,}$$

$$M \setminus N := \{x \in M : x \notin N\} \quad \text{Differenz,}$$

$$M \Delta N := (M \setminus N) \cup (N \setminus M) \quad \text{symmetrische Differenz.}$$

Falls $M \subset G$, so heißt $G \setminus M$ auch *Komplement von M in G* und wird M^c geschrieben. Diese Notation ist hauptsächlich dann in Gebrauch, wenn sich die gesamte Untersuchung innerhalb einer ‚Grundmenge‘ G abspielt, die nicht jedes Mal eigens genannt werden muß. So ist bei (elementaren) zahlentheoretischen Aussagen meist klar, daß alle Mengen als Teilmengen der natürlichen Zahlen angenommen werden.

Zwei Mengen M und N heißen *disjunkt*, falls $M \cap N = \emptyset$.

Es gibt unzählige bekannte mengentheoretische Identitäten, die in Aussagen übersetzt werden, die wiederum z.B. durch Wahrheitstafeln als tautologisch nachgewiesen werden

⁸Eine *Proposition* ist eine mathematische Aussage, der nicht die Wichtigkeit eines *Satzes* (synonym eines *Theorems*) zukommt. Ob etwas noch eine Proposition oder schon ein Theorem ist oder aber ein *Lemma* (Hilfssatz) oder ein *Korollar* (einfache Folgerung aus einem Theorem), liegt im Ermessen der Autorin und ist in vielen Fällen Geschmackssache. Einige bekannte Resultate werden aus historischen Gründen als Lemmata (i.e. Plural von Lemma) bezeichnet, obwohl sie zweifellos den Rang eines Theorems beanspruchen dürfen; so etwa das *Zornsche Lemma* oder das *Lemma von Sperner*.

können. Als Beispiel betrachte die MORGANSchen Gesetze, wonach für beliebige Mengen M und N , die beide Teilmengen einer (Grund-)Menge G sind, gilt

$$(M \cap N)^c = M^c \cup N^c \quad \text{sowie} \quad (M \cup N)^c = M^c \cap N^c.$$

Wir zeigen nur die erste Gleichheit. Nach Extensionalitätsprinzip (Def. 1.5) genügt es zu zeigen: Für jedes $x \in G$ gilt

$$x \in (M \cap N)^c \Leftrightarrow x \in M^c \cup N^c. \quad (1.2)$$

Nun ist $x \in (M \cap N)^c$ nach Definition des Komplements und des Durchschnitts äquivalent zu $\neg(x \in M \wedge x \in N)$, und die rechte Seite $x \in M^c \cup N^c$ ist äquivalent zu $\neg x \in M \vee \neg x \in N$. Mit einer Wahrheitstafel kann aber die Tautologie

$$\neg(A \wedge B) \Leftrightarrow \neg A \vee \neg B$$

nachgewiesen werden, aus der die Äquivalenz (1.2) schließlich folgt.

DEFINITION 1.8 (Potenzmenge). Sei M eine Menge, so heißt die Menge aller Teilmengen von M die *Potenzmenge* von M . Man bezeichnet sie mit $\mathcal{P}(M)$.

Die Potenzmenge von $\{1, 2, 3\}$ lautet etwa

$$\{\emptyset, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{2, 3\}, \{1, 3\}, \{1, 2, 3\}\}.$$

DEFINITION 1.9 (kartesisches Produkt). Seien M und N Mengen.

i) Für $x \in M$ und $y \in N$ bezeichnet

$$(x, y) := \{\{x\}, \{x, y\}\}$$

ein *geordnetes Paar*⁹.

ii) Die Menge aller geordneten Paare von Elementen aus M und N heißt *kartesisches Produkt* (oder einfach *Produktmenge*) von M und N :

$$M \times N := \{(x, y) : x \in M \wedge y \in N\}.$$

PROPOSITION 1.10 (Paaraxiom von PEANO). Seien $(u, v), (x, y) \in M \times N$. Dann ist $(u, v) = (x, y)$ genau dann, wenn $u = x$ und $v = y$.

BEWEIS. Um die Äquivalenz zu zeigen, zeigen wir beide Implikationen:

, $\Rightarrow$ ': Sei $(u, v) = (x, y)$, also nach Definition $\{\{u\}, \{u, v\}\} = \{\{x\}, \{x, y\}\}$. Nach Extensionalitätsprinzip (Def. 1.5) sind die Elemente dieser Mengen gleich; insbesondere ist $\{u\} = \{x\}$ oder $\{u\} = \{x, y\}$. Im ersten Falle folgt $u = x$ wiederum aus der Extensionalität. Im zweiten Falle muß $x = y$ sein, da eine Menge mit nur einem Element nicht gleich einer Menge mit zwei (verschiedenen) Elementen sein kann. Dann gilt aber wieder $\{u\} = \{x, x\} = \{x\}$ und somit $u = x$. In jedem Falle ist also $u = x$.

Da $\{u\} = \{x\}$, muß wegen der Gleichheit $\{\{u\}, \{u, v\}\} = \{\{x\}, \{x, y\}\}$ auch gelten $\{u, v\} = \{x, y\}$. Abermals wegen $u = x$ folgt daraus schließlich $v = y$.

, $\Leftarrow$ ': Diese Richtung ist einfacher: Wenn $u = x$ und $v = y$, so folgt $\{\{u\}, \{u, v\}\} = \{\{x\}, \{x, y\}\}$ sofort aus dem Extensionalitätsprinzip. $\square$

Die Bezeichnung ‚kartesisch‘ verweist auf DESCARTES, der im 17. Jahrhundert im Rahmen seiner analytischen Geometrie Koordinaten (x, y) in der euklidischen Ebene einföhrte. Bei solchen Koordinaten kommt es nicht nur auf die beiden Koordinaten x, y an, sondern auch auf ihre Reihenfolge. Ein geordnetes Paar (x, y) unterscheidet sich also von der Menge $\{x, y\}$ dadurch, daß die Reihenfolge der beiden Elemente festgelegt ist: Zwar gilt $\{x, y\} = \{y, x\}$, aber im Allgemeinen *nicht* $(x, y) = (y, x)$. Die mengentheoretische Formulierung des geordneten Paares stammt von KURATOWSKI.

⁹Das Zeichen ‚ $:=$ ‘ bedeutet ‚definitionsgemäß gleich‘.

Als Beispiel geben wir

$$\{1, 2, 3\} \times \{1, 2\} = \{(1, 1), (1, 2), (2, 1), (2, 2), (3, 1), (3, 2)\}.$$

Ein weiteres Beispiel ist $\mathbb{R} \times \mathbb{R}$, das oft als $\mathbb{R}^2$ geschrieben wird; es handelt sich um alle Paare (x, y) mit reellen Komponenten x und y . Ein solches Paar kann man sich – nach DESCARTES – als Punkt in der Ebene geometrisch veranschaulichen.

1.2.2. Relationen und Abbildungen.

1.2.2.1. Relationen.

DEFINITION 1.11. Eine *Relation* auf einer Menge M ist eine Teilmenge von $M \times M$.

Beispielsweise ist die Relation ‚kleiner gleich‘ auf $\mathbb{R}$ gemäß dieser Definition genau die Menge der Paare $(x, y) \in \mathbb{R}^2$, für die $x \leq y$. Ein weiteres Beispiel: Ist M eine Menge von Personen, so kann man darauf die Relation F definieren durch $(x, y) \in F$ genau dann, wenn x und y befreundet sind. Als drittes Beispiel betrachten wir auf den natürlichen Zahlen $\mathbb{N}$ die Teilbarkeitsrelation $|\colon (n, m) \in |$ genau dann, wenn n ein Teiler von m ist, wenn also ein $k \in \mathbb{N}$ existiert mit $m = kn$.

Für konkrete Relationen schreibt man üblicherweise nicht $(x, y) \in R$, sondern $R(x, y)$ oder xRy . So schreibt man $x \leq y$ und nicht $(x, y) \in \leq$, ebenso $n|m$ statt $(n, m) \in |$.

DEFINITION 1.12. Eine Relation $R \subset M \times M$ heißt

i) *reflexiv*, falls

$$\forall x \in M : (x, x) \in R,$$

ii) *symmetrisch*, falls

$$\forall x, y \in M : (x, y) \in R \Rightarrow (y, x) \in R,$$

iii) *transitiv*, falls

$$\forall x, y, z \in M : (x, y) \in R \wedge (y, z) \in R \Rightarrow (x, z) \in R.$$

Die Relation $\leq$ ist reflexiv (da $x \leq x$ für alle $x \in \mathbb{R}$), aber nicht symmetrisch (da z.B. $3 \leq 5$, aber nicht $5 \leq 3$), wohl aber transitiv (wenn $x \leq y$ und $y \leq z$, dann folgt $x \leq z$). Ebenso sieht es bei der Teilbarkeitsrelation aus: Sie ist reflexiv und transitiv, aber nicht symmetrisch. Eine sehr wichtige Klasse von Relationen sind die *Äquivalenzrelationen*:

DEFINITION 1.13 (Äquivalenzrelation). Eine Relation $R \subset M \times M$ heißt *Äquivalenzrelation*, wenn sie reflexiv, symmetrisch und transitiv ist. Für jedes $x \in M$ ist die *Äquivalenzklasse von x bzgl. R* definiert als die Menge

$$[x] = \{y \in M : (x, y) \in R\}.$$

BEISPIEL 1.14. Sei n eine natürliche Zahl, dann definieren wir auf den ganzen Zahlen $\mathbb{Z}$ eine Relation R folgendermaßen: $(r, s) \in R$ genau dann, wenn $r - s$ durch n teilbar ist (d.h. wenn es $k \in \mathbb{Z}$ gibt mit $r - s = kn$). Statt $(r, s) \in R$ schreiben wir

$$r \equiv s \pmod{n}$$

(sprich: „ r ist gleich s modulo n .“) So ist etwa $11 \equiv 1 \pmod{2}$, $17 \equiv 5 \pmod{12}$, $5 \equiv -1 \pmod{3}$.

Wir zeigen, daß dies eine Äquivalenzrelation ist: Für jedes $r \in \mathbb{Z}$ ist $r - r = 0$, und 0 ist durch jede ganze Zahl teilbar; also gilt $r \equiv r \pmod{n}$, und die Relation ist reflexiv. Sie ist auch symmetrisch, denn wenn $r - s$ durch n teilbar ist, so auch $s - r = -(r - s)$. Gilt schließlich $r \equiv s \pmod{n}$ und $s \equiv q \pmod{n}$, so existieren $k, l \in \mathbb{Z}$ mit $r - s = kn$ und $s - q = ln$. Mithin ist

$$r - q = (r - s) + (s - q) = kn + ln = (k + l)n,$$

also ist $r - q$ durch n teilbar und damit $r \equiv q \pmod{n}$. Das ist genau die behauptete Transitivität.

Schließlich untersuchen wir die Äquivalenzklassen: $[0]$ ist die Menge der durch n teilbaren ganzen Zahlen (also der Vielfachen von n); $[1]$ ist die Menge der ganzen Zahlen, die bei Division durch n den Rest 1 lassen, usw. $[n]$ ist wieder gleich $[0]$, da $n \equiv 0 \pmod{n}$. Die Menge der ganzen Zahlen kann also in die n paarweise disjunkten Äquivalenzklassen $[0], [1], \dots, [n-1]$ zerlegt werden. (Man überzeuge sich zur Übung, daß diese Äquivalenzklassen tatsächlich disjunkt sind und daß es keine weiteren gibt.)

Ein lebensnahes Beispiel ist die Uhr: Hier identifiziert man z.B. 17 Uhr und 5 Uhr, weil die Differenz 12 ist. Wenn es jetzt 10 Uhr ist und man wissen möchte, wie spät es in 39 Stunden ist, so addiert man und erhält ,49 Uhr'; dies ist aber äquivalent zu 1 Uhr, da $49 \equiv 1 \pmod{12}$.

1.2.2.2. Abbildungen.

DEFINITION 1.15 (Abbildung). Seien M, N Mengen. Eine Teilmenge $f \subset M \times N$ heißt *Abbildung* von M nach N , falls

$$\forall x \in M \quad \exists y \in N : (x, y) \in f$$

und

$$\forall x \in M \quad \forall y, z \in N : (x, y) \in f \wedge (x, z) \in f \Rightarrow y = z. \quad (1.3)$$

Man nennt M den *Definitionsbereich* und N den *Wertebereich* der Abbildung f und schreibt $f : M \rightarrow N$.

Statt $(x, y) \in f$ schreibt man $f(x) = y$; dies ist durch (1.3) gerechtfertigt, denn es gibt zu jedem $x \in M$ nur ein $y \in N$ mit $(x, y) \in f$, und dieses eindeutig bestimmte y kann man dann mit $f(x)$ bezeichnen.

Anschaulich gesprochen nimmt eine Abbildung ein Element von M und ,wirft' es auf ein Element von N , und zwar in eindeutiger Weise: Ein Element von M kann nicht auf mehrere Elemente von N abgebildet werden. Man beachte, daß nicht jeder Wert des Wertebereichs angenommen werden muß.

Wir verwenden die Begriffe *Abbildung* und *Funktion* synonym. Manche Autorinnen reservieren den Begriff der Funktion ausschließlich für Abbildungen, deren Wertebereich die reellen oder komplexen Zahlen sind.

DEFINITION 1.16 (Bild und Urbild). Das *Urbild* von $y \in N$ unter der Abbildung $f : M \rightarrow N$ ist definiert als die Menge

$$f^{-1}(y) := \{x \in M : f(x) = y\}.$$

Das Urbild einer Teilmenge $W \subset N$ unter f ist definiert als

$$f^{-1}(W) := \{x \in M : f(x) \in W\}.$$

Das *Bild* einer Teilmenge $U \subset M$ unter f ist definiert als

$$f(U) := \{f(x) : x \in U\} \subset N.$$

Man beachte, daß Urbilder auch leer sein können, Bilder hingegen nicht (außer wenn $M = \emptyset$).

DEFINITION 1.17. Eine Abbildung $f : M \rightarrow N$ heißt

- *injektiv*, falls

$$\forall x, y \in M : f(x) = f(y) \Rightarrow x = y;$$
- *surjektiv*, falls

$$\forall y \in N \quad \exists x \in M : f(x) = y;$$
- *bijektiv*, wenn sie injektiv und surjektiv ist.

BEISPIEL 1.18. Betrachte die Menge $M = N = \{1, 2, 3\}$ und die Abbildungen $f, g : M \rightarrow M$, gegeben durch

$$\begin{aligned} f(1) &= 1, & f(2) &= 3, & f(3) &= 1; \\ g(1) &= 3, & g(2) &= 2, & g(3) &= 1. \end{aligned}$$

Dann ist f weder injektiv (weil $f(1) = f(3)$) noch surjektiv (weil es kein $x \in M$ gibt mit $f(x) = 2$), und erst recht nicht bijektiv. Dagegen ist g bijektiv, denn es ist injektiv (ordnet also verschiedenen Elementen nie den gleichen Wert zu) und surjektiv (jedes Element im Wertebereich wird getroffen).

Sei nun weiterhin $M = \{1, 2, 3\}$, aber $N = \{1, 2\}$. Dann existiert offenbar keine Injektion $M \rightarrow N$, denn wenn drei Elemente auf zwei abgebildet werden sollen, dann muß mindestens ein Wert zweimal angenommen werden (‘Schubfachprinzip’). Die Abbildung $h : M \rightarrow N$, die durch $h(1) = 1$, $h(2) = 2$, $h(3) = 1$ gegeben ist, ist surjektiv. Dagegen ist die Abbildung j mit $j(1) = j(2) = j(3) = 1$ nicht surjektiv.

Betrachte schließlich nach wie vor $M = \{1, 2, 3\}$, aber nun $N = \{1, 2, 3, 4\}$. Es existiert keine surjektive Abbildung $M \rightarrow N$, denn die drei Elemente des Definitionsbereichs können auf höchstens drei verschiedene Werte abgebildet werden, sodaß stets mindestens ein Element von N nicht getroffen wird. Eine Injektion wäre z.B. gegeben durch $k(1) = 1, k(2) = 2, k(3) = 4$.

Wir beobachten an diesem Beispiel, daß zwischen zwei *endlichen Mengen* nur dann eine Bijektion bestehen kann, wenn sie die gleiche Anzahl von Elementen haben.

BEISPIEL 1.19. Sei $\mathbb{N}$ die Menge der natürlichen Zahlen (ohne null) und sei $f : \mathbb{N} \rightarrow \mathbb{N} \setminus \{1\}$ definiert durch $f(n) = n+1$ (man schreibt auch $n \mapsto n+1$). Dann ist f injektiv, denn aus $n+1 = m+1$ folgt (nach Subtraktion von 1 auf beiden Seiten) $m = n$. Die Abbildung ist auch surjektiv, denn für jedes $n \in \mathbb{N} \setminus \{1\}$ gilt $n-1 \in \mathbb{N}$, und $f(n-1) = n$. Es ist also f eine Bijektion zwischen $\mathbb{N}$ und $\mathbb{N} \setminus \{1\}$. Es kann also Bijektionen zwischen einer Menge und einer echten Teilmenge ihrer selbst geben, aber nach obiger Beobachtung nur dann, wenn diese Menge nicht endlich ist.¹⁰

BEISPIEL 1.20. Betrachte $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$. Diese Abbildung ist nicht injektiv, denn für jedes $x \in \mathbb{R}$ ist $x^2 = (-x)^2$. Schränkt man den Definitionsbereich allerdings auf $\mathbb{R}_0^+$ ein (die Menge der nichtnegativen reellen Zahlen), so ist die Abbildung $f \upharpoonright_{\mathbb{R}_0^+}$ (so bezeichnet man die Einschränkung von f auf $\mathbb{R}_0^+$) injektiv. Wählt man als Wertebereich $\mathbb{R}_0^+$ statt $\mathbb{R}$, so erhält man sogar eine surjektive Abbildung; die Abbildung $\mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+, x \mapsto x^2$ ist also bijektiv.

Eine besondere Bijektion zwischen einer Menge M und sich selbst ist die *Identität* $\iota : M \rightarrow M, x \mapsto x$.

DEFINITION 1.21 (Komposition). Seien $f : M \rightarrow N$ und $g : N \rightarrow P$ zwei Abbildungen. Dann heißt die Abbildung

$$g \circ f : M \rightarrow P, \quad M \ni x \mapsto g(f(x)) \in P$$

die *Komposition* (oder *Verknüpfung*) der Abbildungen f und g .

Wir geben zwei Beispiele: Seien $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$ und $g : \mathbb{R} \rightarrow \mathbb{R}, y \mapsto \sin(y)$ gegeben, so ist $f \circ g : \mathbb{R} \rightarrow \mathbb{R}, y \mapsto \sin(y)^2$ und $g \circ f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto \sin(x^2)$. Es handelt sich selbstverständlich um zwei verschiedene Funktionen.

¹⁰Man kann aus diesen Gedanken sogar die folgende Definition motivieren: Eine Menge M heißt *unendlich*, wenn es eine Bijektion zwischen M und einer echten Teilmenge von M gibt. Eine Menge heißt *endlich*, wenn sie nicht unendlich ist.

Ein weiteres Beispiel: Sei $M = \{1, 2, 3\}$ und definiere eine Abbildung $M \rightarrow M$ durch $f(1) = 1, f(2) = 3, f(3) = 2$. Dann ist $f \circ f = \iota$ (die Identität auf M). Abbildungen mit dieser Eigenschaft bezeichnet man als *selbstinvers* oder als *Involutionen*.¹¹

DEFINITION 1.22 (Umkehrabbildung). Sei $f : M \rightarrow N$ eine Abbildung. Eine Abbildung $g : N \rightarrow M$ heißt *Umkehrabbildung* (oder *Umkehrfunktion* oder *inverse Abbildung/Funktion*) zu f , falls

$$f \circ g = \iota_N \quad \text{und} \quad g \circ f = \iota_M.$$

Hierbei bezeichnen ι_M und ι_N die Identität auf der jeweiligen Menge. Ist g Umkehrabbildung zu f , so schreiben wir auch $g = f^{-1}$.

Achtung: f^{-1} ist *nicht* das Gleiche wie $x \mapsto \frac{1}{f(x)}$: Die Funktion $\iota_{\mathbb{R}} : x \mapsto x$ ist invertierbar mit $\iota_{\mathbb{R}}^{-1} = \iota_{\mathbb{R}}$ (ist also eine Involution), aber für jedes $x \in \mathbb{R}$ ist $\frac{1}{\iota_{\mathbb{R}}(x)} = \frac{1}{x}$.

SATZ 1.23 (Genau die Bijektionen sind invertierbar). Eine Abbildung $f : M \rightarrow N$ besitzt genau dann eine Umkehrabbildung, wenn f bijektiv ist. In diesem Falle ist die Umkehrabbildung eindeutig bestimmt.

BEWEIS. Wir zeigen beide Implikationen. Sei zunächst f bijektiv. Definiere die Funktion $g : N \rightarrow M$ durch

$$g(y) = x \quad \text{falls} \quad f(x) = y.$$

Da f surjektiv ist, existiert zu jedem $y \in N$ ein solches x . Da f injektiv ist, existiert auch *nur* ein solches x . Damit ist g als Funktion $N \rightarrow M$ wohldefiniert.¹²

Nun gilt einerseits für jedes $x \in M$ die Gleichheit $g(f(x)) = g(y) = x$, wobei wir $y := f(x)$ gesetzt haben. Andererseits haben wir für jedes $y \in N$ die Gleichheit $f(g(y)) = f(x) = y$, wobei x das eindeutig bestimmte Element von M ist, für das $f(x) = y$. Es folgt, daß g (eine) Umkehrabbildung zu f ist.

Sei nun umgekehrt g eine Umkehrabbildung zu f . Wir müssen zeigen, daß f bijektiv ist. Sei $y \in N$ beliebig, so ist nach Voraussetzung $f(g(y)) = y$, also ist f surjektiv. Sind außerdem $x, x' \in M$ dergestalt, daß $f(x) = f(x') =: y$, so folgt $g(y) = g(f(x)) = x$ und $g(y) = g(f(x')) = x'$, woraus $x = x'$ folgt. Also ist f auch injektiv. Insgesamt folgt die Bijektivität von f .

Schließlich zeigen wir die Eindeutigkeit der Umkehrabbildung. Seien dazu $g, h : N \rightarrow M$ zwei Umkehrabbildungen zu f , und sei $y \in N$. Nach dem gerade Bewiesenen ist f bijektiv, es existiert also genau ein $x \in M$ mit $f(x) = y$. Nach Definition der Umkehrfunktion gilt aber

$$g(y) = g(f(x)) = x = h(f(x)) = h(y),$$

und da $y \in N$ beliebig gewählt war, folgt¹³ $g = h$. □

Umkehrfunktionen kennen Sie aus der Schule: So ist beispielsweise der Logarithmus $\log : \mathbb{R}^+ \rightarrow \mathbb{R}$ die Umkehrfunktion der Exponentialfunktion $\exp : \mathbb{R} \rightarrow \mathbb{R}^+$, oder die Quadratwurzel $\sqrt{\cdot} : \mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$ die Umkehrfunktion der Quadratfunktion $\mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$. Man beachte jeweils die Wahl der richtigen Definitions- und Wertebereiche: So ist etwa $\exp$ bijektiv (und daher invertierbar) als Funktion $\mathbb{R} \rightarrow \mathbb{R}^+$, nicht aber als Funktion $\mathbb{R} \rightarrow \mathbb{R}$.

¹¹Die Abbildung $f \circ f$ wird oft auch mit f^2 bezeichnet. Man unterscheide diese Funktion scharf von der Funktion $x \mapsto f(x)^2$, die also die Werte von f quadriert.

¹²Der Terminus ‚wohldefiniert‘ kann in der Mathematik alles mögliche bedeuten; hier bedeutet er, daß durch die angegebene Vorschrift tatsächlich eine Funktion definiert wird, daß g also jedem $y \in N$ genau einen Wert $x \in M$ zuordnet.

¹³Hier meint $g = h$, daß g und h als Funktionen gleich sind; wir haben bei der letzten Schlußfolgerung implizit verwendet, daß zwei Funktionen $g, h : N \rightarrow M$ gleich sind genau dann, wenn gilt: $\forall y \in N : g(y) = h(y)$. Man überzeuge sich davon, daß dies aus den Definitionen 1.15 und 1.5 folgt.

1.2.3. Familien und unendliche Mengenoperationen.

1.2.3.1. *Definition und Beispiele.* Der Begriff ‚Familie‘ ist (in der Mathematik) ein bloßes Synonym für ‚Abbildung‘. Sein häufiger Gebrauch in bestimmten Kontexten rechtfertigt aber folgende Definition:

DEFINITION 1.24 (Familie). Seien I (die sogenannte *Indexmenge*) und M Mengen, so bezeichnet man jede Abbildung $I \rightarrow M$ als eine *durch I indizierte Familie von Elementen von M* . Diese Abbildung $j \mapsto x_j$ schreibt man als $(x_j)_{j \in I}$.

Als Indexmenge verwendet man meist die natürlichen oder reellen Zahlen oder Teilmengen davon.

BEISPIEL 1.25 (Folgen). Ist I die Menge der natürlichen Zahlen, so wird eine Familie $(x_n)_{n \in \mathbb{N}}$ als *Folge* in M bezeichnet. So ist z.B.

$$\left(\frac{1}{n}\right)_{n \in \mathbb{N}} = \left(1, \frac{1}{2}, \frac{1}{3}, \frac{1}{4}, \dots\right)$$

eine Folge rationaler Zahlen.

BEISPIEL 1.26 (Tupel). Für ein gegebenes $N \in \mathbb{N}$ sei $I = \{1, 2, \dots, N\}$. Eine derart indizierte Familie von Elementen einer Menge M heißt N -Tupel. Zum Beispiel ist $(-1, 0, 4)$ ein Tripel (also 3-Tupel) natürlicher Zahlen. Die Menge aller N -Tupel von Elementen aus M bezeichnet man mit M^N . Die 2-Tupel nennt man (geordnete) Paare und setzt sie somit mit den Objekten von Definition 1.9 gleich. Diese Identifizierung ist dadurch gerechtfertigt, daß die geordneten Paare (nach Definition 1.9) und die 2-Tupel zwar nicht das Gleiche *sind*, aber für alle praktischen Zwecke das Gleiche *tun* (beide Konzepte erfüllen nämlich das Paaraxiom 1.10).¹⁴

BEISPIEL 1.27. Hier ein Beispiel, das Sie ggf. erst in einigen Semestern im Detail verstehen werden: Ein *stochastischer Prozeß* ist, grob gesprochen, eine Familie von Zufallsvariablen. Ein durch $\mathbb{N}$ indizierter stochastischer Prozeß $(X_n)_{n \in \mathbb{N}}$ könnte z.B. den Wert X_n einer Immobilie n Jahre in der Zukunft beschreiben. Ist die Indexmenge $I = \mathbb{R}^+$, spricht man von stochastischen Prozessen *in stetiger Zeit*, etwa wenn X_t den Wert einer Aktie bezeichnet, die zu jedem Zeitpunkt (und nicht nur täglich oder stündlich) gehandelt werden kann.

1.2.3.2. *Unendliche Mengenoperationen.* Sei I eine Indexmenge und $(M_j)_{j \in I}$ eine Familie von Mengen (d.h. jedes M_j ist eine Menge, und in Definition 1.24 würde man $M = \{M_j : j \in I\}$ wählen.)

Wir definieren

$$\bigcup_{j \in I} M_j := \{x : \exists j \in I : x \in M_j\}, \quad \bigcap_{j \in I} M_j := \{x : \forall j \in I : x \in M_j\}.$$

Ist etwa $I = \mathbb{N}$ und $M_j = \left(\frac{1}{j}, 1\right] \subset \mathbb{R}$, so gilt $\bigcup_{j \in \mathbb{N}} M_j = (0, 1]$ und $\bigcap_{j \in \mathbb{N}} M_j = \emptyset$ (da $M_1 = (1, 1] = \emptyset$). Ein weiteres Beispiel: $\mathbb{R} = \bigcup_{x \in \mathbb{R}} \{x\}$ (hier wurde $I = \mathbb{R}$ und $M_x = \{x\}$ gewählt).

Zum Abschluß geben wir (für Ihre mathematische Allgemeinbildung) das *Auswahlaxiom* an. Es besagt: Sei $(M_j)_{j \in I}$ eine Familie von nichtleeren Mengen mit einer Indexmenge I , dann existiert eine ‚Auswahlfunktion‘ $f : I \rightarrow \bigcup_{j \in I} M_j$ mit der Eigenschaft

$$\forall j \in I : f(j) \in M_j.$$

Die Funktion wählt also aus jeder der Mengen M_j ein Element aus. Das Auswahlaxiom scheint auf den ersten Blick sehr plausibel und ist weithin als eine der Grundlagen der modernen Mengenlehre akzeptiert. Dennoch führt es bei sehr ‚großen‘ Indexmengen zu

¹⁴Siehe [11] für eine ausführliche Diskussion der These: „Es kommt nicht darauf an, was mathematische Objekte *sind*, sondern was sie *tun*.“

umstrittenen Konsequenzen (wie dem BANACH-TARSKI-Paradoxon, der Existenz nicht-meßbarer Teilmengen von $\mathbb{R}$ oder der Existenz einer algebraischen Basis für jeden Vektorraum) und wird deshalb von manchen Mathematiker*innen wenn nicht rundheraus abgelehnt, so doch wenn möglich vermieden oder nur mit schlechtem Gewissen verwendet. Als Faustregel gilt allerdings: Alles, was irgendwie von praktischer Bedeutung ist, kann auch ohne das Auswahlaxiom hergeleitet werden.

1.3. Rationale Zahlen

1.3.1. Natürliche Zahlen.

1.3.1.1. PEANO-Axiome. Alle Eigenschaften der natürlichen Zahlen folgen aus den folgenden Axiomen¹⁵, die auf PEANO zurückgehen:

Die *natürlichen Zahlen* sind eine Menge $\mathbb{N}$ zusammen mit einer Abbildung $S : \mathbb{N} \rightarrow \mathbb{N}$ und einem Element $1 \in \mathbb{N}$, sodaß gilt:

- (1) $1 \notin S(\mathbb{N})$;
- (2) S ist injektiv;
- (3) Für jede Teilmenge $N \subset \mathbb{N}$ gilt:

$$1 \in N \wedge (\forall n \in N : S(n) \in N) \Rightarrow N = \mathbb{N}.$$

Man kann sich diese Definition folgendermaßen vorstellen: $\mathbb{N}$ ist die Menge $\{1, 2, 3, 4, \dots\}$ und S die *Nachfolgerfunktion*, die jeder natürlichen Zahl ihren Nachfolger zuordnet (d.h. die Zahl mit 1 addiert). Mit der Konvention, daß 0 keine natürliche Zahl ist¹⁶, ist dann die 1 nicht Nachfolger irgendeiner natürlichen Zahl. Die Nachfolgerfunktion ist injektiv, denn aus $n + 1 = m + 1$ folgt $m = n$. Das letzte Axiom¹⁷ heißt *Induktionsaxiom* und bedeutet: Enthält eine Menge natürlicher Zahlen die 1 und mit jedem Element auch dessen Nachfolger, so enthält die Menge *alle* natürlichen Zahlen. Das Induktionsaxiom formalisiert also den Zählvorgang, wie wir ihn seit früher Kindheit kennen: Mit der 1 beginnend gehen wir zum Nachfolger 2 über, von dort zum Nachfolger 3, usw. Gemäß dem Induktionsaxiom ‚erwischen‘ wir auf diese Weise alle natürlichen Zahlen.

Aus abstrakt-mathematischer Sicht stellt sich die Frage, ob es – über das uns vertraute Konzept hinaus – mehrere unterschiedliche Strukturen geben könnte, die gleichermaßen die PEANO-Axiome erfüllen. DEDEKIND hat gezeigt, daß dies nicht der Fall ist, da alle Modelle, die die Axiome erfüllen, in einem bestimmten Sinne zueinander *isomorph* sind¹⁸. Doch selbst wenn es ‚verschiedene‘ Mengen natürlicher Zahlen gäbe, wäre dies für uns ohne Belang: Wir werden alle Eigenschaften der natürlichen Zahlen aus den PEANO-Axiomen herleiten; etwaige von diesen Axiomen unabhängige Eigenschaften, die zwei Mengen natürlicher Zahlen voneinander unterscheiden würden, wären daher nicht von Interesse.

Neben dieser Eindeutigkeitsfrage kann man auch die Existenzfrage stellen: Gibt es überhaupt ein Modell der PEANO-Axiome, also eine Struktur $(\mathbb{N}, S, 1)$, die die Axiome erfüllt? Wir haben von natürlichen Zahlen eine Anschauung und würden diese Frage daher

¹⁵Ein *Axiom* ist eine grundlegende Aussage, auf der eine mathematische Theorie aufbaut und die definitorisch gesetzt wird. Verschiedene Axiomensysteme können zu unterschiedlichen ‚Mathematiken‘ führen, z.B. zu euklidischer oder nichteuklidischer Geometrie.

¹⁶Man kann ebensogut die Null zu den natürlichen Zahlen dazunehmen – das ist reine Geschmackssache.

¹⁷Strenggenommen handelt es sich um ein *Axiomenschema*, da man für jede Wahl der Teilmenge N eine andere Aussage erhält. Quantifizierung über die Teilmengen von $\mathbb{N}$ ist in der Prädikatenlogik 1. Stufe nicht möglich.

¹⁸Selbstverständlich können die natürlichen Zahlen umbenannt werden, indem man sie z.B. in römischen Ziffern notiert oder in verschiedenen Sprachen zählt etc. Der genannte Satz von DEDEKIND sagt gewissermaßen aus, daß es nur solche mathematisch irrelevanten Varianten von natürlichen Zahlen gibt.

selbstverständlich bejahen. Das Paradigma, daß alle mathematischen Objekte auf Mengen zurückzuführen sein sollen, verlangt aber, die Elemente von $\mathbb{N}$ ihrerseits als Mengen anzugeben. Die folgende Möglichkeit dazu wurde von ZERMELO vorgeschlagen:

Die Menge $\mathbb{N}$ sei die Menge aller Mengen der Gestalt $\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \{\{\{\emptyset\}\}\}$, etc. Als Eins wird die leere Menge verwendet: $1 := \emptyset$; Die Abbildung S bildet ein Element $n \in \mathbb{N}$ auf die Menge $\{n\} \in \mathbb{N}$ ab. Man kann dann zeigen, daß diese Struktur in der Tat die PEANO-Axiome erfüllt, und hat demzufolge

$$\begin{aligned} 1 &= \{\} \\ 2 &= \{\{\}\} \\ 3 &= \{\{\{\}\}\} \\ &\text{etc.} \end{aligned}$$

1.3.1.2. Vollständige Induktion. Das Induktionsaxiom gibt eine nützliche Beweismethode für Aussagen über natürliche Zahlen vor. Ist P ein (einstelliges) Prädikat, das über $\mathbb{N}$ interpretiert wird, so geht man zum Beweis der Aussage $\forall n \in \mathbb{N} : P(n)$ folgendermaßen vor:

Induktionsanfang. Zeige $P(1)$.

Induktionsschritt. Zeige: $\forall n \in \mathbb{N} : P(n) \Rightarrow P(n+1)$.

Anwendung des Induktionsaxioms auf die Teilmenge $N := \{n \in \mathbb{N} : P(n)\}$ ergibt dann sofort die gewünschte Aussage $\forall n \in \mathbb{N} : P(n)$.

BEISPIEL 1.28 (GAUSSsche Summenformel). Wir zeigen: Für jedes $n \in \mathbb{N}$ gilt

$$1 + 2 + 3 + \dots + n = \frac{n(n+1)}{2}.$$

BEWEIS. *Induktionsanfang:* Für $n = 1$ ist die Behauptung wahr, denn $1 = \frac{1 \cdot 2}{2}$.

Induktionsschritt: Sei $n \in \mathbb{N}$ eine natürliche Zahl. Wir nehmen an, die Behauptung gälte für n (Induktionsannahme), und folgern daraus die Behauptung für $n+1$. In der Tat,

$$\begin{aligned} 1 + 2 + 3 + \dots + n + (n+1) &= \frac{n(n+1)}{2} + n+1 \\ &= \frac{n(n+1)}{2} + \frac{2n+2}{2} \\ &= \frac{n^2 + n + 2n + 2}{2} = \frac{(n+1)(n+2)}{2}, \end{aligned}$$

wobei wir bereits im ersten Schritt die Induktionsannahme verwendet haben. Das Ergebnis ist genau die gewünschte Summenformel, in die $n+1$ statt n eingesetzt wurde. $\square$

BEISPIEL 1.29. Wir zeigen: Für alle $n \in \mathbb{N}$ gilt

$$1 + 3 + 5 + \dots + (2n-1) = n^2.$$

BEWEIS. *Induktionsanfang:* Für $n = 1$ gilt die Behauptung wegen $1 = 1^2$.

Induktionsschritt: Für ein $n \in \mathbb{N}$ sei die Formel korrekt. Dann gilt

$$\begin{aligned} 1 + 3 + 5 + \dots + (2n-1) + (2(n+1)-1) &= n^2 + (2(n+1)-1) \\ &= n^2 + 2n + 1 = (n+1)^2. \end{aligned}$$

$\square$

BEISPIEL 1.30 (Mächtigkeit der Potenzmenge). Sei $n \in \mathbb{N} \cup \{0\}$ und M eine Menge mit n Elementen. Dann hat die Potenzmenge $\mathcal{P}(M)$ 2^n Elemente¹⁹.

¹⁹Deshalb wird die Potenzmenge von M manchmal auch mit 2^M bezeichnet.

BEWEIS. *Induktionsanfang:* Hier fangen wir mit $n = 0$ an: Die einzige Menge mit null Elementen ist die leere Menge, und deren Potenzmenge $\{\emptyset\}$ enthält $1 = 2^0$ Elemente.

Induktionsschritt: Sei $n \in \mathbb{N} \cup \{0\}$ und M eine $(n+1)$ -elementige Menge. Bezeichne ein beliebiges der Elemente von M mit x_{n+1} , dann ist $M = M' \cup \{x_{n+1}\}$, wobei M' eine n -elementige Menge ist. Die Teilmengen von M können unterschieden werden in diejenigen, die x_{n+1} enthalten, und diejenigen, die x_{n+1} nicht enthalten.

Die ersteren Teilmengen sind von der Form $V \cup \{x_{n+1}\}$, wobei V eine Teilmenge von M' ist; davon gibt es nach Induktionsannahme genau 2^n Stück. Die letzteren Teilmengen von M sind Teilmengen von M' , wovon es nach Induktionsannahme ebenfalls 2^n Stück gibt.

Insgesamt hat M also $2^n + 2^n = 2^{n+1}$ Teilmengen. $\square$

1.3.1.3. *Arithmetik.* Auf den natürlichen Zahlen sind zwei binäre Operationen $+$ und $\cdot$ definiert, die wir im Beispielmateriale im Laufe der Vorlesung bereits verwendet haben. Wir holen nun die zugehörige Theorie nach.

DEFINITION 1.31 (Addition und Multiplikation in $\mathbb{N}$).

- (1) Die Abbildung $+: \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, $(n, m) \mapsto n + m$, ist wie folgt definiert: Sei $n \in \mathbb{N}$, dann ist
 - (a) $n + 1 := S(n)$;
 - (b) $n + S(m) := S(n + m)$, wenn $n + m$ bereits definiert ist.
- (2) Die Abbildung $\cdot: \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, $(n, m) \mapsto nm$, ist wie folgt definiert: Sei $n \in \mathbb{N}$, dann ist
 - (a) $n \cdot 1 := n$;
 - (b) $nS(m) := nm + n$, wenn nm bereits definiert ist.

Dies ist ein Beispiel einer *rekursiven* Definition: Man definiert eine Abbildung erst für $m = 1$ und dann, unter der Annahme, die Abbildung sei bereits für m erklärt, für den Nachfolger $S(m)$. Nach Induktionsaxiom ist dadurch die Abbildung für alle $m \in \mathbb{N}$ definiert.

SATZ 1.32 (Rechenregeln). Seien $n, m, k \in \mathbb{N}$. Dann gilt

- (1) $(n + m) + k = n + (m + k)$ (*Assoziativgesetz der Addition*)
- (2) $n + m = m + n$ (*Kommutativgesetz der Addition*)
- (3) $(nm)k = n(mk)$ (*Assoziativgesetz der Multiplikation*)
- (4) $nm = mn$ (*Kommutativgesetz der Multiplikation*)
- (5) $n(m + k) = nm + nk$ (*Distributivgesetz*)

BEWEIS. Wir zeigen exemplarisch (1) und (2). Die Beweise der anderen Regeln sind nachdrücklich zur Übung empfohlen.

Beweis von (1). Seien $m, n \in \mathbb{N}$ beliebig. Wir führen Induktion nach k .

Für $k = 1$ haben wir

$$(n + m) + 1 = S(n + m) = n + S(m) = n + (m + 1),$$

was den Induktionsanfang etabliert. Hier haben wir lediglich die Definition der Addition verwendet.

Für den Induktionsschritt gelte (1) für ein $k \in \mathbb{N}$. Dann ist

$$(n + m) + S(k) = S((n + m) + k) = S(n + (m + k)) = n + S(m + k) = n + (m + S(k)),$$

wobei wir im zweiten Schritt die Induktionsannahme und ansonsten jeweils die Definition der Addition benutzt haben.

Beweis von (2). Der Beweis schachtelt zwei Induktionsargumente (nach n und nach m) ineinander. Wir zeigen zunächst durch vollständige Induktion: $\forall n \in \mathbb{N} : n + 1 = 1 + n$.

Für $n = 1$ ist dies klar, denn $1 + 1 = 1 + 1$. Es gelte nun für ein $n \in \mathbb{N}$ $n + 1 = 1 + n$. Dann ist

$$S(n) + 1 = S(S(n)) = S(n + 1) = S(1 + n) = (1 + n) + 1 = 1 + (n + 1) = 1 + S(n).$$

Hier haben wir erst die Definition der Addition, in der dritten Gleichheit die Induktionsannahme, dann wieder die Definition der Addition und im vorletzten Schritt das Assoziativgesetz verwendet.

Nun zeigen wir für festes $n \in \mathbb{N}$: $\forall m \in \mathbb{N} : n + m = m + n$ durch Induktion nach m . Den Induktionsanfang ($m = 1$) haben wir soeben verifiziert.

Unter der Annahme $n + m = m + n$ für ein $m \in \mathbb{N}$ gilt nun

$$\begin{aligned} n + S(m) &= n + (m + 1) = (n + m) + 1 = (m + n) + 1 = 1 + (m + n) \\ &= (1 + m) + n = (m + 1) + n = S(m) + n, \end{aligned}$$

wobei wir wieder die Definition der Addition, das Assoziativgesetz, die Induktionsannahme, und den Induktionsanfang benutzt haben. $\square$

Die Assoziativgesetze garantieren die Sinnhaftigkeit der Ausdrücke $n + m + k$ bzw. nmk ohne Angabe einer bestimmten Klammerung. Allgemeiner kann man aus den Assoziativ- und Kommutativgesetzen folgern, daß man bei Summen mehrerer natürlicher Zahlen Klammern und Reihenfolge beliebig umstellen kann, z.B.

$$(n_1 + n_2) + ((n_3 + n_4) + n_5) = (n_3 + (n_1 + n_4)) + (n_2 + n_5),$$

und ebenso für die Multiplikation. Durch mehrfache Anwendung dieser Gesetze sowie des Distributivgesetzes („Ausmultiplizieren“) gelangt man außerdem zu Identitäten wie

$$(a + b)(c + d) = ac + bc + ad + bd,$$

oder zu den binomischen Formeln $(a + b)^2 = a^2 + 2ab + b^2$ usw.

Zum Abschluß führen wir noch eine übersichtlichere Notation von Summen bzw. Produkten mehrerer Zahlen ein:

DEFINITION 1.33 (Summen- und Produktzeichen). Sei $(a_k)_{k \in \mathbb{N}}$ eine Folge natürlicher Zahlen. Dann definieren wir rekursiv

$$(1) \quad \sum_{k=1}^1 a_k := a_1 \text{ sowie} \\ \sum_{k=1}^{m+1} a_k := \sum_{k=1}^m a_k + a_{m+1} \quad \text{für jedes } m \in \mathbb{N};$$

$$(2) \quad \prod_{k=1}^1 a_k := a_1 \text{ sowie} \\ \prod_{k=1}^{m+1} a_k := \left(\prod_{k=1}^m a_k \right) a_{m+1} \quad \text{für jedes } m \in \mathbb{N}.$$

So könnte man z.B. die GAUSSsche Formel aus Beispiel 1.28 als

$$\sum_{k=1}^n k = \frac{n(n+1)}{2}$$

schreiben und die Formel aus Beispiel 1.29 als

$$\sum_{k=1}^n (2k-1) = n^2.$$

Die Summanden bzw. Faktoren a_k können ebenso als ganze, rationale, reelle oder komplexe Zahlen gewählt werden, sobald wir diese Zahlenmengen eingeführt haben. Man verwendet oft die Konventionen

$$\sum_{k=1}^0 a_k := 0, \quad \prod_{k=1}^0 a_k := 1.$$

Außerdem kann man die Summation bzw. Multiplikation auch erst mit einem größeren Index beginnen, z.B.

$$\sum_{k=5}^7 k^2 = 5^2 + 6^2 + 7^2.$$

1.3.2. Ganze und rationale Zahlen.

1.3.2.1. *Ganze Zahlen.* Gegeben zwei natürliche Zahlen $n, m \in \mathbb{N}$, existiert dann eine Zahl x , sodaß $n + x = m$? Wenn $n < m$, dann existiert ein solches $x \in \mathbb{N}$, und man schreibt $x = m - n$. Ist allerdings $n \geq m$, so existiert keine Lösung dieser Gleichung in $\mathbb{N}$, und wir müssen unseren Zahlenraum erweitern. Da die Subtraktion noch ‚verboten‘ ist (denn z.B. $3 - 5$ ist keine natürliche Zahl), fassen wir ‚3-5‘ stattdessen als geordnetes Paar $(3, 5)$ auf. Dabei müssen wir allerdings beachten, daß diese Darstellung nicht eindeutig ist: Es soll ja etwa $3 - 5 = 7 - 9 = 998 - 1000 = \dots$ sein. Wir wollen also Paare (m, n) und (m', n') miteinander identifizieren, wenn $m - n = m' - n'$. Zur Umgehung des noch nicht definierten Minuszeichens schreiben wir dies als $m + n' = m' + n$ um. Damit ist eine Äquivalenzrelation auf $\mathbb{N} \times \mathbb{N}$ definiert: Es ist offensichtlich $m + n = m + n$ (Reflexivität), $m + n' = m' + n \Rightarrow m' + n = m + n'$ (Symmetrie), und falls $m + n' = m' + n$ und $m' + n'' = m'' + n'$, so folgt $m + n'' = m'' + n$ (Transitivität). Wir fassen diese Gedanken in der folgenden Definition zusammen:

DEFINITION 1.34 (Ganze Zahlen). Auf $\mathbb{N} \times \mathbb{N}$ sei die Äquivalenzrelation $\equiv$ definiert durch $(m, n) \equiv (m', n')$ genau dann, wenn $m + n' = m' + n$. Die Äquivalenzklassen $[(m, n)]$ bezüglich dieser Relation heißen *ganze Zahlen*. Die Menge der ganzen Zahlen wird mit $\mathbb{Z}$ bezeichnet.

Statt $[(m, n)]$ schreiben wir $m - n$ und setzen $-[(m, n)] := [(n, m)]$ (in anderen Worten: $-(m - n) = n - m$). Ist $n \in \mathbb{N}$, so identifizieren wir n mit der ganzen Zahl $[(n + 1, 1)]$. Außerdem definieren wir $0 := [(1, 1)]$.

DEFINITION 1.35 (Addition, Subtraktion, Multiplikation). Die Abbildungen $+, -, \cdot : \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z}$ sind wie folgt definiert: Für $(m, n), (m', n') \in \mathbb{N} \times \mathbb{N}$ ist

- (1) $[(m, n)] + [(m', n')] := [(m + m', n + n')]$;
- (2) $[(m, n)] - [(m', n')] := [(m, n)] + (-[(m', n')])$
- (3) $[(m, n)] \cdot [(m', n')] := [(mm' + nn', nm' + n'm)]$.

Wem diese Definition zunächst unverständlich erscheint, der möge sich z.B. $(m - n)(m' - n') = mm' + nn' - (nm' + n'm)$ vor Augen führen.

An dieser Stelle ist Vorsicht geboten: Wir definieren ja Rechenoperationen für Äquivalenzklassen mittels ihrer *Repräsentanten*, d.h. statt m und n zur Darstellung der ganzen Zahl $[(m, n)]$ hätten wir ebenso $m + 1$ und $n + 1$ oder $m + 27$ und $n + 27$ wählen können. Wir müssen uns also davon überzeugen, daß obige Definition von der Wahl der konkreten Repräsentanten unabhängig ist²⁰. Wir führen dies für (1) vor und überlassen (2) und (3) den Lesenden zur Übung:

Seien $(m_1, n_1) \equiv (m_2, n_2)$ und $(m'_1, n'_1) \equiv (m'_2, n'_2)$, dann müssen wir zeigen $(m_1 + m'_1, n_1 + n'_1) \equiv (m_2 + m'_2, n_2 + n'_2)$.

Nach Definition der Äquivalenzrelation ist $m_1 + n_2 = m_2 + n_1$ sowie $m'_1 + n'_2 = m'_2 + n'_1$; Addition beider Gleichungen ergibt dann bereits

$$m_1 + n_2 + m'_1 + n'_2 = m_2 + n_1 + m'_2 + n'_1,$$

also die Behauptung.

PROPOSITION 1.36 (Rechenregeln). *Addition und Multiplikation sind auf $\mathbb{Z}$ jeweils assoziativ und kommutativ und erfüllen das Distributivgesetz. Für jedes $z \in \mathbb{Z}$ gilt $-z = (-1) \cdot z$.*

²⁰Das Gleiche gilt für die Definition $-[(m, n)] := [(n, m)]$ weiter oben.

BEWEIS. Wir zeigen nur exemplarisch das Kommutativgesetz der Addition und die letzte Aussage.

Zum Kommutativgesetz: Seien $[(m, n)], [(m', n')] \in \mathbb{Z}$, so gilt

$$[(m, n)] + [(m', n')] = [(m + m', n + n')] = [(m' + m, n' + n)] = [(m', n')] + [(m, n)],$$

wobei wir die Definition der Addition und die Kommutativität der Addition natürlicher Zahlen verwendet haben.

Zur letzten Aussage: Sei $z = [(m, n)]$, also $-z = [(n, m)]$. Es ist ja $-1 = -[(2, 1)] = [(1, 2)]$ und daher

$$(-1) \cdot z = [(1, 2)] \cdot [(m, n)] = [(1 \cdot m + 2n, 1 \cdot n + 2m)] = [(m + 2n, n + 2m)] = [(n, m)],$$

da $(m + 2n, n + 2m) \equiv (n, m)$: In der Tat, $m + 2n + m = n + 2m + n$.

□

Aus den Definitionen und dieser Proposition folgen sofort alle bekannten Rechenregeln für ganze Zahlen, z.B. $-(a + b) = -a - b$ oder auch $a - a = 0$ und $a \cdot 0 = 0$. Außerdem ist $\mathbb{Z}$ *nullteilerfrei*: Aus $ab = 0$ folgt $a = 0$ oder $b = 0$. Dies sieht man wie folgt ein: Ist $a = [(m, n)]$ und $b = [(m', n')]$, so bedeutet $ab = 0$: $mm' + nn' = nm' + n'm$ bzw. äquivalent dazu $m'(m - n) = n'(m - n)$. Ist $m - n = 0$, sind wir fertig, denn dann ist $a = 0$. Ist dagegen $m - n \neq 0$, so ist entweder $m - n \in \mathbb{N}$ oder $n - m \in \mathbb{N}$. In beiden Fällen folgt $km' = kn'$ für ein $k \in \mathbb{N}$, und daraus folgt $m' = n'$, wie man durch vollständige Induktion über k sieht. Das aber bedeutet $b = 0$.

Schließlich bemerken wir, daß die Definition der Addition und Multiplikation in $\mathbb{Z}$ mit der in $\mathbb{N}$ konsistent ist, das bedeutet: Werden $n, m \in \mathbb{N}$ als ganze Zahlen $[(n + 1, 1)]$ bzw. $[(m + 1, 1)]$ interpretiert, so ist deren Summe (als ganze Zahlen) gleich $[(n + m + 1, 1)]$, was ja wiederum der natürlichen Zahl $n + m$ entspricht. Analog gilt dies für die Multiplikation.

1.3.2.2. *Rationale Zahlen.* Von den natürlichen zu den ganzen Zahlen sind wir über die Gleichung $n + x = m$ gelangt, die in $\mathbb{N}$ nicht immer eine Lösung hat, in $\mathbb{Z}$ dagegen schon (nämlich $x = m - n$). Ersetzen wir diese Gleichung durch $bx = a$, wobei nun $a, b \in \mathbb{Z}$, erhalten wir in ganz ähnlicher Weise die rationalen Zahlen:

DEFINITION 1.37 (Rationale Zahlen). Auf $\mathbb{Z} \times (\mathbb{Z} \setminus \{0\})$ definieren wir die Äquivalenzrelation: $(a, b) \equiv (a', b')$ genau dann, wenn $ab' = a'b$. Die Äquivalenzklassen $[(a, b)]$ heißen *rationale Zahlen*, die Menge der rationalen Zahlen wird mit $\mathbb{Q}$ bezeichnet.

Die genannte Relation ist in der Tat eine Äquivalenzrelation: Reflexivität und Symmetrie sind klar. Zur Transitivität: Seien $(a, b) \equiv (a', b')$ und $(a', b') \equiv (a'', b'')$, so ist

$$ab' = a'b \quad \text{und} \quad a'b'' = a''b'. \quad (1.4)$$

Multiplikation beider Gleichungen ergibt $ab'a'b'' = a'ba''b'$ bzw., mit $k := a'b' \in \mathbb{Z}$, $k(ab'' - a''b) = 0$. Da $\mathbb{Z}$ nullteilerfrei ist, folgt $k = 0$ oder $ab'' - a''b = 0$. Im letzteren Falle sind wir fertig, da dann $(a, b) \equiv (a'', b'')$. Ist dagegen $k = 0$, so folgt erneut aus der Nullteilerfreiheit und aus $b' \neq 0$ (denn die Relation ist nur auf $\mathbb{Z} \times (\mathbb{Z} \setminus \{0\})$ definiert), daß $a' = 0$. Aus (1.4) und erneut aus $b' \neq 0$ folgt sodann $a = a'' = 0$ und damit schließlich $(a, b) \equiv (a'', b'')$.

Statt $[(a, b)]$ schreibt man wie gewohnt $\frac{a}{b}$. Die Äquivalenzrelation kann man dann als Gleichheit von Brüchen interpretieren, die man durch Erweitern bzw. Kürzen erhält: Z.B. ist $\frac{1}{2} = \frac{3}{6} = \frac{-2}{-4}$ etc. Man identifiziert eine ganze Zahl $z \in \mathbb{Z}$ mit der rationalen Zahl $[(z, 1)]$. Insbesondere gilt $1 = [(1, 1)]$ und $0 = [(0, 1)]$.

DEFINITION 1.38 (Grundrechenarten). Seien $[(a, b)], [(a', b')] \in \mathbb{Q}$. Dann definieren wir

- (1) $[(a, b)] + [(a', b')] := [(ab' + a'b, bb')];$
- (2) $[(a, b)] - [(a', b')] := [(a, b)] + [(-a', b')];$
- (3) $[(a, b)] \cdot [(a', b')] := [(aa', bb')];$
- (4) $[(a, b)] / [(a', b')] := [(ab', ba')] \text{ falls } a' \neq 0.$

Wie bei den ganzen Zahlen muß man auch hier überprüfen, daß diese Definitionen unabhängig von der Wahl der Repräsentanten sind. Wir führen dies exemplarisch für (3) vor und überlassen den Rest den Lesenden zur Übung. Seien also $(a, b) \equiv (c, d)$ und $(a', b') \equiv (c', d')$, so ist zu zeigen: $(aa', bb') \equiv (cc', dd')$. Nach Voraussetzung gilt

$$ad = bc \quad \text{und} \quad a'd' = b'c',$$

und Multiplikation beider Gleichungen liefert bereits das gewünschte Resultat.

PROPOSITION 1.39 ($\mathbb{Q}$ ist ein Körper). *Addition und Multiplikation auf $\mathbb{Q}$ sind assoziativ und kommutativ und erfüllen das Distributivgesetz. Für alle $x \in \mathbb{Q}$ gilt $x + 0 = x$, $1 \cdot x = x$, $x - x = 0$ und, falls $x \neq 0$, $x \cdot \frac{1}{x} = 1$.*

Man nennt eine Struktur mit diesen Eigenschaften einen *Körper*. Sie werden sich in der (Linearen) Algebra noch eingehend damit befassen.

BEWEIS. Wieder greifen wir nur einige Eigenschaften exemplarisch heraus. Das Distributivgesetz wird etwa wie folgt gezeigt: Seien $[(a, b)], [(c, d)], [(e, f)] \in \mathbb{Q}$, so gilt einerseits

$$[(a, b)] \cdot [(c, d)] + [(a, b)] \cdot [(e, f)] = [(a, b)] \cdot [(cf + de, df)] = [(a(cf + de), bdf)]$$

und andererseits

$$[(a, b)] \cdot [(c, d)] + [(a, b)] \cdot [(e, f)] = [(ac, bd)] + [(ae, bf)] = [(acbf + bdae, bdbf)],$$

und dann ist es nicht schwer zu zeigen $(a(cf + de), bdf) \equiv (acbf + bdae, bdbf)$ (denn auf der rechten Seite kann man b kürzen).

Sei $x = [(a, b)] \in \mathbb{Q}$. Die Eigenschaft $x + 0 = x$ folgt sofort wegen $[(a, b)] + [(0, 1)] = [(a \cdot 1 + 0 \cdot b, b \cdot 1)] = [(a, b)]$. Sei zusätzlich $x \neq 0$, dann gilt

$$x \cdot \frac{1}{x} = [(a, b)] \cdot [(b, a)] = [(ab, ba)] = [(1, 1)] = 1.$$

Hier haben wir die Äquivalenz $(ab, ba) = (1, 1)$ benutzt, d.h. wir haben mit ab gekürzt. $\square$

Die Rechenoperationen auf $\mathbb{Q}$ sind, wie man sich leicht überlegt, mit denen auf $\mathbb{Z}$ konsistent (vgl. die Diskussion oben über die Konsistenz von Addition und Multiplikation in $\mathbb{Z}$ und in $\mathbb{N}$). Auch $\mathbb{Q}$ ist nullteilerfrei: Sind nämlich $[(a, b)], [(c, d)] \in \mathbb{Q}$ mit $[(a, b)] \cdot [(c, d)] = 0$, so folgt $(ac, bd) \equiv (0, 1)$, also $ac = 0$. Aus der Nullteilerfreiheit von $\mathbb{Z}$ folgt $a = 0$ oder $c = 0$ und damit $[(a, b)] = 0$ oder $[(c, d)] = 0$.

1.3.2.3. Die Anordnung der rationalen Zahlen. Wir haben gesehen, daß man eine natürliche Zahl n als ganze Zahl auffassen kann (nämlich als $[(n+1, 1)] \in \mathbb{Z}$). Wir nennen eine ganze Zahl $x \in \mathbb{Z}$ *positiv* und schreiben $x > 0$, wenn x eine natürliche Zahl ist. Wir nennen eine rationale Zahl $x = \frac{a}{b} \in \mathbb{Q}$ *positiv* und schreiben ebenfalls $x > 0$, wenn $a, b \in \mathbb{Z}$ beide positiv oder beide negativ sind. Es ist klar, daß diese Eigenschaft nicht von der Darstellung des Bruches $\frac{a}{b}$ abhängt (z.B. ist $\frac{1}{2} = \frac{-3}{-6}$, aber die Eigenschaft, daß Zähler und Nenner das gleiche Vorzeichen haben, bleibt beim Erweitern bzw. Kürzen erhalten).

Eine rationale Zahl $x \in \mathbb{Q}$, die nicht positiv und nicht null ist, heißt *negativ*, und man schreibt $x < 0$. Ist $x \in \mathbb{Q}$ positiv oder null, schreibt man $x \geq 0$, und ist x negativ oder null, schreibt man $x \leq 0$. Offenbar folgt aus $x > 0$, daß $-x < 0$, und umgekehrt.

Sind zwei Zahlen $x, y \in \mathbb{Q}$ gegeben, so ist x *kleiner als* y , falls $y - x > 0$, und man schreibt $x < y$. Analog sind die Relationen $x \leq y$, $x > y$, $x \geq y$ definiert.

Der folgende Satz besagt, daß $\mathbb{Q}$ ein *dichter archimedisch angeordneter Körper* ist, das bedeutet:

SATZ 1.40.

- (1) Für jedes $x \in \mathbb{Q}$ gilt genau eine der drei Aussagen $x > 0$, $x = 0$ oder $x < 0$.
- (2) Für $x, y \in \mathbb{Q}$ folgt aus $x > 0$ und $y > 0$, daß $x + y > 0$ und daß $xy > 0$.
- (3) Für alle $x, y \in \mathbb{Q}$ mit $x, y > 0$ existiert ein $N \in \mathbb{N}$, sodaß $x < Ny$.

(4) Für alle $x, y \in \mathbb{Q}$ mit $x < y$ existiert ein $z \in \mathbb{Q}$ mit $x < z < y$.

Die ersten beiden Aussagen werden als *Trichotomie* und *Abgeschlossenheit unter Addition und Multiplikation* bezeichnet und sind als *Anordnungsaxiome* bekannt. Die dritte Aussage ist das *archimedische Axiom*. Die letzte Aussage ist die *Dichtheit* von $\mathbb{Q}$.

BEWEIS. Zu (1): Dies folgt unmittelbar aus der Definition der Positivität bzw. Negativität.

Zu (2): Seien $\frac{a}{b}$ und $\frac{c}{d}$ positive rationale Zahlen. Wir dürfen annehmen $a, b, c, d > 0$ (sonst erweitere mit -1). Dann aber gilt

$$\frac{a}{b} + \frac{c}{d} = \frac{ad + bc}{bd} \quad \text{sowie} \quad \frac{a}{b} \cdot \frac{c}{d} = \frac{ac}{bd},$$

und da Summen und Produkte natürlicher Zahlen wieder natürliche Zahlen sind (und insbesondere positiv), folgt die Behauptung.

Zu (3): Seien $\frac{a}{b}$ und $\frac{c}{d}$ positive rationale Zahlen, und wir nehmen wieder an $a, b, c, d > 0$. Die Aussage $\frac{a}{b} < N\frac{c}{d}$ ist nach Multiplikation mit bd (unter Berücksichtigung von (2)) äquivalent zu $ad < Nbc$, sodaß die Behauptung durch folgende Aussage impliziert wird:

Sind $m, n \in \mathbb{N}$, so existiert $N \in \mathbb{N}$ mit $m < Nn$. Dies wollen wir nun durch vollständige Induktion nach m (bei fest gewähltem $n \in \mathbb{N}$) beweisen: Der Induktionsanfang ergibt sich z.B. mit der Wahl $N = 2$. Angenommen nun, für ein $m \in \mathbb{N}$ existiere $N \in \mathbb{N}$, sodaß $m < Nn$; dann ist

$$m + 1 < Nn + 1 \leq Nn + n = (N + 1)n,$$

und (3) folgt.

Zu (4): Wähle einfach $z = \frac{x+y}{2}$. □

Eine einfache Folgerung aus diesem Satz ist: Sind $x, y, z \in \mathbb{Q}$ mit $x < y$ und $z < 0$, so gilt $xz > yz$. Bei Multiplikation einer Ungleichung mit einer negativen Zahl dreht sich die Ungleichheit also um – eine berüchtigte Fehlerquelle bei Schüler*innen (und leider auch bei manchen Studierenden!).

Wir fassen weitere Eigenschaften der Anordnung zusammen:

KOROLLAR 1.41. Seien $x, y, a, b \in \mathbb{Q}$. Dann gilt:

- (1) $x < y \wedge a < b \Rightarrow x + a < y + b$;
- (2) $x < y \wedge a > 0 \Rightarrow ax < ay$;
- (3) $0 \leq x < y \wedge 0 \leq a < b \Rightarrow ax < by$
- (4) $x \neq 0 \Rightarrow x^2 > 0$;
- (5) $x > 0 \Rightarrow \frac{1}{x} > 0$;
- (6) $0 < x < y \Rightarrow 0 < \frac{1}{y} < \frac{1}{x}$.

BEWEIS. Zu (1): Zunächst folgt aus $x < y$ auch $x + a < y + a$ (denn $(y + a) - (x + a) = y - x > 0$). Ebenso folgt aus $a < b$, daß $y + a < y + b$. Zusammen ergibt sich $x + a < y + b$.

Zu (2): Da $y - x > 0$ und $a > 0$, ist auch das Produkt $a(y - x) > 0$.

Zu (3): Falls $a > 0$, so ist $ax < ay$ nach (2); ebenso ist $ay < by$ nach (2). Falls dagegen $a = 0$, so folgt die Behauptung sofort aus $by > 0$.

Zu (4): Ist $x > 0$, so ist auch $x^2 = x \cdot x > 0$; ist $x < 0$, so ist $-x > 0$ (s. Bemerkung vor diesem Korollar), und daher $x^2 = (-x) \cdot (-x) > 0$.

Zu (5): Nach (4) ist $\frac{1}{x^2} > 0$. Multiplikation der Ungleichung $x > 0$ mit $\frac{1}{x^2}$ liefert die Behauptung.

Zu (6): Es ist $xy > 0$, also nach (5) auch $\frac{1}{xy} > 0$. Multiplikation von $x < y$ mit $\frac{1}{xy}$ liefert die Behauptung. □

Die Anordnung erlaubt uns außerdem, den Betrag einer rationalen Zahl zu definieren:

DEFINITION 1.42 (Betrag). Sei $x \in \mathbb{Q}$, dann ist der Betrag von x definiert als

$$|x| := \begin{cases} x & \text{falls } x \geq 0, \\ -x & \text{falls } x < 0. \end{cases}$$

PROPOSITION 1.43. Seien $x, y \in \mathbb{Q}$. Dann gilt

- (1) $|x| \geq 0$, und $|x| = 0$ genau dann, wenn $x = 0$;
- (2) $|xy| = |x||y|$;
- (3) $|x + y| \leq |x| + |y|$.

BEWEIS. (1) und (2) folgen sofort aus der Definition des Betrags. Aussage (3) ist als *Dreiecksungleichung* bekannt und wird wie folgt bewiesen: Aus der Definition folgt²¹ $\pm x \leq |x|$ und $\pm y \leq |y|$; Daher ist

$$|x + y| = \pm(x + y) = \pm x \pm y \leq |x| + |y|.$$

□

Zuletzt definieren wir natürliche Potenzen rationaler Zahlen wie folgt: Für $x \in \mathbb{Q}$ sei $x^0 := 1$ und $x^{n+1} := x^n x$. Man nennt x die *Basis* und n den *Exponenten* der Potenz. Durch vollständige Induktion zeigt man leicht die *Potenzgesetze*: Für alle $x, y \in \mathbb{Q}$ und $m, n \in \mathbb{N}$ gilt

$$x^m x^n = x^{n+m}, \quad (xy)^n = x^n y^n, \quad (x^n)^m = x^{nm}.$$

Setzt man für $x \neq 0$ und $n \in \mathbb{N}$ außerdem $x^{-n} := \frac{1}{x^n}$, so kann man leicht sehen, daß die Potenzgesetze für ganzzahlige Exponenten ihre Gültigkeit behalten.

²¹Hier verwenden wir $\pm x \leq |x|$ als Abkürzung für $x \leq |x| \wedge -x \leq |x|$.

KAPITEL 2

Folgen und Vollständigkeit

Die bisher behandelten Grundlagen waren nicht spezifisch für die Analysis. Letztere wird häufig als die Mathematik der Grenzwerte oder der infinitesimalen Größen beschrieben. Wir führen zuerst mit dem Begriff des Grenzwerts einer Folge *den* zentralen Begriff der Analysis ein. Die Frage, ob bestimmte Folgen (die sogenannten Cauchyfolgen) einen Grenzwert besitzen, führt uns sodann zu den reellen Zahlen, die wir als *Vervollständigung* der rationalen Zahlen erhalten. Ein weiteres wichtiges Beispiel eines vollständigen Körpers sind die komplexen Zahlen, die wir im Anschluß behandeln.

2.1. Konvergenz

DEFINITION 2.1 (Konvergenz einer Folge). Eine Folge $(x_n)_{n \in \mathbb{N}}$ rationaler Zahlen heißt *konvergent* gegen ein $x \in \mathbb{Q}$, wenn für alle $\epsilon > 0$ ein $N \in \mathbb{N}$ existiert, sodaß $|x_n - x| < \epsilon$ für alle $n \geq N$.

In diesem Falle heißt x der *Grenzwert* oder *Limes*¹ der Folge. In Prädikatenlogik kann man schreiben: Eine Folge $(x_n)_{n \in \mathbb{N}}$ konvergiert gegen x , wenn

$$\forall \epsilon > 0 \exists N \in \mathbb{N} \forall n \in \mathbb{N} : n \geq N \Rightarrow |x_n - x| < \epsilon.$$

Man schreibt auch $\lim_{n \rightarrow \infty} x_n = x$ oder $x_n \rightarrow x$ für $n \rightarrow \infty$ ².

BEISPIEL 2.2. (1) Sei $x \in \mathbb{Q}$, so ist die konstante Folge $(x)_{n \in \mathbb{N}} = (x, x, x, \dots)$ konvergent mit Grenzwert x .

(2) Die Folge $(\frac{1}{n})_{n \in \mathbb{N}}$ konvergiert gegen 0: Sei nämlich $\epsilon > 0$, so gibt es nach dem archimedischen Axiom (Satz 1.40 (3)) ein $N \in \mathbb{N}$ mit der Eigenschaft $1 < N\epsilon$, also $\frac{1}{N} < \epsilon$; dies gilt dann (nach Korollar 1.41 (6)) auch für jedes $n \geq N$.

PROPOSITION 2.3 (Eindeutigkeit des Limes). *Der Grenzwert einer konvergenten Folge ist eindeutig bestimmt.*

BEWEIS. Die Folge $(x_n)_{n \in \mathbb{N}}$ konvergiere sowohl gegen $x \in \mathbb{Q}$ als auch gegen $y \in \mathbb{Q}$. Wir müssen zeigen $x = y$. Angenommen, dies wäre nicht der Fall, dann wählen wir $\epsilon := \frac{|x-y|}{2} > 0$ und erhalten ein $N_1 \in \mathbb{N}$, sodaß für alle $n \geq N_1$ gilt $|x_n - x| < \epsilon$. Ebenso gibt es $N_2 \in \mathbb{N}$, sodaß für alle $n \geq N_2$ gilt $|x_n - y| < \epsilon$. Damit ist aber für alle $n \geq \max\{N_1, N_2\}$ nach Dreiecksungleichung

$$|x - y| = |(x - x_n) + (x_n - y)| \leq |x_n - x| + |x_n - y| < 2\epsilon = |x - y|,$$

Widerspruch! □

DEFINITION 2.4 (Beschränktheit). Eine Folge $(x_n)_{n \in \mathbb{N}}$ rationaler Zahlen heißt *nach oben beschränkt*, wenn es ein $M > 0$ gibt, sodaß $x_n \leq M$ für alle $n \in \mathbb{N}$. Sie heißt *nach unten beschränkt*, wenn $(-x_n)_{n \in \mathbb{N}}$ nach oben beschränkt ist. Sie heißt *beschränkt*, wenn sie von oben und von unten beschränkt ist.

¹Plural *Limites*.

²Das Symbol ∞ bezeichnet hier nicht etwa eine Zahl oder ein anderes wohldefiniertes mathematisches Objekt, sondern ist lediglich Teil einer Schreibweise, mit der die Konvergenz der Folge ausgedrückt werden kann.

Die Folge $((-1)^n)_{n \in \mathbb{N}} = (-1, 1, -1, 1, \dots)$ ist etwa beschränkt (z.B. durch -1 von unten und durch 1 von oben), die Folge $(n)_{n \in \mathbb{N}}$ dagegen nicht, denn nach dem archimedischen Axiom existiert für jedes $Q \ni M > 0$ ein $n \in \mathbb{N}$, sodaß $n > M$.

Offenbar ist $(x_n)_{n \in \mathbb{N}}$ beschränkt genau dann, wenn $(|x_n|)_{n \in \mathbb{N}}$ (nach oben) beschränkt ist.

PROPOSITION 2.5. *Jede konvergente Folge ist beschränkt.*

BEWEIS. Es gelte $\lim_{n \rightarrow \infty} x_n = x$, dann existiert ein $N \in \mathbb{N}$, sodaß $|x_n - x| < 1$ für alle $n \geq N$ (hier haben wir $\epsilon = 1$ gesetzt). Insbesondere gilt für solche n : $|x_n| < 1 + |x|$; denn aus der Dreiecksungleichung folgt $|x_n - x| + |x| \geq |x_n|$. Setze

$$M := \max\{|x_1|, |x_2|, \dots, |x_{N-1}|, 1 + |x|\},$$

wobei wir mit dem Ausdruck auf der rechten Seite die größte unter den aufgeführten Zahlen meinen. Dann ist $|x_n| \leq M$ für alle $n \in \mathbb{N}$. $\square$

BEMERKUNG 2.6. Man beachte: Die Menge, über die im obigen Beweis das Maximum gebildet wird, ist endlich; deshalb existiert das Maximum überhaupt. Das Maximum über eine unendliche Menge braucht im Allgemeinen nicht zu existieren: Das Maximum über $\mathbb{N}$ existiert zum Beispiel nicht (bzw. ist „unendlich“), da es keine größte natürliche Zahl gibt.

Die Umkehrung gilt nicht: Die Folge $((-1)^n)_{n \in \mathbb{N}}$ ist, wie wir gesehen haben, beschränkt. Sie ist aber nicht konvergent, denn wäre x der Grenzwert, so gäbe es $N \in \mathbb{N}$ mit $|(-1)^n - x| < 1$ für alle $n \geq N$; insbesondere wäre $|1 - x| < 1$ und $|-1 - x| < 1$. Nach Dreiecksungleichung folgt aber

$$2 = |2| = |-1 - 1| = |(-1 - x) + (x - 1)| \leq |-1 - x| + |x - 1| < 1 + 1 = 2,$$

und $2 < 2$ ist ein Widerspruch zur Annahme der Konvergenz.

DEFINITION 2.7 (bestimmte Divergenz). Eine Folge $(x_n)_{n \in \mathbb{N}}$ heißt *divergent*, wenn sie nicht konvergent ist. Sie heißt *bestimmt divergent* gegen $+\infty$, wenn es zu jedem $M > 0$ ein $N \in \mathbb{N}$ gibt, sodaß für alle $n \geq N$ gilt: $x_n > M$. Sie heißt bestimmt divergent gegen $-\infty$, wenn $(-x_n)_{n \in \mathbb{N}}$ bestimmt gegen $+\infty$ divergiert. Eine Folge, die weder konvergiert noch bestimmt divergiert, heißt *unbestimmt divergent*.

Insbesondere ist jede beschränkte nicht konvergente Folge unbestimmt divergent, so z.B. $((-1)^n)_{n \in \mathbb{N}}$. Aber auch die unbeschränkte Folge $((-1)^n n)_{n \in \mathbb{N}}$ ist unbestimmt divergent. Die Folge $(n^2)_{n \in \mathbb{N}}$ divergiert dagegen bestimmt gegen $+\infty$.

BEISPIEL 2.8. Ist $x \in \mathbb{Q}$ mit $x > 1$, so divergiert $(x^n)_{n \in \mathbb{N}}$ bestimmt gegen $+\infty$. Um dies zu sehen, setze $y := x - 1 > 0$ und verwende die BERNOULLI-Ungleichung (Übungsblatt 4):

$$x^n = (1 + y)^n \geq 1 + ny.$$

Nach dem archimedischen Axiom existiert zu jedem $M > 0$ ein $N \in \mathbb{N}$, sodaß $1 + ny > M$ für jedes $n \geq N$, was zu zeigen war.

Andererseits gilt für $0 < x < 1$: $\lim_{n \rightarrow \infty} x^n = 0$. Dies läßt sich leicht verifizieren, indem man den ersten Teil dieses Beispiels auf $\frac{1}{x} > 1$ anwendet.

SATZ 2.9 (Limites und Grundrechenarten). Seien $(x_n)_{n \in \mathbb{N}}$ und $(y_n)_{n \in \mathbb{N}}$ konvergente Folgen mit Limites x bzw. y . Dann sind auch $(x_n \pm y_n)_{n \in \mathbb{N}}$ und $(x_n y_n)_{n \in \mathbb{N}}$ konvergent mit Limites $x \pm y$ bzw. xy . Ist $y \neq 0$, so existiert $N \in \mathbb{N}$, sodaß auch $y_n \neq 0$ für $n \geq N$, und die Folge $\left(\frac{x_n}{y_n}\right)_{n \geq N}$ konvergiert gegen $\frac{x}{y}$.

BEWEIS. Sei $\epsilon > 0$, so gibt es ein $N \in \mathbb{N}$, sodaß $|x_n - x| < \frac{\epsilon}{2}$ und $|y_n - y| < \frac{\epsilon}{2}$ für alle $n \geq N$. Nach Dreiecksungleichung ist dann für solche n

$$|(x_n + y_n) - (x + y)| = |(x_n - x) + (y_n - y)| \leq |x_n - x| + |y_n - y| < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon.$$

Die Aussage über Differenzen folgt mit der Beobachtung, daß $\lim_{n \rightarrow \infty} y_n = y$ auch $\lim_{n \rightarrow \infty} (-y_n) = -y$ impliziert (denn $|y_n - y| = |-y_n - (-y)|$).

Für das Produkt benutzen wir Proposition 2.5, derzufolge ein $M > 0$ existiert, sodaß $|x_n|, |y_n| \leq M$ für alle $n \in \mathbb{N}$. Damit gilt

$$|x_n y_n - xy| = |x_n y_n - x y_n + x y_n - xy| \leq |(x_n - x)y_n| + |x(y_n - y)| \leq M|x_n - x| + M|y_n - y|.$$

Wählen wir zu gegebenem $\epsilon > 0$ also $N \in \mathbb{N}$ so groß, daß $|x_n - x| < \frac{\epsilon}{2M}$ und $|y_n - y| < \frac{\epsilon}{2M}$, dann folgt in der Tat $|x_n y_n - xy| < \epsilon$.

Es bleibt die Aussage über die Quotienten zu beweisen. Ist $y \neq 0$, so wählen wir $\epsilon := \frac{|y|}{2} > 0$, um ein $N \in \mathbb{N}$ zu erhalten, sodaß $y_n \neq 0$ für alle $n \geq N$. Nach der eben bewiesenen Verträglichkeit des Grenzwerts mit Produkten genügt es zu zeigen, daß $\frac{1}{y_n} \rightarrow \frac{1}{y}$ für $n \rightarrow \infty$. Dazu rechnen wir

$$\left| \frac{1}{y_n} - \frac{1}{y} \right| = \left| \frac{y - y_n}{y y_n} \right| = \frac{|y - y_n|}{|y| |y_n|} < \frac{2}{|y|^2} |y - y_n|,$$

wo wir im letzten Schritt verwendet haben $|y_n| > \frac{|y|}{2}$ für $n \geq N$. Ist also $\epsilon > 0$ gegeben, so finden wir ein $N^* \geq N$, sodaß für alle $n \geq N^*$ gilt

$$|y - y_n| < \frac{\epsilon |y|^2}{2},$$

und es folgt wie gewünscht

$$\left| \frac{1}{y_n} - \frac{1}{y} \right| < \epsilon.$$

□

Die Nützlichkeit dieses Satzes erkennen wir an folgenden Beispielen:

BEISPIEL 2.10. Jede Linearkombination konvergenter Folgen ist wieder konvergent, das heißt: Sind die N Folgen $(x_n^1)_{n \in \mathbb{N}}, (x_n^2)_{n \in \mathbb{N}}, \dots, (x_n^N)_{n \in \mathbb{N}}$ konvergent mit Limites $x^1, x^2, \dots, x^N$, und sind $a^1, a^2, \dots, a^N \in \mathbb{Q}$, so gilt

$$\lim_{n \rightarrow \infty} \sum_{j=1}^N a^j x_n^j = \sum_{j=1}^N a^j x^j.$$

Dies folgt unmittelbar aus dem Satz.

BEISPIEL 2.11. Ist $k \in \mathbb{N}$ und $(x_n)_{n \in \mathbb{N}}$ konvergent mit $\lim_{n \rightarrow \infty} x_n = x$, so ist auch $(x_n^k)_{n \in \mathbb{N}}$ konvergent mit Grenzwert x^k . Dies folgt durch $(k-1)$ -malige Anwendung der Regel über Produkte.

BEISPIEL 2.12. Wir betrachten $(x_n)_{n \in \mathbb{N}}$ mit

$$x_n = \frac{3n^7 - n^2 + 2n + 1}{1 - n^7}.$$

Zunächst ist klar, daß für $n > 1$ der Nenner ungleich null ist. Durch Kürzen mit n^7 erhalten wir

$$x_n = \frac{3 - n^{-5} + 2n^{-6} + n^{-7}}{n^{-7} - 1}.$$

Gemäß dem vorigen Beispiel konvergieren negative Potenzen von n gegen null, und nach Satz kann man den Limes in Summen, Quotienten etc. ‚reinziehen‘. Es folgt

$$\lim_{n \rightarrow \infty} x_n = -3.$$

SATZ 2.13. Seien $(x_n)_{n \in \mathbb{N}}$ und $(y_n)_{n \in \mathbb{N}}$ konvergente Folgen mit Limites x bzw. y , und es gelte $x_n \leq y_n$ für alle $n \in \mathbb{N}$. Dann gilt auch $x \leq y$.

BEWEIS. Angenommen $x > y$, dann wäre $\epsilon := \frac{x-y}{2} > 0$. Für dieses ϵ finden wir $N \in \mathbb{N}$, sodaß $|x_n - x|, |y_n - y| < \epsilon$ für alle $n \geq N$. Nach Wahl von ϵ gilt für solche n insbesondere $x_n > \frac{x+y}{2}$ sowie $y_n < \frac{x+y}{2}$, also $x_n > y_n$ im Widerspruch zur Voraussetzung. $\square$

Der Satz bleibt offensichtlich gültig, wenn statt $x_n \leq y_n$ für *alle* $n \in \mathbb{N}$ lediglich $x_n \leq y_n$ für *fast alle*, d.h. für alle bis auf endlich viele $n \in \mathbb{N}$ gefordert wird.

Man darf in diesem Satz nicht $\leq$ durch $<$ ersetzen: Für jedes $n \in \mathbb{N}$ ist zwar $\frac{1}{n} > 0$, aber die Grenzwert ist (nicht etwa größer, sondern gleich) null.

SATZ 2.14 (Sandwich-Theorem). Es seien $(x_n)_{n \in \mathbb{N}}$, $(y_n)_{n \in \mathbb{N}}$ und $(z_n)_{n \in \mathbb{N}}$ mit $x_n \leq y_n \leq z_n$ für alle $n \in \mathbb{N}$, und es gelte $\lim_{n \rightarrow \infty} x_n = \lim_{n \rightarrow \infty} z_n =: x$. Dann gilt auch $\lim_{n \rightarrow \infty} y_n = x$.

BEWEIS. Sei $\epsilon > 0$, dann gibt es $N \in \mathbb{N}$ mit $x_n > x - \epsilon$ und $z_n < x + \epsilon$ für alle $n \geq N$. Da $x_n \leq y_n \leq z_n$, folgt für solche n

$$x - \epsilon < y_n < x + \epsilon,$$

also $|x - y_n| < \epsilon$. $\square$

Beispielsweise konvergiert die Folge $((-1)^n \frac{1}{n})_{n \in \mathbb{N}}$, da sie zwischen den beiden gegen null konvergierenden Folgen $(\pm \frac{1}{n})_{n \in \mathbb{N}}$ eingeklemmt ist.

DEFINITION 2.15 (Cauchyfolgen). Eine Folge $(x_n)_{n \in \mathbb{N}}$ rationaler Zahlen heißt *Cauchyfolge*, wenn es zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$ gibt, sodaß für alle $m, n \geq N$ gilt $|x_n - x_m| < \epsilon$.

SATZ 2.16. Jede konvergente Folge ist Cauchy.

BEWEIS. Sei $(x_n)_{n \in \mathbb{N}}$ konvergent gegen $x \in \mathbb{Q}$, und sei $\epsilon > 0$. Wir wählen $N \in \mathbb{N}$ so groß, daß $|x_n - x| < \frac{\epsilon}{2}$ für alle $n \geq N$. Dann gilt für alle $n, m \geq N$ nach Dreiecksungleichung:

$$|x_n - x_m| \leq |x_n - x| + |x_m - x| < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon.$$

$\square$

Man sieht leicht, daß jede Cauchyfolge beschränkt ist (wähle etwa $\epsilon = 1$ und verwende, daß ab einem gewissen Index $N \in \mathbb{N}$ alle Glieder sich von x_N um höchstens 1 unterscheiden).

2.2. Reelle Zahlen

2.2.1. Definition. Die Umkehrung von Satz 2.16 gilt in $\mathbb{Q}$ nicht: Die (rationale) Folge $(1; 1, 4; 1, 41; 1, 414; \dots)$ der Dezimaldarstellungen von $\sqrt{2}$ bis zu einer gewissen Nachkommastelle ist Cauchy (die Differenz zweier Folgenglieder ab Index N ist nämlich höchstens 10^{1-N}), aber die Folge konvergiert nicht in $\mathbb{Q}$ (denn der Grenzwert, wenn er existierte, wäre $\sqrt{2}$, und wir haben bereits gesehen, daß das nicht rational ist). Um die Äquivalenz zwischen Konvergenz und Cauchy-Eigenschaft zu erhalten, erweitern wir im Folgenden die rationalen zu den reellen Zahlen.

Eine gegen null konvergente Folge rationaler Zahlen bezeichnen wir kurz als *Nullfolge*.

Wir führen auf der Menge der rationalen Cauchyfolgen die Äquivalenzrelation $\equiv$ wie folgt ein:

$$(x_n)_{n \in \mathbb{N}} \equiv (y_n)_{n \in \mathbb{N}} \quad \text{genau dann, wenn } (x_n - y_n)_{n \in \mathbb{N}} \text{ eine Nullfolge ist.}$$

Zum Beispiel sind die Folge $(1 + \frac{1}{n})_{n \in \mathbb{N}}$ und die konstante Folge $(1)_{n \in \mathbb{N}}$ äquivalent, weil ihre Differenz $(\frac{1}{n})_{n \in \mathbb{N}}$ eine Nullfolge ist.

Es handelt sich tatsächlich um eine Äquivalenzrelation: Jede Folge minus sich selbst ist konstant null und damit eine Nullfolge; ist $(x_n - y_n)_{n \in \mathbb{N}}$ eine Nullfolge, so auch $(y_n - x_n)_{n \in \mathbb{N}}$; und wenn sowohl $(x_n - y_n)_{n \in \mathbb{N}}$ als auch $(y_n - z_n)_{n \in \mathbb{N}}$ Nullfolgen sind, so auch $(x_n - z_n)_{n \in \mathbb{N}}$ (denn $x_n - z_n = (x_n - y_n) + (y_n - z_n)$, was nach Satz 2.9 wieder gegen null konvergiert).

DEFINITION 2.17 (Reelle Zahlen). Eine Äquivalenzklasse bzgl. der Äquivalenzrelation $\equiv$ nennt man *reelle Zahl*. Die Menge der reellen Zahlen wird mit $\mathbb{R}$ bezeichnet.

Strenggenommen müßten wir also eine reelle Zahl als $[(x_n)_{n \in \mathbb{N}}]$ schreiben, wobei $(x_n)_{n \in \mathbb{N}}$ eine Cauchyfolge rationaler Zahlen ist. Wir vereinfachen im Folgenden die Notation und schreiben nur (x_n) .

Ist $q \in \mathbb{Q}$ eine rationale Zahl, so identifizieren wir sie mit der reellen Zahl, die durch die konstante Folge $(q)_{n \in \mathbb{N}}$ dargestellt wird. Konstante Folgen sind selbstverständlich Cauchy.

2.2.2. Arithmetik in $\mathbb{R}$.

DEFINITION 2.18. Seien $x = (x_n), y = (y_n) \in \mathbb{R}$. Dann setzen wir

- (1) $x + y := (x_n + y_n)$,
- (2) $x - y := (x_n - y_n)$,
- (3) $xy := (x_n y_n)$,
- (4) $\frac{x}{y} := \left(\frac{x_n}{y_n}\right)$, falls $y \neq 0$.

Hier gilt es wie immer nachzuprüfen, ob die angegebenen Operationen überhaupt wohldefiniert sind. Zunächst müssen wir uns überzeugen, daß die Folgen auf der rechten Seite jeweils Cauchy sind. Der Beweis, daß Summen, Differenzen, Produkte und Quotienten von Cauchyfolgen wieder Cauchy sind, folgt ähnlich wie in Satz 2.9. Zum Beispiel für das Produkt: Da Cauchyfolgen beschränkt sind (Beweis wie in Proposition 2.5), gibt es $M > 0$, sodaß $|x_n|, |y_n| \leq M$ für alle $n \in \mathbb{N}$. Damit gilt

$$\begin{aligned} |x_n y_n - x_m y_m| &= |x_n y_n - x_m y_n + x_m y_n - x_m y_m| \\ &\leq |(x_n - x_m) y_n| + |x_m (y_n - y_m)| \leq M|x_n - x_m| + M|y_n - y_m|. \end{aligned}$$

Wählen wir zu gegebenem $\epsilon > 0$ also $N \in \mathbb{N}$ so groß, daß $|x_n - x_m| < \frac{\epsilon}{2M}$ und $|y_n - y_m| < \frac{\epsilon}{2M}$, dann folgt in der Tat $|x_n y_n - x_m y_m| < \epsilon$.

Bei den Quotienten ist wie in Satz 2.9 zu beachten, daß im Falle $y \neq 0$ fast alle Glieder der darstellenden Cauchyfolge ihrerseits ungleich null sind. In der Tat: $y \neq 0$ bedeutet, daß die Cauchyfolge (y_n) nicht gegen null konvergiert. Es gibt also ein $\epsilon > 0$, sodaß $|y_n| > \epsilon$ für unendlich viele $n \in \mathbb{N}$. Da die Folge Cauchy ist, gibt es zu diesem ϵ ein $N \in \mathbb{N}$, sodaß für alle $n, m \geq N$ gilt $|y_n - y_m| < \frac{\epsilon}{2}$. Sei $n \geq N$ so gewählt, daß $|y_n| > \epsilon$; dann gilt für alle $m \geq n$:

$$|y_m| \geq |y_n| - |y_m - y_n| > \epsilon - \frac{\epsilon}{2} = \frac{\epsilon}{2} > 0.$$

Weiterhin ist zu beachten, daß die Rechenoperationen von der konkreten Darstellung durch eine Cauchyfolge unabhängig sind. Es genügt dazu zu zeigen: Sind (x_n) und (y_n) Cauchyfolgen und $(x'_n), (y'_n)$ dazu äquivalente Cauchyfolgen (sodaß also $x_n - x'_n \rightarrow 0$ und $y_n - y'_n \rightarrow 0$ für $n \rightarrow \infty$), so sind $(x_n + y_n - (x'_n + y'_n))$, $(x_n - y_n - (x'_n - y'_n))$, $(x_n y_n - x'_n y'_n)$ sowie $\left(\frac{x_n}{y_n} - \frac{x'_n}{y'_n}\right)$ ebenfalls Nullfolgen.

Wir zeigen dies exemplarisch für das Produkt: Wir haben

$$x_n y_n - x'_n y'_n = (x_n y_n - x'_n y_n) + (x'_n y_n - x'_n y'_n) = y_n(x_n - x'_n) + x'_n(y_n - y'_n);$$

Da aber (y_n) und (x'_n) als Cauchyfolgen beschränkt sind, und da $(x_n - x'_n)$ und $(y_n - y'_n)$ Nullfolgen sind, und da schließlich das Produkt einer beschränkten Folge mit einer Nullfolge wieder eine Nullfolge ist (Übung!), folgt in der Tat $x_n y_n - x'_n y'_n \rightarrow 0$ für $n \rightarrow \infty$.

Die Grundrechenarten in $\mathbb{R}$ sind in der angegebenen Weise also wohldefiniert. Offenbar sind sie mit denen in $\mathbb{Q}$ konsistent.

SATZ 2.19 ($\mathbb{R}$ ist ein Körper). Addition und Multiplikation auf $\mathbb{R}$ sind assoziativ und kommutativ und erfüllen das Distributivgesetz. Für alle $x \in \mathbb{R}$ gilt $x + 0 = x$, $1 \cdot x = x$, $x - x = 0$ und, falls $x \neq 0$, $x \cdot \frac{1}{x} = 1$.

BEWEIS. Dies folgt aus den entsprechenden Gesetzen in $\mathbb{Q}$. Zum Beispiel ist für $x, y \in \mathbb{R}$

$$x + y = (x_n + y_n) = (y_n + x_n) = y + x.$$

□

2.2.3. Die Anordnung von $\mathbb{R}$.

DEFINITION 2.20. Eine reelle Zahl $x = (x_n)$ heißt *positiv*, wenn es ein $\epsilon > 0$ gibt, so daß $x_n > \epsilon$ für fast alle (also alle bis auf endlich viele) $n \in \mathbb{N}$. Ist x weder positiv noch null, so heißt x *negativ*.

Man beachte, daß es für die Positivität von x nicht ausreicht, daß $x_n > 0$ für (fast) alle $n \in \mathbb{N}$: Die durch die Nullfolge $(\frac{1}{n})$ dargestellte reelle Zahl ist null, obwohl $\frac{1}{n} > 0$ für alle $n \in \mathbb{N}$.

Auch diese Definition ist unabhängig von der Wahl der darstellenden Cauchyfolge: Ist (x_n) Cauchy mit $x_n > \epsilon > 0$ für fast alle $n \in \mathbb{N}$ und ist (p_n) eine Nullfolge, so gibt es ein $N \in \mathbb{N}$, sodaß $|p_n| < \frac{\epsilon}{2}$ für alle $n \geq N$; dann aber ist für hinreichend große $n \in \mathbb{N}$

$$x_n + p_n > \frac{\epsilon}{2} > 0,$$

erfüllt also ebenfalls die Positivitätsbedingung. Wie gehabt schreiben wir $x < y$ falls $y - x > 0$ und analog für $\geq, <, \leq$.

SATZ 2.21 ($\mathbb{R}$ ist ein archimedisch angeordneter Körper).

- (1) Für jedes $x \in \mathbb{R}$ gilt genau eine der drei Aussagen $x > 0$, $x = 0$ oder $x < 0$.
- (2) Für $x, y \in \mathbb{R}$ folgt aus $x > 0$ und $y > 0$, daß $x + y > 0$ und daß $xy > 0$.
- (3) Für alle $x, y \in \mathbb{R}$ mit $x, y > 0$ existiert ein $N \in \mathbb{N}$, sodaß $x < Ny$.

BEWEIS. Zu (1): Es ist lediglich zu zeigen, daß es keine reelle Zahl gibt, die positiv und null ist. Dies folgt aber unmittelbar aus der Definition der Positivität und der Definition der Konvergenz gegen null.

Zu (2): Seien $x = (x_n)$ und $y = (y_n)$ positiv. Dann existiert $\epsilon > 0$, sodaß $x_n, y_n > \epsilon$ für fast alle $n \in \mathbb{N}$. Mit Korollar 1.41 folgt $x_n + y_n > 2\epsilon$ sowie $x_n y_n > \epsilon^2$ für fast alle $n \in \mathbb{N}$, also sind $x + y, xy > 0$.

Zu (3): Seien $x = (x_n)$ und $y = (y_n)$ positiv. Dann existiert $\epsilon > 0$, sodaß $x_n, y_n > \epsilon$ für fast alle $n \in \mathbb{N}$. Da (x_n) Cauchy ist, ist es auch beschränkt, also existiert $M \in \mathbb{Q}$ mit $x_n \leq M - 1$ für alle $n \in \mathbb{N}$. Nach dem archimedischen Axiom in $\mathbb{Q}$ gibt es ein $N \in \mathbb{N}$, sodaß $M < N\epsilon$. Insgesamt ergibt sich für fast alle $n \in \mathbb{N}$:

$$Ny_n > N\epsilon > M \geq x_n + 1,$$

also $Ny_n - x_n \geq 1$ für fast alle $n \in \mathbb{N}$. Es folgt nun $x < Ny$ aus der Definition der Positivität. □

Genau wie in Definition 1.42 können wir den Betrag als Abbildung $\mathbb{R} \rightarrow \mathbb{R}$ definieren.

SATZ 2.22. Korollar 1.41 und Proposition 1.43 bleiben gültig, wenn man $\mathbb{Q}$ durch $\mathbb{R}$ ersetzt. Das gilt auch für den gesamten Inhalt von Abschnitt 2.1.

BEWEIS. Für Formulierung und Beweis der genannten Definitionen und Aussagen haben wir lediglich die Eigenschaft von $\mathbb{Q}$ benutzt, ein archimedisch angeordneter Körper zu sein. Da auch $\mathbb{R}$ diese Eigenschaft besitzt, behalten alle Aussagen dort ihre Gültigkeit. □

SATZ 2.23 ($\mathbb{Q}$ liegt dicht in $\mathbb{R}$). Für jedes $x \in \mathbb{R}$ und jedes reelle $\epsilon > 0$ existiert ein $q \in \mathbb{Q}$ mit $|x - q| < \epsilon$.

BEWEIS. Sei x dargestellt durch die Cauchyfolge (x_n) , und sei $N \in \mathbb{N}$ so groß, daß $|x_n - x_m| < \frac{\epsilon}{2}$ für alle $n, m \geq N$. Setze $q := x_N \in \mathbb{Q}$. Dann gilt für alle $n \geq N$:

$$|x_n - q| < \frac{\epsilon}{2}$$

und daher einerseits $x_n - q < \frac{\epsilon}{2}$ und andererseits $q - x_n < \frac{\epsilon}{2}$. Es folgt mit der Definition der Positivität in $\mathbb{R}$: $x - q < \epsilon$ und $q - x < \epsilon$, zusammen also $|x - q| < \epsilon$ wie gewünscht. $\square$

2.3. Vollständigkeit

2.3.1. Cauchy-Vollständigkeit.

SATZ 2.24 ($\mathbb{R}$ ist vollständig). In $\mathbb{R}$ konvergiert jede Cauchyfolge.

BEWEIS. Sei $(x_n)_{n \in \mathbb{N}}$ eine Cauchyfolge reeller Zahlen. Nach Satz 2.23 können wir zu jedem $n \in \mathbb{N}$ ein $x'_n \in \mathbb{Q}$ finden, sodaß $|x_n - x'_n| < \frac{1}{n}$. Wir zeigen, daß die Folge $(x'_n)_{n \in \mathbb{N}}$ ihrerseits Cauchy ist: Sei nämlich $\epsilon > 0$ und $N \in \mathbb{N}$ so groß, daß $|x_n - x_m| < \frac{\epsilon}{3}$ für alle $n, m \geq N$. Wir dürfen annehmen, daß zudem $\frac{1}{N} < \frac{\epsilon}{3}$. Dann gilt für alle $n, m \geq N$:

$$|x'_n - x'_m| \leq |x'_n - x_n| + |x_n - x_m| + |x_m - x'_m| < \frac{1}{N} + \frac{\epsilon}{3} + \frac{1}{N} < \frac{\epsilon}{3} + \frac{\epsilon}{3} + \frac{\epsilon}{3} = \epsilon,$$

also ist (x'_n) eine Cauchyfolge rationaler Zahlen. Diese Folge definiert daher eine reelle Zahl x . Wir zeigen, daß $\lim_{n \rightarrow \infty} x_n = x$: Sei dazu abermals $\epsilon > 0$ und $N \in \mathbb{N}$ so groß, daß $|x'_n - x'_m| < \frac{\epsilon}{3}$ für alle $n, m \geq N$, und daß $\frac{1}{N} < \frac{\epsilon}{3}$. Aus $|x'_n - x'_m| < \frac{\epsilon}{3}$ folgt nach Definition 2.20 $|x - x'_n| < \frac{2}{3}\epsilon$, und daraus ergibt sich für alle $n \geq N$

$$|x - x_n| \leq |x - x'_n| + |x'_n - x_n| < \frac{2}{3}\epsilon + \frac{1}{N} < \frac{2}{3}\epsilon + \frac{1}{3}\epsilon = \epsilon.$$

$\square$

Wir werden bald sehen, daß die Vollständigkeit von $\mathbb{R}$ die Existenz einer reellen Zahl impliziert, deren Quadrat gleich 2 ist.

2.3.2. Intervallschachtelung. Die Vollständigkeit manifestiert sich aber auch in vielfältiger anderer Weise, zum Beispiel in Gestalt des Intervallschachtelungsprinzips. Wir definieren dazu das *abgeschlossene Intervall* zwischen zwei reellen Zahlen $a \leq b$ als

$$[a, b] := \{x \in \mathbb{R} : a \leq x \leq b\},$$

und die *Länge* des Intervalls als $|I| := b - a$.

SATZ 2.25 (Intervallschachtelung). Sei $(I_n)_{n \in \mathbb{N}}$ eine Folge abgeschlossener Intervalle in $\mathbb{R}$, sodaß $I_n \supset I_{n+1}$ für alle $n \in \mathbb{N}$ und $\lim_{n \rightarrow \infty} |I_n| = 0$. Dann enthält die Menge $\bigcap_{n \in \mathbb{N}} I_n$ genau eine reelle Zahl. Mit anderen Worten: Es gibt genau ein $x \in \mathbb{R}$, sodaß $x \in I_n$ für alle $n \in \mathbb{N}$.

BEWEIS. Ist $I_n = [a_n, b_n]$, so behaupten wir, daß $(a_n)_{n \in \mathbb{N}}$ Cauchy ist. Sei dazu $\epsilon > 0$ und wähle $N \in \mathbb{N}$ so groß, daß $b_n - a_n < \epsilon$ für alle $n \geq N$. Seien $n, m \geq N$ mit $m \geq n$, so ist nach Voraussetzung $a_m \geq a_n$ sowie $b_m \leq b_n$ und somit

$$|a_m - a_n| = a_m - a_n \leq b_m - a_n \leq b_n - a_n < \epsilon.$$

Da also $(a_n)_{n \in \mathbb{N}}$ Cauchy ist, konvergiert diese Folge nach Satz 2.24 gegen ein $x \in \mathbb{R}$. Da für alle $n \leq m$ gilt $a_n \leq a_m \leq b_m \leq b_n$, gilt nach Satz 2.13 $a_n \leq x \leq b_n$, also $x \in I_n$, für alle $n \in \mathbb{N}$. Ist $y \in \mathbb{R}$ eine weitere Zahl mit dieser Eigenschaft, so gilt $|I_n| \geq |x - y|$ für alle $n \in \mathbb{N}$, was nach der Voraussetzung $\lim_{n \rightarrow \infty} |I_n| = 0$ nur dann möglich ist, wenn $x = y$. $\square$

2.3.3. Suprema und Infima von Teilmengen von $\mathbb{R}$.

DEFINITION 2.26. Sei $U \subset \mathbb{R}$ eine Teilmenge von $\mathbb{R}$. Ein $S \in \mathbb{R}$ heißt *obere Schranke* für U , wenn $x \leq S$ für alle $x \in U$. Ein $s \in \mathbb{R}$ heißt *untere Schranke* für U , wenn $x \geq s$ für alle $x \in U$. Existiert eine obere bzw. untere Schranke für U , so heißt U *beschränkt*.

Existiert eine *kleinste obere Schranke* für U , also eine obere Schranke $S \in \mathbb{R}$ mit $S' \geq S$ für alle oberen Schranken S' , so bezeichnet man dieses S als *Supremum*³ von U und schreibt dafür $\sup U$.

Existiert eine *größte untere Schranke* für U , so nennt man sie das *Infimum* von U und schreibt dafür $\inf U$.

Ist $\mathbb{R} \ni \sup U \in U$ bzw. $\mathbb{R} \ni \inf U \in U$, so heißt $\sup U$ auch *Maximum* bzw. $\inf U$ *Minimum* von U , und man schreibt dafür $\max U$ bzw. $\min U$.

Aus der Definition ist klar, daß Supremum und Infimum einer Menge im Falle ihrer Existenz eindeutig bestimmt sind. Ist eine Teilmenge von $\mathbb{R}$ nach oben bzw. nach unten unbeschränkt, so schreibt man $\sup U = +\infty$ bzw. $\inf U = -\infty$. Die leere Menge ist zwar beschränkt, besitzt aber weder Supremum noch Infimum; man schreibt aber manchmal $\sup \emptyset = -\infty$ und $\inf \emptyset = +\infty$ (warum?). Für nichtleere Mengen U ist dagegen stets $\inf U \leq \sup U$.

BEISPIEL 2.27. Für ein abgeschlossenes Intervall $I = [a, b]$ ist $\sup I = b = \max I$ und $\inf I = a = \min I$. Für die Menge $U = \{\frac{1}{n} : n \in \mathbb{N}\}$ gilt $\inf U = 0$ (aber die Menge besitzt kein Minimum, da $0 \notin U$) und $\sup U = \max U = 1$. Supremum und Infimum können also, müssen aber nicht in der Menge selbst enthalten sein.

Die Menge $\{x \in \mathbb{Q} : x^2 < 2\}$ besitzt in $\mathbb{Q}$ weder Infimum noch Supremum, in $\mathbb{R}$ aber wohl (nämlich $\pm\sqrt{2}$; wir kommen noch darauf zurück).

Das letzte Beispiel zeigt, daß in $\mathbb{Q}$ eine beschränkte Menge kein Infimum oder Supremum zu haben braucht. Die reellen Zahlen haben dagegen dank ihrer Vollständigkeit diese Eigenschaft:

SATZ 2.28. Sei eine nichtleere Teilmenge $U \subset \mathbb{R}$ nach oben beschränkt. Dann existiert $\sup U$ (als reelle Zahl). Ist $U \subset \mathbb{R}$ nach unten beschränkt, so existiert $\inf U$.

BEWEIS. Wir zeigen nur die Aussage über das Supremum, die über das Infimum folgt analog.

Sei $S_1 \in \mathbb{R}$ eine obere Schranke für U , und wähle irgendein Element $x_1 \in U$. Dann ist $x_1 \leq S_1$. Wir definieren rekursiv zwei Folgen $(S_n)_{n \in \mathbb{N}}$ und $(x_n)_{n \in \mathbb{N}}$ mit Startwerten S_1 bzw. x_1 wie folgt: Seien S_n und x_n mit $x_n \leq S_n$ bereits definiert, so setze

$$x_{n+1} := \begin{cases} x_n & \text{falls } \frac{x_n + S_n}{2} \text{ obere Schranke für } U, \\ \frac{x_n + S_n}{2} & \text{sonst} \end{cases}$$

und

$$S_{n+1} := \begin{cases} S_n & \text{falls } \frac{x_n + S_n}{2} \text{ keine obere Schranke für } U, \\ \frac{x_n + S_n}{2} & \text{sonst.} \end{cases}$$

Offenbar ist jedes S_n obere Schranke für U , und für jedes n existiert $y \in U$, sodaß $x_n \leq y$. Mit $I_n := [x_n, S_n]$ ist außerdem klar, daß $I_n \supset I_{n+1}$ für alle $n \in \mathbb{N}$, und da sich in jedem Schritt die Intervalllänge halbiert, gilt $|I_n| = 2^{1-n}|I_1|$. Mit Beispiel 2.8 sehen wir $\lim_{n \rightarrow \infty} |I_n| = 0$. Die Voraussetzungen von Satz 2.25 sind also erfüllt, und dieser liefert ein eindeutiges $x \in \mathbb{R}$, sodaß $x \in I_n$ für alle $n \in \mathbb{N}$.

Wir zeigen, daß dieses x das gesuchte Supremum ist. Dazu ist zweierlei zu zeigen: x ist obere Schranke für U , und jede obere Schranke ist nicht kleiner als x .

³Plural *Suprema*.

Für die erste Aussage sei $\epsilon > 0$ und $n \in \mathbb{N}$ so groß, daß $|I_n| = S_n - x_n < \epsilon$. Da $x_n \leq x \leq S_n$, ist auch $S_n - x < \epsilon$. Für jedes $y \in U$ ist aber $y \leq S_n$ (da S_n obere Schranke ist) und daher auch $y - x < \epsilon$. Da $\epsilon > 0$ beliebig war, folgt $y \leq x$.

Für ϵ und n wie eben gilt $0 \leq x - x_n < \epsilon$. Da aber ein $y \in U$ mit $y \geq x_n$ existiert, folgt $x - y < \epsilon$. Jede obere Schranke S für U erfüllt aber $y \leq S$, also haben wir $x - S < \epsilon$. Da $\epsilon > 0$ beliebig war, gilt sogar $x \leq S$, und damit ist alles bewiesen. $\square$

DEFINITION 2.29 (Monotonie von Folgen). Eine Folge $(x_n)_{n \in \mathbb{N}}$ reeller Zahlen heißt *monoton wachsend*, wenn $x_{n+1} \geq x_n$ für alle $n \in \mathbb{N}$. Gilt sogar $x_{n+1} > x_n$ für alle $n \in \mathbb{N}$, so heißt die Folge *streng monoton wachsend*.

Eine Folge $(x_n)_{n \in \mathbb{N}}$ heißt *(streng) monoton fallend*, wenn $(-x_n)_{n \in \mathbb{N}}$ (streng) monoton wachsend ist.

KOROLLAR 2.30. Jede von oben beschränkte monoton wachsende Folge konvergiert in $\mathbb{R}$. Ebenso ist jede von unten beschränkte monoton fallende Folge in $\mathbb{R}$ konvergent.

BEWEIS. Sei $(x_n)_{n \in \mathbb{N}}$ von oben beschränkt und monoton wachsend. Dann ist die Menge $\{x_n : n \in \mathbb{N}\}$ nichtleer und von oben beschränkt und besitzt nach Satz 2.28 ein Supremum $x \in \mathbb{R}$. Wir zeigen, daß die Folge gegen dieses x konvergiert: Sei nämlich $\epsilon > 0$, dann existiert ein $N \in \mathbb{N}$, sodaß $0 \leq x - x_N < \epsilon$ (denn andernfalls wäre $x - \epsilon$ eine kleinere obere Schranke für $\{x_n : n \in \mathbb{N}\}$, mithin wäre x nicht das Supremum). Aufgrund der Monotonie der Folge ist dann aber auch für jedes $n \geq N$

$$|x - x_n| = x - x_n \leq x - x_N < \epsilon.$$

Die Aussage über von unten beschränkte monoton fallende Folgen geht analog. $\square$

BEISPIEL 2.31. Die Folge $(x_n)_{n \in \mathbb{N}}$ sei rekursiv definiert durch

$$x_1 := x^* \in \mathbb{R}, \quad x_{n+1} := x_n^2 - x_n + 1.$$

Wir zeigen, daß die Folge konvergent ist, sofern $0 \leq x^* \leq 1$: Sie ist monoton wachsend, denn für jedes $x \in \mathbb{R}$ gilt $x^2 - x + 1 \geq x$ (denn dies ist äquivalent zu $0 \leq x^2 - 2x + 1 = (x - 1)^2$). Außerdem gilt $0 \leq x_n \leq 1$ für alle $n \in \mathbb{N}$: Für x_1 ist dies nach Voraussetzung der Fall; Ist dagegen bereits $0 \leq x_n \leq 1$, so gilt $0 \leq x_n^2 \leq x_n$ und deshalb $0 \leq x_n^2 - x_n + 1 \leq 1$.

Als beschränkte monotone Folge ist $(x_n)_{n \in \mathbb{N}}$ nach Korollar 2.30 konvergent. Wir können auch den Grenzwert ermitteln: Sei $x = \lim_{n \rightarrow \infty} x_n$, dann gilt nach Satz 2.9

$$x = \lim_{n \rightarrow \infty} x_{n+1} = \lim_{n \rightarrow \infty} (x_n^2 - x_n + 1) = x^2 - x + 1.$$

Der Grenzwert erfüllt also $(x - 1)^2 = 0$, und wir erhalten $x = 1$.

Für $x = 0,5$ erhalten wir zum Beispiel die Folge (auf zwei Nachkommastellen gerundet):

$$(0,5; 0,75; 0,85; 0,87; 0,89; 0,90; 0,91; 0,92; 0,92; \dots).$$

Die Konvergenz ist in diesem Falle so langsam, daß man den Grenzwert nicht ohne Weiteres durch ‚Ausprobieren‘ ermitteln kann.

2.4. Teilfolgen und Häufungspunkte

2.4.1. Häufungspunkte. Eine Teilfolge ergibt sich aus einer gegebenen Folge, indem nur bestimmte Indizes ausgewählt werden. So sind z.B. $(\frac{1}{2n+1})_{n \in \mathbb{N}}$ oder $(\frac{1}{n^2})_{n \in \mathbb{N}}$ Teilfolgen von $(\frac{1}{n})_{n \in \mathbb{N}}$. Genauer definieren wir:

DEFINITION 2.32 (Teilfolge). Sei $(x_n)_{n \in \mathbb{N}}$ eine reelle Folge und sei $(n_k)_{k \in \mathbb{N}}$ eine streng monoton wachsende Folge natürlicher Zahlen, dann heißt die Folge $(x_{n_k})_{k \in \mathbb{N}}$ eine *Teilfolge* von $(x_n)_{n \in \mathbb{N}}$.

Mit $x_n = \frac{1}{n}$ und $n_k = 2k + 1$ bzw. $n_k = k^2$ erhalten wir die beiden einführenden Beispiele.

PROPOSITION 2.33. *Konvergiert eine Folge, so konvergiert auch jede Teilfolge.*

BEWEIS. Gelte $\lim_{n \rightarrow \infty} x_n = x$ und sei $(x_{n_k})_{k \in \mathbb{N}}$ eine Teilfolge. Sei $\epsilon > 0$ und $N \in \mathbb{N}$ so groß, daß $|x_n - x| < \epsilon$ für alle $n \geq N$. Da $(n_k)_{k \in \mathbb{N}}$ streng monoton wächst, ist $n_k \geq k$ für jedes $k \in \mathbb{N}$, also gilt für $k \geq N$ erst recht $|x_{n_k} - x| < \epsilon$. $\square$

Ebenso gilt: Divergiert eine Folge bestimmt gegen $\pm\infty$, so auch jede Teilfolge.

DEFINITION 2.34 (Häufungspunkt einer Folge). Ein $x \in \mathbb{R}$ heißt *Häufungspunkt* der Folge $(x_n)_{n \in \mathbb{N}}$, wenn es eine Teilfolge gibt, die gegen x konvergiert. Auch $\pm\infty$ ist Häufungspunkt der Folge, sofern eine Teilfolge bestimmt gegen $\pm\infty$ divergiert.

Betrachte etwa die Folge $((-1)^n)_{n \in \mathbb{N}}$. Diese besitzt genau die beiden Häufungspunkte ± 1 , die sich z.B. aus der Wahl der Teilfolgen $n_k = 2k$ und $n_k = 2k + 1$ ergeben.

Die Folge $(n)_{n \in \mathbb{N}}$ besitzt als Häufungspunkt genau $+\infty$, da sie (und damit jede Teilfolge) bestimmt gegen $+\infty$ divergiert. Die Folge $((-1)^n)_{n \in \mathbb{N}}$ besitzt genau die beiden Häufungspunkte $\pm\infty$, da jede ihrer Teilfolgen, die nicht unbestimmt divergiert, bestimmt gegen $\pm\infty$ divergiert.

Eine äquivalente Formulierung der Definition *reeller*⁴ Häufungspunkte wäre: $x \in \mathbb{R}$ ist Häufungspunkt der Folge $(x_n)_{n \in \mathbb{N}}$, wenn es zu jedem $\epsilon > 0$ unendlich viele Folgenglieder gibt mit $|x_n - x| < \epsilon$.

BEISPIEL 2.35. Eine Folge kann unendlich viele Häufungspunkte haben. Betrachte dazu die Folge

$$(1, 1, 2, 1, 2, 3, 1, 2, 3, 4, 1, 2, 3, 4, 5, \dots).$$

Offenbar kommt jede natürliche Zahl in dieser Folge unendlich oft vor, ist also Häufungspunkt. Zudem ist $+\infty$ Häufungspunkt, da $\{1, 2, 3, \dots\}$ eine bestimmt divergente Teilfolge ist. Wir werden später sehen, daß es Folgen gibt, die *jede reelle Zahl* als Häufungspunkt haben.

Der folgende Satz ist von zentraler Bedeutung in der Analysis, weil er den Prototyp eines sogenannten *Kompaktheitssatzes* darstellt:

SATZ 2.36 (BOLZANO-WEIERSTRASS). Jede beschränkte Folge besitzt einen reellen Häufungspunkt.

Mit anderen Worten: Jede beschränkte Folge besitzt eine konvergente Teilfolge.

BEWEIS. Sei die Folge durch $M > 0$ beschränkt, d.h. alle Folgenglieder liegen im Intervall $I_0 := [-M; M]$. Wir konstruieren eine Folge abgeschlossener Intervalle $I_n \subset I$, sodaß $I_{n+1} \subset I_n$ und $|I_n| = 2^{-n+1}M$ für alle $n \in \mathbb{N}$, und sodaß jedes I_n unendlich viele Folgenglieder enthält.

Zu diesem Zwecke definieren wir für ein abgeschlossenes Intervall $I = [a, b]$ die Teilintervalle

$$I^+ := \left[\frac{a+b}{2}, b \right], \quad I^- := \left[a, \frac{a+b}{2} \right].$$

Da die Folge ihre Werte in $I_0 = [-M, M]$ annimmt, liegen unendlich viele Glieder in I_0^+ oder in I_0^- . (Andernfalls hätte die Folge insgesamt nur endlich viele Glieder.) Definiere I_1 als I_0^+ oder I_0^- derart, daß I_1 unendlich viele Glieder enthält.

Ist I_n bereits definiert, so enthält wiederum I_n^+ oder I_n^- unendlich viele Folgenglieder, und I_{n+1} ist eines von beiden, das so gewählt wird, daß unendlich viele Folgenglieder in I_{n+1} liegen.

Da I^+ und I^- Teilmengen von I mit der halben Länge sind, erhalten wir auf diese Weise tatsächlich eine Folge von Intervallen mit $I_{n+1} \subset I_n$ und $|I_n| = 2^{-n+1}M$ für alle $n \in \mathbb{N}$.

⁴im Gegensatz zu $\pm\infty$.

Nach Satz 2.25 von der Intervallschachtelung existiert (genau) ein $x \in \mathbb{R}$, das in jedem I_n enthalten ist. Dieses x ist der gewünschte Häufungspunkt: Sei nämlich $\epsilon > 0$ und $N \in \mathbb{N}$ so groß, daß $2^{-N+1}M < \epsilon$. Dann gilt $|x - y| < \epsilon$ für jedes $y \in I_N$, und I_N enthält unendlich viele Folgenglieder. $\square$

Man beachte, daß in diesem Beweis die Wahl von I_n^+ oder I_n^- als dem neuen I_{n+1} nicht eindeutig sein muß: Es ist ja möglich, daß sowohl I_n^+ als auch I_n^- unendlich viele Folgenglieder enthalten. Dies entspricht der Möglichkeit, daß eine Folge durchaus mehrere Häufungspunkte haben kann.

KOROLLAR 2.37. Jede Folge besitzt einen Häufungspunkt (in $\mathbb{R}$ oder $\pm\infty$).

BEWEIS. Ist die Folge beschränkt, so folgt die Behauptung aus dem Satz von BOLZANO-WEIERSTRASS. Ist die Folge unbeschränkt, so ist sie von oben oder von unten unbeschränkt. Ist sie von oben unbeschränkt, so gibt es eine bestimmt gegen $+\infty$ divergente Teilfolge (warum?). Dann ist $+\infty$ Häufungspunkt. Ist sie von unten unbeschränkt, so ist analog $-\infty$ Häufungspunkt. $\square$

BEMERKUNG 2.38. Es mag der Leserin aufgefallen sein, daß in den vorangehenden Ausführungen die Konvergenz gegen eine reelle Zahl und die bestimmte Divergenz gegen $\pm\infty$ die gleiche Rolle gespielt haben. Man könnte daher beide Begriffe unter den Begriff der Konvergenz gegen eine Zahl in $\mathbb{R} \cup \{\pm\infty\}$ subsumieren und damit die Definition des Häufungspunkts und einige andere Aussagen ökonomischer formulieren. Es gibt in der Tat eine Möglichkeit, diese Idee in der gebotenen mathematischen Strenge auszuführen: Man bezeichnet diese Methode als *Kompaktifizierung* von $\mathbb{R}$. Ob man dies tun möchte oder nicht, hängt maßgeblich davon ab, was man mit der Struktur $\mathbb{R} \cup \{\pm\infty\}$ anstellen möchte: Die *algebraischen* Eigenschaften (vulgo Grundrechenarten) lassen sich nicht darauf erweitern⁵, die *topologischen* (mit Konvergenz zusammenhängenden) dagegen schon.

2.4.2. Limites inferiores et superiores.

DEFINITION 2.39 (Limites inferiores et superiores). Sei $(x_n)_{n \in \mathbb{N}}$ eine reelle Folge und H die Menge ihrer Häufungspunkte (ggf. inklusive $\pm\infty$). Dann heißt

$$\liminf_{n \rightarrow \infty} x_n := \inf H \quad \text{bzw.} \quad \limsup_{n \rightarrow \infty} x_n := \sup H$$

der *limes inferior* bzw. *limes superior* der Folge.

BEISPIEL 2.40. Es gilt zum Beispiel

$$\begin{aligned} \liminf_{n \rightarrow \infty} (-1)^n &= -1, & \limsup_{n \rightarrow \infty} (-1)^n &= 1, \\ \liminf_{n \rightarrow \infty} n &= \limsup_{n \rightarrow \infty} n = +\infty, \\ \liminf_{n \rightarrow \infty} (-1)^n n &= -\infty, & \limsup_{n \rightarrow \infty} (-1)^n n &= +\infty. \end{aligned}$$

Ist die Folge konvergent oder bestimmt divergent, so hat sie genau einen Häufungspunkt (nämlich ihren Grenzwert bzw. $\pm\infty$), und dann fallen die Begriffe $\lim$, $\liminf$, $\limsup$ zusammen.

⁵Denn angenommen, man definierte $\infty = \frac{1}{0}$. Dann würde einerseits gelten $\infty \cdot 0 = 1$ und andererseits nach Distributivgesetz $\infty \cdot (0 + 0) = \infty \cdot 0 + \infty \cdot 0 = 2$, also folgt mit $0 + 0 = 0$, daß $1 = 2$, Widerspruch.

KAPITEL 3

Stetige Funktionen

Wir beschränken uns nun auf Funktionen, die eine Teilmenge von $\mathbb{R}$ nach $\mathbb{R}$ abbilden. Zentral ist dabei die Eigenschaft der *Stetigkeit*: Kleine Änderungen im Input führen auch nur kleine Änderungen im Output mit sich. Wir geben verschiedene äquivalente Definitionen der Stetigkeit an und zeigen, daß stetige Funktionen besonders gute Eigenschaften haben.

3.1. Charakterisierungen von Stetigkeit

3.1.1. Definition und Beispiele. Sei $U \subset \mathbb{R}$ eine Teilmenge und $f : U \rightarrow \mathbb{R}$ eine Funktion.

DEFINITION 3.1 (Stetigkeit). Die Funktion f heißt *stetig* im Punkt $x \in U$, wenn gilt:

$$\forall \epsilon > 0 \quad \exists \delta > 0 \quad \forall y \in U : \quad |x - y| < \delta \implies |f(x) - f(y)| < \epsilon.$$

Ist f in jedem Punkt $x \in U$ stetig, so heißt sie (überall/auf ihrem gesamten Definitionsbereich) *stetig*.

BEISPIEL 3.2. (1) Die Funktion $\mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto x^2$ ist überall stetig. Sei nämlich $x \in \mathbb{R}$ beliebig und $\epsilon > 0$, so ist

$$|y^2 - x^2| = |(y + x)(y - x)| = |y + x||y - x|. \quad (3.1)$$

Wir wollen dies kleiner als ϵ bekommen, indem wir y sehr nah an x wählen. Dazu beobachten wir zunächst, daß $|y + x| < 2|x| + 1$ falls $|y - x| < 1$, denn dann gilt

$$|y + x| = |x + (y - x) + x| \leq 2|x| + |y - x| < 2|x| + 1.$$

Wähle nun $\delta := \min \left\{ 1, \frac{\epsilon}{2|x|+1} \right\}$ (min bezeichnet den kleineren der beiden Werte), dann ist für $|y - x| < \delta$ einerseits $|y + x| < 2|x| + 1$ und andererseits $|y - x| < \frac{\epsilon}{2|x|+1}$, also folgt mit (3.1):

$$|y^2 - x^2| = |y + x||y - x| < (2|x| + 1) \frac{\epsilon}{2|x| + 1} = \epsilon,$$

also die behauptete Stetigkeit.

- (2) Die Betragsfunktion $\mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto |x|$, ist stetig: Sei $x \in \mathbb{R}$, $\epsilon > 0$, und $\delta = \epsilon$. Ist dann $|y - x| < \delta$, so gilt mithilfe der Ungleichung $|y - x| \geq ||y| - |x||$ (die leicht aus der Dreiecksungleichung folgt):

$$||y| - |x|| \leq |y - x| < \delta = \epsilon.$$

- (3) Die konstante Funktion $f : x \mapsto c$ ist gewissermaßen die ‚stetigste‘ von allen, denn hier gilt $|f(x) - f(y)| = 0 < \epsilon$ für jede beliebige Wahl von x, y, ϵ, δ .
- (4) Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$,

$$x \mapsto \begin{cases} 0 & \text{falls } x \leq 0; \\ 1 & \text{falls } x > 0 \end{cases}$$

ist stetig an jedem Punkt $x \neq 0$, denn für $0 < \delta < |x|$ und $|y - x| < \delta$ gilt $f(y) - f(x) = 0$; sie ist aber unstetig im Punkt $x = 0$: Sei dazu $\epsilon = 1$ und $\delta > 0$ beliebig, dann

gilt für $y = \frac{\delta}{2}$, daß einerseits $|y - x| < \delta$, aber andererseits $|f(y) - f(0)| = |1 - 0| = 1$, was nicht kleiner ist als ϵ .

- (5) Betrachte $f: \mathbb{R} \rightarrow \mathbb{R}$,

$$x \mapsto \begin{cases} 1 & \text{falls } x \in \mathbb{Q}; \\ 0 & \text{falls } x \in \mathbb{R} \setminus \mathbb{Q}. \end{cases}$$

Wir haben bereits gesehen, daß $\mathbb{Q}$ dicht in $\mathbb{R}$ liegt, und man überlegt sich leicht (sofern man z.B. $\sqrt{2}$ als bekannt voraussetzt), daß auch $\mathbb{R} \setminus \mathbb{Q}$ dicht in $\mathbb{R}$ liegt¹. Daraus folgt: f ist nirgends, also in keinem Punkt, stetig. Sei nämlich $x \in \mathbb{R}$ beliebig und $\epsilon = 1$. Falls $x \in \mathbb{R} \setminus \mathbb{Q}$, so ist $f(x) = 0$, und für jedes $\delta > 0$ existiert $y \in \mathbb{Q}$ mit $|y - x| < \delta$ (dies ist genau die Dichtheitseigenschaft von $\mathbb{Q}$). Dann ist aber $f(y) = 1$ und somit $|f(y) - f(x)| = 1$, was nicht kleiner als ϵ ist. Der Fall $x \in \mathbb{Q}$ geht analog.

- (6) Nun wählen wir den Definitionsbereich $U = \mathbb{Q}$ und betrachten darauf die Funktion $f: \mathbb{Q} \rightarrow \mathbb{R}$,

$$x \mapsto \begin{cases} 0 & \text{falls } x < \sqrt{2}; \\ 1 & \text{falls } x > \sqrt{2}. \end{cases}$$

Wir behaupten die kontraintuitive Aussage, daß f überall (also in jedem Punkt seines Definitionsbereichs) stetig ist. In der Tat: Sei $\epsilon > 0$ und $x \in \mathbb{Q}$ beliebig, dann ist $\delta := |x - \sqrt{2}| > 0$ (denn $\sqrt{2}$ ist irrational), und für jedes $y \in \mathbb{Q}$ mit $|x - y| < \delta$ ist daher $f(y) = f(x)$.

- (7) Die Funktion $f: \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$, $x \mapsto \frac{1}{x}$ ist überall stetig: Sei nämlich $\epsilon > 0$ und $x \neq 0$, so gilt die Gleichheit

$$\left| \frac{1}{x} - \frac{1}{y} \right| = \frac{|y - x|}{|x||y|}. \quad (3.2)$$

Wir wählen $\delta := \min \left\{ \frac{|x|}{2}, \frac{1}{2}|x|^2\epsilon \right\}$; ist dann nämlich $|y - x| < \delta$, so ist einerseits $|y| > \frac{|x|}{2}$ und andererseits $|y - x| < \frac{1}{2}|x|^2\epsilon$, also folgt nach (3.2):

$$\left| \frac{1}{x} - \frac{1}{y} \right| = \frac{|y - x|}{|x||y|} < \frac{\frac{1}{2}|x|^2\epsilon}{\frac{|x|^2}{2}} = \epsilon.$$

Man beachte, daß man (für gegebenes ϵ) δ immer kleiner wählen muß, je näher x an die Null rückt. Gewissermaßen verschlechtert sich die Stetigkeit der Funktion, je näher man der Null kommt.

Beachte außerdem, daß es keine stetige Fortsetzung von f auf ganz $\mathbb{R}$ gibt. Das bedeutet: Die Funktion $\tilde{f}: \mathbb{R} \rightarrow \mathbb{R}$,

$$x \mapsto \begin{cases} \frac{1}{x} & \text{falls } x \neq 0; \\ c & \text{falls } x = 0 \end{cases}$$

ist für jede Wahl von $c \in \mathbb{R}$ im Punkte $x = 0$ unstetig (wie man sich überlegen möge).

Die Intuition, eine Funktion sei stetig, wenn man ihren Graphen zeichnen könne, ohne den Stift abzusetzen, ist also nicht ganz verkehrt, aber auch nicht ganz korrekt, wie die letzten beiden Beispiele zeigen.

¹Nämlich so: Seien $x, y \in \mathbb{R}$ mit $x < y$. Da $\mathbb{Q}$ dicht in $\mathbb{R}$ liegt, gibt es $x', y' \in \mathbb{Q}$ mit $x \leq x' < y' \leq y$. Setze $z := x' + \frac{\sqrt{2}}{2}(y' - x')$. Offenbar ist $z \notin \mathbb{Q}$, da sonst $\sqrt{2} \in \mathbb{Q}$ wäre. Da außerdem $0 < \frac{\sqrt{2}}{2} < 1$, gilt $x < z < y$.

3.1.2. Folgenstetigkeit. Die vorangehenden Beispiele zeigen, daß die ϵ - δ -Definition der Stetigkeit nicht besonders handlich ist: Selbst für recht einfache Funktionen wie $x \mapsto x^2$ oder $x \mapsto \frac{1}{x}$ war ein gewisser Aufwand nötig, um die Stetigkeit zu zeigen. Wir geben eine einfachere Möglichkeit an, die Stetigkeit einer Funktion zu prüfen. Dazu benötigen wir einen Konvergenzbegriff für Funktionen:

DEFINITION 3.3 (Grenzwerte für Funktionen). Sei $f : U \rightarrow \mathbb{R}$ und $x^0 \in \mathbb{R}$ dergestalt, daß eine Folge $(x_n)_{n \in \mathbb{N}} \subset U$ mit $\lim_{n \rightarrow \infty} x_n = x^0$ existiert². Wir schreiben

$$\lim_{x \rightarrow x^0} f(x) = a,$$

wenn für jede Folge $(x_n)_{n \in \mathbb{N}} \subset U$, für die $\lim_{n \rightarrow \infty} x_n = x^0$, auch $\lim_{n \rightarrow \infty} f(x_n) = a$ gilt. Hierbei sind auch die Werte $a = \pm \infty$ zugelassen.

Man schreibt $\lim_{x \rightarrow \pm \infty} f(x) = a$, falls für jede bestimmt gegen $\pm \infty$ konvergente Folge $(x_n)_{n \in \mathbb{N}} \subset U$ gilt: $\lim_{n \rightarrow \infty} f(x_n) = a$. Auch hier ist $a = \pm \infty$ zulässig.

Man kann Grenzwerte auch ‚von rechts‘ oder ‚von links‘ nehmen, d.h. in obiger Definition betrachtet man nur solche Folgen, die in monoton fallender oder steigender Weise gegen x^0 konvergieren. Man schreibt dafür dann

$$\lim_{x \searrow x^0} f(x) \quad \text{bzw.} \quad \lim_{x \nearrow x^0} f(x).$$

BEISPIEL 3.4. Es ist

$$\lim_{x \rightarrow \pm \infty} \frac{1}{x} = 0, \quad \lim_{x \searrow 0} \frac{1}{x} = +\infty, \quad \lim_{x \nearrow 0} \frac{1}{x} = -\infty, \quad \lim_{x \rightarrow 1} \frac{1}{x} = 1.$$

Man sieht dies folgendermaßen ein: Ist etwa $(x_n)_{n \in \mathbb{N}}$ eine bestimmt gegen $+\infty$ divergente Folge, so konvergiert die Folge $\left(\frac{1}{x_n}\right)_{n \in \mathbb{N}}$ gegen null. (Denn für $\epsilon > 0$ sei M so groß, daß $\frac{1}{M} < \epsilon$; dann gibt es wegen der bestimmten Divergenz ein $N \in \mathbb{N}$, sodaß $x_n > M$ für alle $n \geq N$, und daher auch $\frac{1}{x_n} < \frac{1}{M} < \epsilon$ für solche n gemäß Korollar 1.41.)

Dies zeigt die erste Aussage, die anderen beiden sieht man ähnlich ein. Ist $(x_n)_{n \in \mathbb{N}}$ eine gegen 1 konvergente Folge, so konvergiert auch $\left(\frac{1}{x_n}\right)_{n \in \mathbb{N}}$ gegen 1 nach Satz 2.9.

SATZ 3.5 (Folgenstetigkeit). Sei $f : U \rightarrow \mathbb{R}$ eine Funktion und $x^0 \in U$. Dann ist f genau dann stetig in x^0 , wenn

$$\lim_{x \rightarrow x^0} f(x) = f(x^0).$$

Man kann also einen Limes in eine stetige Funktion ‚reinziehen‘.

BEWEIS. ‚ $\Rightarrow$ ‘: Sei f stetig in x^0 und $(x_n)_{n \in \mathbb{N}} \subset U$ eine gegen x^0 konvergente Folge. Sei $\epsilon > 0$, so gibt es nach Annahme ein $\delta > 0$, sodaß $|x^0 - y| < \delta$ impliziert: $|f(x^0) - f(y)| < \epsilon$. Da die Folge gegen x^0 konvergiert, gibt es aber ein $N \in \mathbb{N}$, sodaß $|x_n - x^0| < \delta$ für alle $n \geq N$. Für solche n gilt also auch $|f(x^0) - f(x_n)| < \epsilon$, und dies zeigt, da $\epsilon > 0$ beliebig war, daß $\lim_{x \rightarrow x^0} f(x) = f(x^0)$.

‚ $\Leftarrow$ ‘: Sei umgekehrt $\lim_{x \rightarrow x^0} f(x) = f(x^0)$. Angenommen f wäre nicht stetig in x^0 , so gäbe es ein $\epsilon > 0$ und für jedes $n \in \mathbb{N}$ ein $x_n \in U$, sodaß zwar $|x_n - x^0| < \frac{1}{n}$, aber $|f(x_n) - f(x^0)| \geq \epsilon$. Die Folge $(x_n)_{n \in \mathbb{N}}$ konvergiert nach Wahl gegen x^0 , aber $(f(x_n))_{n \in \mathbb{N}}$ konvergiert nicht gegen $f(x^0)$, im Widerspruch zur Annahme. $\square$

Sind f und g zwei Funktionen mit gemeinsamem Definitionsbereich $U \subset \mathbb{R}$, so definieren wir die Funktionen $f \pm g$ und fg punktweise, das heißt durch die Abbildungsvorschriften

$$(f \pm g)(x) := f(x) \pm g(x), \quad (fg)(x) := f(x)g(x).$$

²In diesem Fall nennt man x^0 einen *Häufungspunkt* von U , und die Menge der Häufungspunkte von U heißt (*topologischer*) *Abschluß* von U . Der Abschluß einer Menge ist offenbar eine Obermenge ihrer selbst.

Auf dem eingeschränkten Definitionsbereich $\{x \in U : g(x) \neq 0\}$ definiert man außerdem

$$\left(\frac{f}{g}\right)(x) := \frac{f(x)}{g(x)}.$$

KOROLLAR 3.6. Seien $f, g : U \rightarrow \mathbb{R}$ stetig im Punkt $x^0 \in U$. Dann sind auch $f \pm g$, fg und, falls $g(x^0) \neq 0$, auch $\frac{f}{g}$ stetig im Punkt x^0 .

BEWEIS. Dies folgt unmittelbar aus Sätzen 3.5 und 2.9. $\square$

Eine Funktion $\mathbb{R} \rightarrow \mathbb{R}$ der Form $x \mapsto \sum_{k=0}^N a_k x^k$ heißt *Polynomfunktion*, wobei die a_k reelle Koeffizienten sind. Sind f und g Polynomfunktionen und $U := \{x \in \mathbb{R} : g(x) \neq 0\}$, so heißt $\frac{f}{g} : U \rightarrow \mathbb{R}$ *rationale Funktion*. Durch mehrmalige Anwendung von Korollar 3.6 und unter Berücksichtigung der Stetigkeit der konstanten Funktionen und der Identität $x \mapsto x$ erhalten wir sofort:

KOROLLAR 3.7. Rationale Funktionen sind in ihrem Definitionsbereich überall stetig.

Dies betrifft insbesondere alle Polynomfunktionen und Potenzfunktionen $x \mapsto x^k$, $k \in \mathbb{Z}$.

KOROLLAR 3.8 (Kompositionen stetiger Funktionen sind stetig). Seien $U, V \subset \mathbb{R}$ und $f : U \rightarrow V$ stetig in $x \in U$ sowie $g : V \rightarrow \mathbb{R}$ stetig in $f(x) \in V$. Dann ist $g \circ f : U \rightarrow \mathbb{R}$ stetig in x .

BEWEIS. Sei $(x_n)_{n \in \mathbb{N}} \subset U$ eine Folge mit $\lim_{n \rightarrow \infty} x_n = x$, so gilt wegen der Stetigkeit von f in x und Satz 3.5 auch $\lim_{n \rightarrow \infty} f(x_n) = f(x)$; erneute Anwendung des Satzes auf g ergibt dann $\lim_{n \rightarrow \infty} g(f(x_n)) = g(f(x))$, also wiederum nach Satz 3.5 die Stetigkeit von $g \circ f$ in x . $\square$

3.1.3. Stetigkeit und offene Mengen. Wir geben eine dritte Charakterisierung der Stetigkeit, die die allgemeinste ist, weil sie nicht nur in $\mathbb{R}$, sondern in beliebigen *topologischen Räumen* formuliert werden kann. Wir diskutieren hier (noch) nicht topologische Räume, sondern bleiben in $\mathbb{R}$.

Sei $x \in \mathbb{R}$, dann nennen wir die Menge $B_\epsilon(x) := \{y \in \mathbb{R} : |y - x| < \epsilon\}$ die ϵ -*Umgebung* von x . Es handelt sich in anderen Worten um das offene Intervall $(x - \epsilon, x + \epsilon)$.

DEFINITION 3.9 (offene Teilmengen). Sei $U \subset \mathbb{R}$ beliebig. Eine Teilmenge $V \subset U$ heißt *offen* bezüglich U , wenn es zu jedem $x \in V$ ein $\epsilon > 0$ gibt, sodaß $B_\epsilon(x) \cap U \subset V$.

So sind offene Intervalle der Form $(a, b) := \{x \in \mathbb{R} : a < x < b\}$ auch in diesem Sinne offen bezüglich $\mathbb{R}$: Wähle für $x \in (a, b)$ einfach $\epsilon := \min\{x - a, b - x\} > 0$.

Abgeschlossene Intervalle der Form $[a, b] := \{x \in \mathbb{R} : a \leq x \leq b\}$ sind dagegen nicht offen bezüglich $\mathbb{R}$, denn jede Umgebung eines der Randpunkte a oder b enthält Punkte, die nicht in $[a, b]$ enthalten sind.

Dagegen ist jede Menge offen bezüglich sich selbst, und die leere Menge ist offen bezüglich jeder Menge.

Als weiteres Beispiel betrachte $U = [0, 2) := \{x \in \mathbb{R} : 0 \leq x < 2\}$ und $V = [0, 1]$ bzw. $W = [0, 1]$. Dann ist V bezüglich U offen (obwohl es, als Intervall in $\mathbb{R}$ betrachtet, nur als ‚halboffen‘ bezeichnet wird), aber W ist bezüglich U nicht offen. Bezüglich $\mathbb{R}$ sind weder V noch W offen.

SATZ 3.10 (Topologische Charakterisierung der Stetigkeit). Sei $U \subset \mathbb{R}$. Eine Funktion $f : U \rightarrow \mathbb{R}$ ist überall in U stetig genau dann, wenn das Urbild jeder offenen Menge bzgl. $\mathbb{R}$ unter f seinerseits offen bzgl. U ist.

BEWEIS. „ $\Rightarrow$ “: Sei f überall stetig und $V \subset \mathbb{R}$ offen (bzgl. $\mathbb{R}$). Wir müssen zeigen, daß dann $f^{-1}(V) = \{x \in U : f(x) \in V\}$ selbst offen ist. Sei dazu $x \in f^{-1}(V)$, das bedeutet, $f(x) \in V$. Da V offen ist, gibt es ein $\epsilon > 0$, sodaß $B_\epsilon(f(x)) \subset V$. Da aber f in x stetig

ist, gibt es nach Definition 3.1 zu diesem ϵ ein $\delta > 0$, sodaß aus $y \in B_\delta(x) \cap U$ folgt: $f(y) \in B_\epsilon(f(x)) \subset V$. Also ist für $y \in B_\delta(x) \cap U$ auch $y \in f^{-1}(V)$, und somit ist $f^{-1}(V)$ offen.

, $\Leftarrow$ ': Sei nun umgekehrt das Urbild jeder offenen Menge unter f offen und sei $x \in U$. Wir müssen zeigen, daß f in x stetig ist. Sei dazu $\epsilon > 0$, dann ist nach Annahme die Menge

$$W := f^{-1}(B_\epsilon(f(x)))$$

offen bzgl. U . Es existiert also ein $\delta > 0$, sodaß $B_\delta(x) \cap U \subset W$. Dies bedeutet: ist $y \in B_\delta(x) \cap U$, so ist $f(y) \in B_\epsilon(f(x))$, und gemäß Definition 3.1 heißt dies gerade, daß f in x stetig ist. $\square$

Man beachte: Ersetzt man in obigem Satz ‚Urbild‘ durch ‚Bild‘, so entsteht eine falsche Aussage. Die konstante Funktion $x \mapsto 1$ etwa ist auf ganz $\mathbb{R}$ stetig, aber das Bild der offenen Menge $\mathbb{R}$ unter dieser Funktion ist die einelementige Menge $\{1\}$, die *nicht* offen ist.

3.2. Eigenschaften stetiger Funktionen

3.2.1. Zwischenwertsatz und Wurzelziehen.

3.2.1.1. Zwischenwertsatz.

SATZ 3.11 (Zwischenwertsatz). Sei f stetig auf dem abgeschlossenen Intervall $[a, b]$, wobei $a, b \in \mathbb{R}$ und $a < b$. Ist $f(a) \geq 0$ und $f(b) \leq 0$, so existiert eine Nullstelle $x^0 \in [a, b]$, das heißt $f(x^0) = 0$.

Das Gleiche gilt, falls $f(a) \leq 0$ und $f(b) \geq 0$.

BEWEIS. Wir zeigen nur die erste Aussage, da dann die zweite sofort durch Übergang zu $-f$ folgt.

Wir verwenden Intervallschachtelung. Sei dazu (wie im Beweis von Satz 2.36) für ein abgeschlossenes Intervall $I := [c, d]$

$$I^+ := \left[\frac{c+d}{2}, d \right], \quad I^- := \left[c, \frac{c+d}{2} \right].$$

Wir konstruieren eine Folge $(I_n)_{n \in \mathbb{N}}$ von Intervallen, sodaß $I_n \supset I_{n+1}$ und $|I_n| \rightarrow 0$ für $n \rightarrow \infty$, und so, daß für $I_n := [a_n, b_n]$ gilt: $f(a_n) \geq 0$ und $f(b_n) \leq 0$.

Dazu setzen wir $I_0 := [a, b]$ und, falls I_n bereits bekannt ist für ein $n \in \mathbb{N}$, $I_{n+1} := I_n^+$ oder $I_{n+1} := I_n^-$, je nachdem, ob

$$f\left(\frac{a_n + b_n}{2}\right) \geq 0 \quad \text{oder} \quad f\left(\frac{a_n + b_n}{2}\right) \leq 0$$

(im Falle der Gleichheit ist die Wahl von I_n^+ und I_n^- gleichermaßen zulässig).

Es ist klar, daß die so konstruierte Folge die gewünschten Eigenschaften hat. Nach Intervallschachtelungsprinzip (Satz 2.25) existiert genau ein $x^0 \in [a, b]$, sodaß $x^0 \in I_n$ für alle $n \in \mathbb{N}$. Es bleibt zu zeigen $f(x^0) = 0$.

Sei dazu $\epsilon > 0$ und $\delta > 0$ so klein, daß aus $|x^0 - y| < \delta$ folgt $|f(x^0) - f(y)| < \epsilon$. Solch ein δ existiert, da f stetig ist. Wähle außerdem $N \in \mathbb{N}$ so groß, daß $|I_N| < \delta$. Dann gilt insbesondere $|a_N - x^0| < \delta$ und $|x^0 - b_N| < \delta$ und somit

$$0 \leq f(a_N) \leq f(x^0) + \epsilon \quad \text{sowie} \quad 0 \geq f(b_N) \geq f(x^0) - \epsilon,$$

also $|f(x^0)| < \epsilon$. Da aber $\epsilon > 0$ beliebig war, folgt wie gewünscht $f(x^0) = 0$. $\square$

Beispielsweise hat jede Polynomfunktion ungerader Ordnung, also jede Funktion der Form

$$x \mapsto \sum_{j=0}^n a_j x^j, \quad n \text{ ungerade,} \quad a_j \in \mathbb{R} \quad \text{für alle } j = 1, \dots, n, \quad a_n \neq 0,$$

mindestens eine reelle Nullstelle (Übung). Die Polynomfunktion $x \mapsto x^2 + 1$ hat dagegen keine reelle Nullstelle, wohl aber (wie wir sehen werden) eine komplexe. Es hat sogar *jede* Polynomfunktion eine komplexe Nullstelle – dies ist die Aussage des *Fundamentalsatzes der Algebra*, den Sie in einer fortgeschrittenen Vorlesung kennenlernen werden (z.B. Funktionentheorie).

KOROLLAR 3.12. Sei f stetig auf dem abgeschlossenen Intervall $[a, b]$, wobei $a, b \in \mathbb{R}$ und $a < b$. Dann nimmt f jeden Wert zwischen $f(a)$ und $f(b)$ an.

BEWEIS. Sei $p \in \mathbb{R}$ eine Zahl zwischen $f(a)$ und $f(b)$ (d.h. im Intervall $[f(a), f(b)]$, falls $f(a) \leq f(b)$, und andernfalls im Intervall $[f(b), f(a)]$). Die Funktion $g := f - p$ ist stetig in $[a, b]$ und erfüllt $g(a) \geq 0$ und $g(b) \leq 0$, oder $g(a) \leq 0$ und $g(b) \geq 0$. Nach Zwischenwertsatz existiert $x \in [a, b]$ mit $g(x) = 0$, also $f(x) = p$. $\square$

3.2.1.2. Wurzeln. Zur Vorbereitung geben wir die folgende einleuchtende Definition:

DEFINITION 3.13 (Monotonie). Eine Funktion $f: \mathbb{R} \supset U \rightarrow \mathbb{R}$ heißt *monoton wachsend* (oder *steigend*), falls aus $x, y \in U$ und $x < y$ folgt, daß $f(x) \leq f(y)$. Sie heißt *monoton fallend*, falls $-f$ monoton wachsend ist.

Sie heißt sogar *streng monoton wachsend*, falls aus $x, y \in U$ und $x < y$ folgt, daß $f(x) < f(y)$, und *streng monoton fallend*, falls $-f$ streng monoton wachsend ist.

SATZ 3.14 (Wurzeln). Sei $k \in \mathbb{N}$ und $a \in \mathbb{R}$ mit $a \geq 0$. Dann existiert genau eine reelle Zahl x mit der Eigenschaft $x^k = a$. Diese wird als k -te Wurzel aus a bezeichnet, und man schreibt $x = \sqrt[k]{a}$.

BEWEIS. Betrachte die Funktion $f: [0, \infty) \rightarrow \mathbb{R}$, $x \mapsto x^k - a$. Da es sich um eine Polynomfunktion handelt, ist sie nach Korollar 3.7 überall stetig. Es gilt $f(0) = -a \leq 0$; da außerdem $\lim_{x \rightarrow \infty} f(x) = +\infty$, existiert ein $b > 0$, sodaß $f(b) \geq 0$. Nach Zwischenwertsatz existiert eine Nullstelle $x \geq 0$, und diese erfüllt natürlich $x^k = a$.

Es bleibt die Eindeutigkeit zu zeigen. Dazu zeigen wir zunächst, daß f streng monoton wächst. Seien dazu $y, z \geq 0$ mit $y < z$. Die k -malige Anwendung von Korollar 1.41(3) ergibt $y^k < z^k$, und somit auch $y^k - a < z^k - a$.

Daraus folgt nun unmittelbar die Eindeutigkeit der Nullstelle, d.h. der k -ten Wurzel: Denn für jedes $0 \leq x' < x$ gilt $f(x') < f(x) = 0$ und für jedes $x < x'$ gilt $f(x') > f(x) = 0$, also ist x die einzige nichtnegative Zahl, für die gilt $x^k = a$. $\square$

Wie steht es mit etwaigen *negativen* Lösungen der Gleichung $x^k = a$? Hier muß man unterscheiden, ob k gerade oder ungerade ist. Falls k gerade, so ist $(-x)^k = x^k$ für alle $x \in \mathbb{R}$, und somit gibt es genau zwei Lösungen der Gleichung $x^k = a$, nämlich $\pm \sqrt[k]{a}$. Ist k dagegen ungerade, so ist $(-x)^k = -x^k$ negativ für alle negativen x , und es existiert daher keine negative Lösung von $x^k = a$ (wir hatten ja $a \geq 0$ vorausgesetzt).

3.2.2. Weitere Eigenschaften stetiger Funktionen.

3.2.2.1. Intervalle. Wir haben bereits mit Intervallen gearbeitet, klären jetzt aber nochmal genau die Terminologie. Ein *Intervall* ist eine *zusammenhängende* Teilmenge $I \subset \mathbb{R}$, das bedeutet: Sind $x, z \in I$ und ist $x \leq y \leq z$, so folgt auch $y \in I$. Offenbar hat jedes Intervall eine der Formen

$$\begin{aligned} [a, b] &:= \{x \in \mathbb{R} : a \leq x \leq b\}, & (a, b] &:= \{x \in \mathbb{R} : a < x \leq b\}, \\ [a, b) &:= \{x \in \mathbb{R} : a \leq x < b\}, & (a, b) &:= \{x \in \mathbb{R} : a < x < b\}, \end{aligned}$$

wobei a, b entweder reelle Zahlen mit $a \leq b$ sind oder $a = -\infty$ oder $b = +\infty$; man meint damit zum Beispiel $(-\infty, 0] = \{x \in \mathbb{R} : x \leq 0\}$. Falls eine der Intervallgrenzen $\pm\infty$ ist, ist $\pm\infty$ also nicht Element des Intervalls³, und deshalb schreibt man zum Beispiel $(-\infty, 0]$ und

³Das wäre auch gar nicht möglich, denn $\pm\infty$ sind nicht als mathematische Objekte definiert; siehe aber Bemerkung 2.38.

nicht $[-\infty, 0]$. Es gilt natürlich $(-\infty, +\infty) = \mathbb{R}$. Ist $a \in \mathbb{R}$, so ist auch $[a, a]$ zulässig (es ist dies die Menge mit dem einzigen Element a).

Sind a und b reelle Zahlen, nennt man die Intervalle $[a, b]$, $(a, b]$, $[a, b)$, (a, b) *beschränkt*; ist mindestens eine der Intervallgrenzen unendlich, so heißt das Intervall *unbeschränkt*. Sind $a, b \in \mathbb{R}$, so heißen die Intervalle $[a, b]$, $(-\infty, b]$, $[a, +\infty)$ und $(-\infty, +\infty)$ *abgeschlossen* und die Intervalle (a, b) , $(-\infty, b)$, $(a, +\infty)$ und $(-\infty, +\infty)$ *offen*, was konsistent ist mit Definition 3.9 (mit $U = \mathbb{R}$). Man beachte, daß $(-\infty, +\infty)$ offen und abgeschlossen zugleich ist. Alle anderen Intervalltypen heißen *halboffen*.

Ist $a > b$, so vereinbaren wir⁴ $(a, b) := (b, a)$ und analog für abgeschlossene und halboffene Intervalle.

Ist ein Intervall abgeschlossen und beschränkt, so heißt es *kompakt*. So sind etwa $(-\infty, 0]$ und $[-1, 0]$ beide abgeschlossen, aber nur das Letztere ist kompakt.

Abgeschlossene Intervalle haben folgende gute Eigenschaft: Ist I abgeschlossen und $(x_n)_{n \in \mathbb{N}} \subset I$ eine konvergente Folge, so liegt der Grenzwert seinerseits in I . Dies folgt aus Satz 2.13 (wobei eine der beiden Folgen konstant gleich einem Intervallrand gewählt wird). Mengen, die nicht abgeschlossen sind, haben diese Eigenschaft nicht: Zum Beispiel ist $\frac{1}{n} \in (0, 1]$ für alle n , aber der Grenzwert Null ist nicht in $(0, 1]$ enthalten.

Es mag verwundern, warum wir für solche einfachen Konzepte soviel Terminologie einführen; dies wird später klarer werden, wenn wir im Rahmen der Topologie Begriffe wie offen, abgeschlossen oder kompakt auf andere Mengen als $\mathbb{R}$ verallgemeinern.

PROPOSITION 3.15. *Ist $I \subset \mathbb{R}$ ein Intervall und $f : I \rightarrow \mathbb{R}$ stetig, so ist das Bild $f(I)$ wieder ein Intervall.*

BEWEIS. Seien $p, q \in f(I)$, das heißt, es gibt $x, y \in I$ mit $f(x) = p$ und $f(y) = q$. Dann ist $[x, y] \subset I$, und nach Korollar 3.12 nimmt f jeden Wert in $[p, q]$ an. Damit ist $f(I)$ ein Intervall. $\square$

3.2.2.2. Stetige Funktionen auf kompakten Intervallen.

SATZ 3.16 (Stetige Bilder kompakter Intervalle sind kompakt). Sei $[a, b]$ ein kompaktes Intervall und $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann ist $f(I)$ wieder ein kompaktes Intervall.

BEWEIS. Nach Proposition 3.15 ist $f(I)$ ein Intervall. Wir müssen zeigen, daß es abgeschlossen und beschränkt ist.

Zur Beschränktheit: Angenommen, dies wäre nicht der Fall, und der rechte Intervallrand von $f(I)$ wäre $+\infty$. Dann gäbe es eine Folge $(x_n)_{n \in \mathbb{N}} \subset I$, sodaß $\lim_{n \rightarrow \infty} f(x_n) = +\infty$. Da die Folge $(x_n)_{n \in \mathbb{N}}$ beschränkt ist (denn I ist ja beschränkt), gibt es nach dem Satz von BOLZANO-WEIERSTRASS eine konvergente Teilfolge $(x_{n_k})_{k \in \mathbb{N}}$ mit Limes x , der selbst in I liegt (da I abgeschlossen ist). Nach Stetigkeit von f gilt dann aber

$$\lim_{k \rightarrow \infty} f(x_{n_k}) = f(x) \in \mathbb{R},$$

im Widerspruch zu $\lim_{n \rightarrow \infty} f(x_n) = +\infty$. Also ist die rechte Intervallgrenze von $f(I)$ endlich. Analog zeigt man, daß auch die linke Intervallgrenze endlich ist.

Wir zeigen nun, daß $f(I)$ auch abgeschlossen ist. Sei dazu $p \in \mathbb{R}$ der rechte Intervallrand von $f(I)$, das heißt $p := \sup f(I)$, und sei $(p_n)_{n \in \mathbb{N}} \subset f(I)$ eine Folge mit $\lim_{n \rightarrow \infty} p_n = p$. Eine solche Folge existiert stets, siehe Übungsblatt 6 Aufgabe 3. Dann gibt es zu jedem p_n ein $x_n \in I$ mit $f(x_n) = p_n$, und nach BOLZANO-WEIERSTRASS und der Abgeschlossenheit von I existiert eine Teilfolge $(x_{n_k})_{k \in \mathbb{N}} \subset I$, die gegen ein $x \in I$ konvergiert. Da f stetig ist, folgt

$$f(x) = \lim_{k \rightarrow \infty} f(x_{n_k}) = \lim_{k \rightarrow \infty} p_{n_k} = p$$

⁴ganz unter uns, denn diese Konvention ist nicht Standard.

und insbesondere $p \in f(I)$. Die rechte Intervallgrenze ist also Element von $f(I)$, und für die linke Intervallgrenze zeigt man dies analog. Damit ist alles gezeigt. $\square$

Man vergleiche mit Satz 3.10: Urbilder offener Mengen unter stetigen Funktionen sind stets offen, wohingegen Bilder kompakter Mengen unter stetigen Funktionen stets kompakt sind. Wer Schwierigkeiten hat, sich das zu merken, denke an eine konstante Funktion.

Wir nennen eine Funktion $f : I \rightarrow \mathbb{R}$ *beschränkt*, wenn die Menge $f(I)$ beschränkt ist, wenn es also ein $M > 0$ gibt mit $|f(x)| \leq M$ für alle $x \in I$.

KOROLLAR 3.17 (Maximum und Minimum stetiger Funktionen auf kompakten Intervallen). Sei $I = [a, b]$ ein kompaktes Intervall und $f : I \rightarrow \mathbb{R}$ stetig. Dann ist f beschränkt und nimmt sein Maximum und sein Minimum an, d.h. es existiert ein $\bar{x} \in I$ mit $f(\bar{x}) = \max f(I)$ und ein $\underline{x} \in I$ mit $f(\underline{x}) = \min f(I)$.

BEWEIS. Da nach Satz 3.16 $f(I)$ ein kompaktes Intervall ist, ist $f(I)$ insbesondere beschränkt, und es gilt $\sup f(I) \in f(I)$, also existiert $\bar{x} \in I$ mit $f(\bar{x}) = \sup f(I) = \max f(I)$. Analog für das Minimum. $\square$

Die Kompaktheit des Intervalls ist entscheidend: Betrachte etwa die Funktionen $f, g : (0, 1) \rightarrow \mathbb{R}$, $f : x \mapsto x^2$, $g : x \mapsto \frac{1}{x}$. Dann ist $\sup f((0, 1)) = 1$, aber es existiert keine Zahl $x \in (0, 1)$, für die $x^2 = 1$. Das Supremum ist also kein Maximum. Die Funktion g ist noch nicht einmal beschränkt.

Zur Notation: Anstatt $\max f(I)$ schreibt man oft auch $\max_{x \in I} f(x)$, und analog für $\min$, $\sup$, $\inf$.

3.2.3. Gleichmäßige Stetigkeit. Ganz zu Beginn der Vorlesung ist darauf hingewiesen worden, daß Existenz- und Allquantoren in prädikatenlogischen Aussagen nicht vertauscht werden dürfen (vgl. das Beispiel mit den Töpfen und Deckeln). Der Unterschied zwischen Stetigkeit und gleichmäßiger Stetigkeit beruht genau auf dieser Vertauschung:

DEFINITION 3.18 (gleichmäßige Stetigkeit). Sei $U \subset \mathbb{R}$ und $f : U \rightarrow \mathbb{R}$. Die Funktion f heißt auf U *gleichmäßig stetig*, wenn gilt:

$$\forall \epsilon > 0 \quad \exists \delta > 0 \quad \forall x \in U \quad \forall y \in U : \quad |x - y| < \delta \implies |f(x) - f(y)| < \epsilon.$$

Man vergleiche dies mit der Definition der Stetigkeit: Eine Funktion ist stetig in U , falls

$$\forall \epsilon > 0 \quad \forall x \in U \quad \exists \delta > 0 \quad \forall y \in U : \quad |x - y| < \delta \implies |f(x) - f(y)| < \epsilon.$$

Der Unterschied besteht also ‚nur‘ in der Vertauschung von $\forall x$ und $\exists \delta$. Für die gleichmäßige Stetigkeit muß δ unabhängig von x (also ‚gleichmäßig in x ‘) gewählt werden, wohingegen das δ im Falle gewöhnlicher Stetigkeit durchaus von x abhängen darf.⁵ Insbesondere ist jede gleichmäßig stetige Funktion stetig, aber nicht umgekehrt:

BEISPIEL 3.19. In Beispiel 3.2 hatten wir die Funktion $f : \mathbb{R} \setminus \{0\}$, $x \mapsto \frac{1}{x}$ als stetig identifiziert. Wir wählten dazu $\delta := \min \left\{ \frac{|x|}{2}, \frac{1}{2}|x|^2 \epsilon \right\}$. Diese Wahl von δ ist offenbar von x abhängig: Je näher x bei null liegt, desto kleiner wird dieses δ .

Wir zeigen, daß diese Funktion *nicht* gleichmäßig stetig in $\mathbb{R} \setminus \{0\}$ ist. Wähle dazu $\epsilon = 1$ und sei $\delta > 0$ beliebig. Wenn wir zu diesem δ zwei Zahlen $x_\delta, y_\delta \in \mathbb{R} \setminus \{0\}$ finden, sodaß $|x_\delta - y_\delta| < \delta$, aber $\left| \frac{1}{x_\delta} - \frac{1}{y_\delta} \right| \geq 1$, sind wir fertig.

⁵Ebenso müßte für die Aussage „Es gibt einen Deckel, der auf jeden Topf paßt“ ein passender Deckel unabhängig vom gewählten Topf gefunden werden, wohingegen für die Aussage „Auf jeden Topf paßt ein Deckel“ der Deckel natürlich je nach Topf unterschiedlich gewählt werden wird.

Man prüft leicht nach, daß die Wahl $x_\delta := \delta$, $y_\delta := \frac{1}{2}\delta$ das gewünschte Ergebnis liefert: Es ist nämlich einerseits $|x_\delta - y_\delta| = \frac{\delta}{2} < \delta$ und andererseits

$$\left| \frac{1}{x_\delta} - \frac{1}{y_\delta} \right| = \frac{1}{\delta} \geq 1,$$

sofern $\delta \leq 1$. Ist dagegen $\delta > 1$, so wähle einfach $x_\delta = \frac{1}{2}$, $y_\delta = \frac{3}{2}$. Dann ist nämlich $|x_\delta - y_\delta| = 1 < \delta$ und

$$\left| \frac{1}{x_\delta} - \frac{1}{y_\delta} \right| = \frac{4}{3} \geq 1.$$

Anschaulich entspricht die mangelnde Gleichmäßigkeit der Stetigkeit dem in der Nähe von null immer steiler werdenden Graphen der Funktion: Minimale Unterschiede in den Eingangsdaten (x und y) führen zu beträchtlichen Änderungen in den Funktionswerten.

Wieder ist es so, daß stetige Funktionen auf kompakten Intervallen sich besonders gut verhalten:

SATZ 3.20 (Stetige Funktionen auf kompakten Intervallen sind gleichmäßig stetig). Sei $I \subset \mathbb{R}$ ein kompaktes Intervall und $f : I \rightarrow \mathbb{R}$ stetig. Dann ist f auf I sogar gleichmäßig stetig.

BEWEIS. Angenommen dies wäre nicht der Fall, so gäbe es ein $\epsilon > 0$, sodaß für jede Wahl von $\delta > 0$ ein $x_\delta \in I$ und ein $y_\delta \in I$ existierte mit $|x_\delta - y_\delta| < \delta$, aber $|f(x_\delta) - f(y_\delta)| \geq \epsilon$. Wähle $\delta := \frac{1}{n}$ und schreibe $x_n := x_\delta$, $y_n := y_\delta$. Auf diese Weise erhalten wir zwei Folgen $(x_n)_{n \in \mathbb{N}}, (y_n)_{n \in \mathbb{N}} \subset I$. Nach dem Satz von BOLZANO-WEIERSTRASS existiert eine konvergente Teilfolge $(x_{n_k})_{k \in \mathbb{N}}$ mit Grenzwert $x \in I$. Dann konvergiert $(y_{n_k})_{k \in \mathbb{N}}$ ebenfalls gegen x wegen $|x_n - y_n| < \frac{1}{n}$.

Nun ist f in x stetig, also existiert $\eta > 0$, sodaß aus $|x - y| < \eta$ folgt $|f(x) - f(y)| < \frac{\epsilon}{2}$. Wähle k so groß, daß $|x - x_{n_k}| < \eta$ und $|x - y_{n_k}| < \eta$, so gilt

$$|f(x_{n_k}) - f(y_{n_k})| \leq |f(x_{n_k}) - f(x)| + |f(x) - f(y_{n_k})| < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon,$$

im Widerspruch zu $|f(x_{n_k}) - f(y_{n_k})| \geq \epsilon$. □

KAPITEL 4

Differentiation und Integration

Ihren historischen Ursprung hat die Analysis im Differential- und Integralkalkül von LEIBNIZ und NEWTON¹. Während dieser den Kalkül zur Formulierung seiner Mechanik entwickelte, war jener durch geometrische Fragestellungen motiviert. Ob letztlich NEWTON oder LEIBNIZ den Infinitesimalkalkül² zuerst entwickelte und wer dementsprechend von wem plagiiert hatte, war lange umstritten und führte zu lächerlichen Auseinandersetzungen zwischen britischen und kontinentalen Wissenschaftlern; heute geht man davon aus, daß beide unabhängig voneinander arbeiteten.

Wir stellen hier die grundlegende Theorie der Differential- und Integralrechnung in einer Dimension vor und exemplifizieren sie an den uns bereits bekannten Typen von Funktionen (Polynome, rationale Funktionen, Wurzeln). Weiteres Beispielmateriale liefern uns dann im nächsten Kapitel die *Potenzreihen*, die besonders angenehm zu differenzieren und integrieren sind, und zu denen die bekannten transzendenten Funktionen wie die Exponentialfunktion, der Logarithmus, Sinus und Kosinus etc. gehören.

Ziel dieser Vorlesung ist es allerdings nicht, Ihnen eine möglichst virtuose Rechentechnik für Ableitungen und Integrale explizit gegebener Funktionen anzutrainieren; solche Fertigkeiten sind im Laufe des vergangenen Jahrhunderts dank numerischer Verfahren und Computeralgebra immer mehr obsolet geworden.

4.1. Ableitungen

4.1.1. Definition und Beispiele.

4.1.1.1. Definition.

DEFINITION 4.1 (Ableitung). Sei $U \subset \mathbb{R}$ eine Teilmenge und $x \in U$ ein Häufungspunkt von $U \setminus \{x\}$, d.h. es existiert eine Folge $(x_n)_{n \in \mathbb{N}} \subset U \setminus \{x\}$ mit $\lim_{n \rightarrow \infty} x_n = x$.

Dann heißt eine Funktion $f: U \rightarrow \mathbb{R}$ *differenzierbar* im Punkt x , falls der Grenzwert

$$f'(x) := \lim_{x' \rightarrow x} \frac{f(x) - f(x')}{x - x'} \quad (4.1)$$

existiert, und $f'(x)$ heißt *Ableitung* von f an der Stelle x .

Falls f in jedem Punkt $x \in U$ differenzierbar ist, so heißt f in U differenzierbar, und die Ableitung $x \mapsto f'(x)$ kann ihrerseits als Funktion $U \rightarrow \mathbb{R}$ betrachtet werden.

Selbstverständlich kann man auch schreiben

$$f'(x) := \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h},$$

mit der Konvention, daß nur Werte $h \neq 0$ mit $x+h \in U$ in der Limesbildung zulässig sind.

¹Die Exhaustionsmethode des ARCHIMEDES zur Berechnung der Kreisfläche enthält allerdings bereits die Grundidee der Integralrechnung

²In der Mathematik und Logik ist *Kalkül* maskulin, in der Umgangssprache neutral.

4.1.1.2. *Interpretation.* Wir bieten drei Interpretationen der Ableitung an: Eine geometrische, eine physikalische und eine approximationstheoretische. Alle drei zeigen auf, daß die Ableitung die *infinitesimale Änderung* des Funktionswerts bei infinitesimalen Änderungen des Arguments angibt.

Zunächst die geometrische Interpretation: Gegeben den Graphen einer Funktion, möchte man die *Tangente* an den Graphen im Punkt $(x, f(x))$ berechnen. Dazu muß man ihre Steigung kennen. Die Steigung m einer linearen Funktion der Form $g: x \mapsto mx + b$ läßt sich bekanntlich als

$$m = \frac{g(x_2) - g(x_1)}{x_2 - x_1}$$

darstellen, was nützlich ist, wenn zwei Funktionswerte $g(x_1)$, $g(x_2)$ bekannt sind. Die Funktion f unseres Interesses ist im Allgemeinen allerdings nicht linear. Man nimmt aber an (was anschaulich plausibel ist), daß sich die Tangentensteigung durch die Steigungen der *Sekanten* durch $(x, f(x))$ und einen nahegelegenen Punkt auf dem Graphen $(x+h, f(x+h))$ annähern läßt. Die Sekantensteigung ist also gegeben durch

$$\frac{\Delta f}{\Delta x} := \frac{f(x+h) - f(x)}{h},$$

und im Limes $h \rightarrow 0$ (falls er existiert) erhält man dann die Tangentensteigung $f'(x)$. Aus dieser Anschauung heraus schreibt man häufig $\frac{df}{dx}$ statt f' bzw. $\left.\frac{df}{dx}\right|_x$ statt $f'(x)$. Die sogenannten ‚Differential‘ df und dx haben dabei keine eigenständige Bedeutung, sondern treten nur als Quotient (Differentialquotient) auf, der wiederum über den Grenzwert (4.1) interpretiert wird.³

Zur physikalischen Interpretation: Es gibt viele physikalische Gesetzmäßigkeiten, die sich (nur) mithilfe von Ableitungen formulieren lassen. Wir diskutieren hier nur die Begriffe der Geschwindigkeit und der Beschleunigung. Angenommen, ein Körper bewegt sich geradlinig (z.B. annähernd ein Zug, oder ein Auto auf einer geraden Autobahn) mit nichtkonstanter Geschwindigkeit. Die landläufige Definition besagt „Geschwindigkeit = Weg durch Zeit“, wobei ‚Weg‘ die zwischen Zeitpunkt t_1 und t_2 zurückgelegte Strecke $\Delta s = s(t_2) - s(t_1)$ und ‚Zeit‘ genauer die Zeitdifferenz $\Delta t = t_2 - t_1$ meint. Die so gewonnene Größe

$$\bar{v} := \frac{\Delta s}{\Delta t}$$

gibt dann die *Durchschnittsgeschwindigkeit* des Körpers im Zeitintervall $[t_1, t_2]$ an. Sind wir jedoch daran interessiert, welche Geschwindigkeit *genau zum Zeitpunkt* t_1 vorlag, messen wir Weg- und Zeitdifferenzen in immer kleineren Abständen von t_1 , nehmen also den Grenzwert

$$v(t_1) = \lim_{t \rightarrow t_1} \frac{s(t_1 - t)}{t_1 - t}.$$

Dieser ist dann (definitionsgemäß) gleich der *Momentangeschwindigkeit*, die (idealerweise) auf dem Tachometer angezeigt wird. In der Mechanik ist also Geschwindigkeit die Ableitung des Weges nach der Zeit, und nur im Falle der gleichförmigen Bewegung ist die Geschwindigkeit konstant und ist dann gerade der Quotient aus Weg und Zeit.

Die *Beschleunigung* ist landläufig bekannt als „Änderung der Geschwindigkeit pro Zeiteinheit“ und ist, aufgrund ähnlicher Erwägungen wie für die Geschwindigkeit selbst, gegeben

³Bis ins 19. Jahrhundert waren die Differentiale, mit denen Mathematiker und Physiker recht erfolgreich jonglierten, Gegenstand erbitterter Diskussionen. Sie wurden bisweilen interpretiert als Größen, die größer null, aber kleiner als jede positive Zahl sein sollen – offenbar ein Widerspruch zur Struktur der reellen Zahlen (jede reelle Zahl, die kleiner als jede positive Zahl ist, ist null oder negativ). Die ‚Entmythologisierung‘ der Differentiale erfolgte erst mit der mathematisch rigorosen Einführung des heute gebräuchlichen Limesbegriffs. Eine eigenständige, mathematisch einwandfreie Bedeutung erhalten Differentiale in der Theorie der *Differentialformen*, die Ihnen vielleicht in Analysis III begegnen werden.

(bzw. definiert) als Ableitung der Geschwindigkeit nach der Zeit, mithin als zweite Ableitung des Weges nach der Zeit.

In der Physik schreibt man Zeitableitungen gerne mit einem Punkt, also z.B. $v = \dot{s}$ oder $m\ddot{s} = F$ (Zweites Newtonsches Gesetz).

Schließlich geben wir eine dritte Interpretation der Ableitung an, nämlich als *lineare Approximation*:

SATZ 4.2. Sei $U \subset \mathbb{R}$ und $x \in U$ ein Häufungspunkt von $U \setminus \{x\}$. Eine Funktion $f : U \rightarrow \mathbb{R}$ ist im Punkte $x \in U$ differenzierbar genau dann, wenn ein $m \in \mathbb{R}$ und eine Funktion $\phi : U \rightarrow \mathbb{R}$ existieren, sodaß

$$f(\xi) = f(x) + m(\xi - x) + \phi(\xi) \quad (4.2)$$

für alle $\xi \in U$, und

$$\lim_{\xi \rightarrow x} \frac{\phi(\xi)}{\xi - x} = 0. \quad (4.3)$$

In diesem Falle gilt $m = f'(x)$.

BEWEIS. Sei f in x differenzierbar, und setze $\phi(\xi) := f(\xi) - f(x) - f'(x)(\xi - x)$. Damit gilt natürlich (4.2) mit $m = f'(x)$, und wir müssen noch (4.3) zeigen. Es gilt aber

$$\lim_{\xi \rightarrow x} \frac{\phi(\xi)}{\xi - x} = \lim_{\xi \rightarrow x} \frac{f(\xi) - f(x) - f'(x)(\xi - x)}{\xi - x} = \lim_{\xi \rightarrow x} \frac{f(\xi) - f(x)}{\xi - x} - f'(x) = 0$$

nach Definition der Ableitung.

Gelte nun umgekehrt (4.2) und (4.3). Dann gilt

$$\lim_{\xi \rightarrow x} \frac{f(\xi) - f(x)}{\xi - x} = m + \lim_{\xi \rightarrow x} \frac{\phi(\xi)}{\xi - x} = m,$$

also ist f in x differenzierbar mit Ableitung m . □

Die Charakterisierung der Differenzierbarkeit in diesem Satz besagt, daß man f durch die linear-affine Funktion⁴ $\xi \mapsto f(x) + f'(x)(\xi - x)$ annähern kann, sofern f in x differenzierbar ist, und daß der Fehler $\phi(\xi)$, den man dabei macht, in der Nähe von x gegenüber dem linearen Term $\xi - x$ vernachlässigbar ist (siehe (4.3)). Man sagt auch, man entwickle f um den Punkt x bis zu erster Ordnung. Eine naheliegende Verallgemeinerung ist die Entwicklung einer Funktion bis zu höherer Ordnung. Dies werden wir später tun, wenn wir den Satz von TAYLOR diskutieren.

KOROLLAR 4.3 (Differenzierbare Funktionen sind stetig). Ist $f : U \rightarrow \mathbb{R}$ in $x \in \mathbb{R}$ differenzierbar, so ist f dort auch stetig.

BEWEIS. Sei $(x_n)_{n \in \mathbb{N}} \subset U$ eine Folge mit $\lim_{n \rightarrow \infty} x_n = x$. O.B.d.A. ist $x_n \neq x$ für alle $n \in \mathbb{N}$. Dann gilt nach dem vorigen Satz für ein $\phi : U \rightarrow \mathbb{R}$ wie in (4.3):

$$f(x_n) = f(x) + f'(x)(x_n - x) + \phi(x_n) = f(x) + f'(x)(x_n - x) + \frac{\phi(x_n)}{x_n - x}(x_n - x) \rightarrow f(x)$$

für $n \rightarrow \infty$, also ist f in x stetig nach Satz 3.5. □

⁴Die Terminologie ist uneinheitlich: Manche Autorinnen bezeichnen nur Funktionen der Form $x \mapsto mx$ als linear und solche der Form $x \mapsto mx + b$ als affin; andere nennen auch letztere linear. Wir schreiben linear-affin, um zu verdeutlichen, daß es uns um Funktionen der Form $x \mapsto mx + b$ geht, also um solche, deren Graph eine Gerade beschreibt.

4.1.1.3. *Beispiele.*

- (1) Die konstante Funktion $\mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto c$ ist überall differenzierbar mit Ableitung null, denn für alle $x \in \mathbb{R}$ und alle $h \neq 0$ ist

$$\frac{f(x+h) - f(x)}{h} = 0.$$

- (2) Die Identität $\iota : \mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto x$ ist überall differenzierbar mit Ableitung 1, denn für alle $x \in \mathbb{R}$ und $h \neq 0$ ist

$$\frac{\iota(x+h) - \iota(x)}{h} = 1.$$

- (3) Sei $k \in \mathbb{N}$ und $f : \mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto x^k$. Mit dem binomischen Lehrsatz (Übungsblatt 3, Aufgabe 7b) berechnen wir für jedes $x \in \mathbb{R}$ und $h \neq 0$:

$$\frac{(x+h)^k - x^k}{h} = \frac{\sum_{j=0}^k \binom{k}{j} x^{k-j} h^j - x^k}{h} = \sum_{j=1}^k \binom{k}{j} x^{k-j} h^{j-1} \rightarrow kx^{k-1}$$

für $h \rightarrow 0$, da alle Terme h^{j-1} bis auf $j = 1$ gegen null konvergieren, und da $\binom{k}{1} = k$. Also erhalten wir die aus der Schule bekannte Ableitungsregel

$$\frac{dx^k}{dx} = kx^{k-1}, \quad k \in \mathbb{N}.$$

- (4) Für $f : \mathbb{R} \setminus \{0\} \rightarrow \mathbb{R}$, $x \mapsto \frac{1}{x}$ und $x, h \neq 0$ ergibt sich

$$\frac{\frac{1}{x+h} - \frac{1}{x}}{h} = \frac{x - (x+h)}{hx(x+h)} = -\frac{1}{x(x+h)} \rightarrow -\frac{1}{x^2}$$

für $h \rightarrow 0$, also ist $\frac{d}{dx} \frac{1}{x} = -\frac{1}{x^2}$.

- (5) Sei $f : \mathbb{R}_0^+ \rightarrow \mathbb{R}$, $x \mapsto \sqrt{x}$ (wobei $\mathbb{R}_0^+$ die Menge der nichtnegativen reellen Zahlen bezeichnet). Wir berechnen zunächst für $x > 0$ und $h \neq 0$

$$\frac{1}{h}(\sqrt{x+h} - \sqrt{x}) = \frac{(\sqrt{x+h} - \sqrt{x})(\sqrt{x+h} + \sqrt{x})}{h(\sqrt{x+h} + \sqrt{x})} = \frac{1}{\sqrt{x+h} + \sqrt{x}} \rightarrow \frac{1}{2\sqrt{x}},$$

also ist die Quadratwurzelfunktion in jedem $x > 0$ differenzierbar mit Ableitung $\frac{1}{2\sqrt{x}}$. In $x = 0$ ist sie allerdings nicht differenzierbar, denn für $h \searrow 0$ gilt⁵

$$\frac{\sqrt{h} - 0}{h - 0} = \frac{1}{\sqrt{h}} \rightarrow +\infty.$$

- (6) Nach Korollar 4.3 sind differenzierbare Funktionen stetig. Die Umkehrung gilt nicht, wie das folgende Beispiel zeigt: Die Betragsfunktion $x \mapsto |x|$ ist zwar auf ganz $\mathbb{R}$ stetig, aber bei $x \neq 0$ nicht differenzierbar. Denn betrachte die Folge $(x_n)_{n \in \mathbb{N}} = ((-1)^n \frac{1}{n})_{n \in \mathbb{N}}$, so gilt

$$\frac{|x_n| - |0|}{x_n - 0} = (-1)^n,$$

und diese Folge ist nicht konvergent.

Definiert man allerdings die *rechts-* bzw. *linksseitige Ableitung* einer Funktion als

$$f'(x+) := \lim_{h \searrow 0} \frac{f(x+h) - f(x)}{h}, \quad f'(x-) := \lim_{h \nearrow 0} \frac{f(x+h) - f(x)}{h},$$

⁵Der letzte Schritt folgt aus $\lim_{h \searrow 0} \sqrt{h} = 0$, was wie folgt begründet werden kann: Für $\epsilon > 0$ wähle $h < \epsilon^2$, dann folgt durch Wurzelziehen auf beiden Seiten wie gewünscht $\sqrt{h} < \epsilon$. Aber warum darf man in einer Ungleichung auf beiden Seiten die Wurzel ziehen? Seien $0 \leq a < b$ und angenommen, $\sqrt{a} \geq \sqrt{b}$, so wäre nach Quadrieren auf beiden Seiten (das ist ja erlaubt!) auch $a \geq b$, Widerspruch.

so haben wir für die Betragsfunktion $f'(0+) = 1$ und $f'(0-) = -1$.

4.1.2. Ableitungsregeln. Wie wir in den Beispielen gesehen haben, ist es mitunter recht mühsam, die Ableitung einer Funktion direkt aus der Definition zu bestimmen. Abhilfe schaffen die hier vorgestellten Regeln.

Wir haben bereits gesehen, daß die Ableitung einer auf $U \subset \mathbb{R}$ differenzierbaren Funktion f selbst wieder als Funktion aufgefaßt werden kann, und wir nennen diese Funktion f' . Summen, Produkte etc. von Funktionen sind wie üblich punktweise definiert, d.h. $(f+g)(x) := f(x) + g(x)$ oder $(cf)(x) := cf(x)$ (vgl. Abschnitt 3.1.2).

4.1.2.1. **LEIBNIZ- und Kettenregel.** Die folgende Beobachtung folgt sofort aus der Definition der Ableitung und Satz 2.9:

PROPOSITION 4.4 (Linearität der Ableitung). *Seien $f, g : U \rightarrow \mathbb{R}$ in $x \in U$ differenzierbar und $\alpha, \beta \in \mathbb{R}$. Dann ist auch die Funktion $\alpha f + \beta g$ in x differenzierbar, und es gilt*

$$(\alpha f + \beta g)'(x) = \alpha f'(x) + \beta g'(x).$$

Damit können wir nun mithilfe der besprochenen Beispiele alle Polynomfunktionen ableiten:

$$\left(\sum_{k=0}^n a_k x^k \right)' = \sum_{k=1}^n k a_k x^{k-1} = \sum_{k=0}^{n-1} (k+1) a_{k+1} x^k.$$

Insbesondere ist die Ableitung eines Polynoms n -ten Grades nur noch ein Polynom $(n-1)$ -ten Grades.

SATZ 4.5 (LEIBNIZ-Regel⁶). *Seien $f, g : U \rightarrow \mathbb{R}$ differenzierbar in $x \in U$. Dann ist auch fg differenzierbar in x , und es gilt*

$$(fg)'(x) = f'(x)g(x) + f(x)g'(x).$$

BEWEIS. Es gilt für $h \neq 0$

$$\begin{aligned} \frac{f(x+h)g(x+h) - f(x)g(x)}{h} &= \frac{f(x+h)g(x+h) - f(x+h)g(x)}{h} + \frac{f(x+h)g(x) - f(x)g(x)}{h} \\ &= f(x+h) \frac{g(x+h) - g(x)}{h} + g(x) \frac{f(x+h) - f(x)}{h} \\ &\rightarrow f(x)g'(x) + g(x)f'(x) \end{aligned}$$

für $h \rightarrow 0$. Beim Grenzübergang für den ersten Summanden haben wir die Stetigkeit von f in x verwendet (Korollar 4.3). $\square$

SATZ 4.6 (Kettenregel). *Seien $U, V \subset \mathbb{R}$, $f : U \rightarrow V$ differenzierbar in $x \in U$ und $g : V \rightarrow \mathbb{R}$ differenzierbar in $f(x) \in V$. Dann ist auch $g \circ f : U \rightarrow \mathbb{R}$ differenzierbar in x , und es gilt*

$$(g \circ f)'(x) = g'(f(x))f'(x).$$

BEWEIS. Wir setzen für $y := f(x)$:

$$g^*(z) := \begin{cases} \frac{g(z) - g(y)}{z - y} & \text{falls } z \neq y, \\ g'(y) & \text{falls } z = y. \end{cases}$$

Beachte, daß wegen der Differenzierbarkeit von g in y gilt $\lim_{z \rightarrow y} g^*(z) = g'(y)$.

Nun gilt für $h \neq 0$:

$$\frac{g(f(x+h)) - g(f(x))}{h} = g^*(f(x+h)) \cdot \frac{f(x+h) - f(x)}{h}.$$

⁶oft auch einfach ‚Produktregel‘.

Im Limes $h \rightarrow 0$ konvergiert (da f in x differenzierbar, also auch stetig ist) $f(x+h)$ gegen $f(x)$ und daher, weil g^* in $y = f(x)$ stetig ist, auch $g^*(f(x+h))$ gegen $g^*(f(x)) = g'(f(x))$. Der zweite Faktor konvergiert gegen $f'(x)$ nach Annahme der Differenzierbarkeit von f in x . Insgesamt konvergiert der Differenzenquotient also wie behauptet gegen $g'(f(x))f'(x)$. $\square$

Ein sorgloser Umgang mit Differentialen verleitet zu folgendem einfachen ‚Beweis‘ der Kettenregel:

$$\frac{dg}{dx} = \frac{dg}{df} \frac{df}{dx}.$$

Auch wenn dies natürlich kein mathematisch korrekter Beweis ist, so enthält er doch die Kernidee des richtigen Beweises, nämlich die Erweiterung des Differenzenquotienten mit $\Delta f := f(x+h) - f(x)$.

BEISPIEL 4.7. Betrachte für $k \in \mathbb{N}$ die Funktion $\mathbb{R}^+ \rightarrow \mathbb{R}$, $x \mapsto x^{-k}$. Sie kann aufgefaßt werden als Verknüpfung der Funktionen $f: \mathbb{R}^+ \rightarrow \mathbb{R}^+$, $x \mapsto \frac{1}{x}$, und $g: \mathbb{R}^+ \rightarrow \mathbb{R}$, $y \mapsto y^k$. Wir haben bereits gesehen, daß $f'(x) = -\frac{1}{x^2}$ und $g'(y) = ky^{k-1}$. Mit der Kettenregel ergibt sich somit

$$\frac{d}{dx}(x^{-k}) = (g \circ f)'(x) = g'(f(x))f'(x) = -kx^{1-k} \frac{1}{x^2} = -kx^{-k-1}.$$

Die Regel $\frac{d}{dx}x^n = nx^{n-1}$ gilt also sogar für alle $n \in \mathbb{Z}$.

SATZ 4.8 (Quotientenregel). Seien $f, g: U \rightarrow \mathbb{R}$ beide in $x \in U$ differenzierbar und gelte $g(x) \neq 0$. Dann ist auch $\frac{f}{g}$ in x differenzierbar, und es gilt

$$\left(\frac{f}{g}\right)'(x) = \frac{f'(x)g(x) - f(x)g'(x)}{g(x)^2}.$$

BEWEIS. Eine implizite Voraussetzung für die Differenzierbarkeit einer Funktion in $x \in U$ war die Existenz einer gegen x konvergenten Folge in $U \setminus \{x\}$. Da der Definitionsbereich von $\frac{f}{g}$ womöglich kleiner ist als U , müssen wir diese Voraussetzung hier überprüfen. Nach Annahme der Differenzierbarkeit von g in x existiert eine Folge $(x_n)_{n \in \mathbb{N}} \subset U$ mit $x_n \neq x$ und $\lim_{n \rightarrow \infty} x_n = x$. Da $g(x) \neq 0$ und g in x stetig ist, ist g auch in einer Umgebung von x ungleich null (Übung!), sodaß fast alle x_n in $\{x \in U : g(x) \neq 0\}$, also dem Definitionsbereich von $\frac{f}{g}$, liegen.

Nach diesem Prolegomenon nun zum eigentlichen Beweis: Sei $G := \frac{1}{g}$, so gilt wegen $\left(\frac{1}{x}\right)' = -\frac{1}{x^2}$ und der Kettenregel, daß

$$G'(x) = -\frac{1}{g(x)^2}g'(x),$$

und nach LEIBNIZ-Regel

$$\begin{aligned} \left(\frac{f}{g}\right)'(x) &= (fG)'(x) = f'(x)G(x) + f(x)G'(x) \\ &= \frac{f'(x)}{g(x)} - \frac{f(x)g'(x)}{g(x)^2} = \frac{f'(x)g(x) - f(x)g'(x)}{g(x)^2}. \end{aligned}$$

$\square$

BEISPIEL 4.9.

$$\frac{d}{dx} \frac{x^2 - 1}{x^2 + 1} = \frac{2x(x^2 + 1) - 2x(x^2 - 1)}{(x^2 + 1)^2} = \frac{4x}{(x^2 + 1)^2}.$$

4.1.2.2. *Ableitung der Umkehrfunktion.* Sei $I \subset \mathbb{R}$ ein Intervall. Wir erinnern uns (Satz 1.23): Ist $f : I \rightarrow f(I)$ bijektiv, so besitzt f eine eindeutig bestimmte Umkehrfunktion $f^{-1} : f(I) \rightarrow I$. Eine streng monoton wachsende oder fallende Funktion auf einem Intervall ist insbesondere bijektiv (wenn der Wertebereich als $f(I)$ gewählt wird), da aus $x \neq y$ auch $f(x) \neq f(y)$ folgt.

SATZ 4.10 (Ableitung der Umkehrfunktion). Sei $I \subset \mathbb{R}$ ein Intervall und $f : I \rightarrow \mathbb{R}$ streng monoton. Ist f im Punkt $x \in I$ differenzierbar mit $f'(x) \neq 0$, so ist die Umkehrfunktion $f^{-1} : f(I) \rightarrow I$ im Punkt $y := f(x)$ differenzierbar, und es gilt

$$(f^{-1})'(y) = \frac{1}{f'(x)}.$$

BEWEIS. Für die Differenzierbarkeit müssen wir zunächst wieder zeigen, daß überhaupt eine gegen y konvergente Folge $(y_n)_{n \in \mathbb{N}} \subset f(I)$ existiert mit $y_n \neq y$ für alle $n \in \mathbb{N}$. Sei dazu $(x_n)_{n \in \mathbb{N}} \subset I$ eine gegen x konvergente Folge mit $x_n \neq x$. Da f in x differenzierbar und damit auch stetig ist, folgt $y_n := f(x_n) \rightarrow f(x) = y$ für $n \rightarrow \infty$, offenbar ist $y_n \in f(I)$ für jedes n , und aufgrund der Bijektivität folgt aus $x_n \neq x$ auch $y_n \neq y$.

Sei nun also $(y_n)_{n \in \mathbb{N}} \subset f(I) \setminus \{y\}$ eine gegen y konvergente Folge und setze $x_n := f^{-1}(y_n)$. Dann gilt $x_n \neq x$ und $\lim_{n \rightarrow \infty} x_n = x$, da f^{-1} stetig ist (Übung 8, Aufgabe 9). Daher gilt

$$\frac{f^{-1}(y_n) - f^{-1}(y)}{y_n - y} = \frac{x_n - x}{f(x_n) - f(x)} \rightarrow \frac{1}{f'(x)}$$

für $n \rightarrow \infty$, da $f'(x) \neq 0$. Damit ist die Behauptung gezeigt. $\square$

Eine Eselsbrücke für diese Ableitungsregel (und ihren Beweis) ist die ‚Gleichheit‘

$$\frac{dy}{dx} = \frac{1}{\frac{dx}{dy}}.$$

BEISPIEL 4.11 (Ableitung von Wurzeln). Betrachte für $k \in \mathbb{N}$ die Abbildung $\mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$, $x \mapsto x^k$. Nach Abschnitt 3.2.1.2 ist diese Abbildung invertierbar mit Umkehrfunktion $\mathbb{R}_0^+ \rightarrow \mathbb{R}_0^+$, $y \mapsto \sqrt[k]{y}$. Für alle $x > 0$ gilt

$$\frac{d}{dx}(x^k) = kx^{k-1} \neq 0,$$

also ist die k -te Wurfelfunktion in jedem $y > 0$ differenzierbar, und es gilt nach der Regel für die Ableitung der Umkehrfunktion (wir schreiben $y = x^k$):

$$\frac{d}{dy} \sqrt[k]{y} = \frac{1}{kx^{k-1}} = \frac{1}{k \sqrt[k]{y}^{k-1}}.$$

Im nächsten Beispiel werden wir sehen, daß diese Regel als Spezialfall der Potenzregel aufgefaßt werden kann:

BEISPIEL 4.12 (Rationale Potenzen). Betrachte für $p \in \mathbb{Z}$, $q \in \mathbb{N}$ die Funktion $\mathbb{R}^+ \rightarrow \mathbb{R}$, $x \mapsto x^{p/q} := \sqrt[q]{x^p}$. Man prüft leicht nach, daß auch für so definierte rationale Exponenten die Potenzgesetze aus Abschnitt 1.3.2.3 weiterhin gelten. Wir können diese Funktion auffassen als Verknüpfung von $f : \mathbb{R}^+ \rightarrow \mathbb{R}^+$, $f(x) = x^p$, und $g : \mathbb{R}^+ \rightarrow \mathbb{R}$, $g(y) = \sqrt[q]{y}$. Nach Kettenregel und dem vorigen Beispiel haben wir

$$\frac{d}{dx} x^{p/q} = (g \circ f)'(x) = g'(f(x))f'(x) = \frac{px^{p-1}}{q \sqrt[q]{x^{p(q-1)}}} = \frac{p}{q} x^{p-1-p+\frac{p}{q}} = \frac{p}{q} x^{\frac{p}{q}-1}.$$

Damit ist gezeigt, daß die Ableitungsregel $(x^\alpha)' = \alpha x^{\alpha-1}$ sogar für alle $\alpha \in \mathbb{Q}$ gilt.

4.1.3. Höhere Ableitungen. Ist $I \subset \mathbb{R}$ ein Intervall und $f : I \rightarrow \mathbb{R}$ überall differenzierbar, so ist die Ableitung $f' : I \rightarrow \mathbb{R}$ ebenfalls eine Funktion. Ist diese stetig in I , so heißt f dort *stetig differenzierbar*. Nicht jede differenzierbare Funktion ist stetig differenzierbar:

BEISPIEL 4.13. Wir verwenden in diesem Beispiel die Ihnen aus der Schule bekannte Sinusfunktion, deren Ableitung der Kosinus ist. Später werden wir noch systematisch über diese speziellen Funktionen sprechen.

Betrachte also die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$, definiert durch

$$f(x) := \begin{cases} x^2 \sin\left(\frac{1}{x}\right) & \text{falls } x \neq 0, \\ 0 & \text{falls } x = 0. \end{cases}$$

In $x \neq 0$ ist f als Produkt bzw. Komposition differenzierbarer Funktionen selbst differenzierbar mit

$$f'(x) = 2x \sin\left(\frac{1}{x}\right) - x^2 \cos\left(\frac{1}{x}\right) \cdot \left(-\frac{1}{x^2}\right) = 2x \sin\left(\frac{1}{x}\right) + \cos\left(\frac{1}{x}\right).$$

In $x = 0$ ist f aber ebenfalls differenzierbar mit

$$f'(0) = \lim_{x \rightarrow 0} \frac{x^2 \sin\left(\frac{1}{x}\right)}{x} = 0,$$

da $\sin$ beschränkt ist.

Allerdings ist f' in 0 nicht stetig: Betrachte etwa die durch $x_n = \frac{1}{2\pi n}$ gegebene Nullfolge, für die gilt

$$f'(x_n) = 2x_n \sin\left(\frac{1}{x_n}\right) + \cos\left(\frac{1}{x_n}\right) \rightarrow 1 \neq 0 = f'(0).$$

für $n \rightarrow \infty$, da $\cos(2\pi n) = 1$.

Also ist f auf ganz $\mathbb{R}$ differenzierbar, aber nicht stetig differenzierbar.

Ist für eine differenzierbare Funktion $f : I \rightarrow \mathbb{R}$ auch die Ableitungsfunktion f' wiederum auf I differenzierbar, so kann man die Ableitung von f' bilden. Diese bezeichnet man als *zweite Ableitung* von f und bezeichnet sie mit f'' oder $\frac{d^2 f}{dx^2}$.

Durch k -malige Ableitung erhält man, sofern existent, die *k-te Ableitung*, die man mit f^k oder $\frac{d^k f}{dx^k}$ bezeichnet. Eine Funktion, die k -mal differenzierbar ist und deren k -te Ableitung auf ganz I stetig ist, heißt *k-mal stetig differenzierbar*. Nach Konvention ist die nullte Ableitung einer Funktion die Funktion selbst.

Wir verwenden folgende Notation: Für ein Intervall I bezeichne $C^k(I)$ die Menge der k -mal stetig differenzierbaren Funktionen $I \rightarrow \mathbb{R}$. Insbesondere ist $C^0(I)$ die Menge der stetigen Funktionen auf I (statt $C^0(I)$ schreibt man manchmal auch einfach $C(I)$). Die Menge der beliebig oft differenzierbaren Funktionen heißt schließlich $C^\infty(I)$.

SATZ 4.14. Seien $I \subset \mathbb{R}$ ein Intervall und $k \in \mathbb{N}$. Mit der punktweisen Addition $(f + g)(x) := f(x) + g(x)$ und der skalaren Multiplikation $(cf)(x) := cf(x)$ ist $C^k(I)$ ein $\mathbb{R}$ -Vektorraum.

Für $f, g \in C^k(I)$ ist auch das (punktweise definierte) Produkt fg in $C^k(I)$ ⁷, und es gilt die allgemeine LEIBNIZ-Regel

$$(fg)^{(k)} = \sum_{j=0}^k \binom{k}{j} f^{(j)} g^{(k-j)}.$$

BEWEIS. Für die Vektorraumstruktur ist nur zu zeigen, daß Summen und skalare Vielfache k -mal stetig differenzierbarer Funktionen wieder k -mal stetig differenzierbar sind; dies folgt aber durch k -malige Anwendung der Linearität der Ableitung (Proposition 4.4), und aus der Stetigkeit der Summe und des Produkts stetiger Funktionen (Korollar 3.6).

⁷Man sagt auch, $C^k(I)$ bilde eine *Algebra*.

Der zweite Teil wird den Studierenden zur Übung überlassen. $\square$

Die Räume $C^k(I)$ sind abstrakte Vektorräume (d.h. sie erfüllen die Vektorraumaxiome, lassen sich aber nicht so einfach veranschaulichen wie etwa $\mathbb{R}^3$). Vektorräume, deren Elemente Funktionen sind, heißen *Funktionsräume*; die Auffassung von Mengen von Funktionen als Vektorräume ist in der höheren Analysis sehr fruchtbar, z.B. in der Theorie der Fourierreihen und -transformationen (Analysis II oder III oder Funktionalanalysis).

Funktionsräume sind typischerweise unendlichdimensional: Betrachte etwa ein (aus mehr als einem Punkt bestehendes) I und den Raum $C(I)$ der stetigen Funktionen. Wir geben eine unendliche Menge von linear unabhängigen Elementen von $C(I)$ an, etwa die Funktionen $f_k : x^k$ für $k \in \mathbb{N} \cup \{0\}$. Linearkombinationen dieser Funktionen sind Polynomfunktionen, und eine Polynomfunktion ist genau dann identisch null, wenn alle Koeffizienten null sind: Betrachte nämlich das Polynom

$$p(x) = \sum_{k=0}^N a_k x^k,$$

und sei $p(x) = 0$ für alle $x \in I$. Wir nehmen an $0 \in I$ (es ist nicht schwer zu zeigen, daß diese Annahme keine Einschränkung darstellt). Für die k -te Ableitung von p gilt $p^{(k)}(0) = k!a_k$ für $k = 0, \dots, N$ (Übung). Da aber alle Ableitungen der Nullfunktion selbst wieder null sind, gilt wie behauptet $a_k = 0$ für alle $k = 0, \dots, N$.⁸

4.2. Monotonie und Konvexität

Mithilfe des Ableitungsbegriffes lassen sich die wichtigsten Eigenschaften einer gegebenen Funktion ermitteln: Monotonie, lokale Extrema, Krümmung, Wendepunkte usw. Dies haben Sie in der gymnasialen Oberstufe extensiv eingeübt. Der Zusammenhang zwischen Ableitung und Monotonie wird mit dem wichtigen Mittelwertsatz der Differentialrechnung hergestellt, den wir nun vorstellen.

4.2.1. Mittelwertsatz und Monotonie.

4.2.1.1. *Mittelwertsatz der Differentialrechnung.* Sei I ein offenes Intervall und $f : I \rightarrow \mathbb{R}$. Man sagt, f habe im Punkt $x \in I$ ein *lokales Maximum (Minimum)*, wenn es eine Umgebung $B_\epsilon(x) \subset I$ (siehe Abschnitt 3.1.3) gibt, sodaß für alle $y \in B_\epsilon(x)$ gilt

$$f(y) \leq f(x) \quad (\text{bzw. } f(y) \geq f(x)).$$

Lokale Maxima und Minima bezeichnet man zusammenfassend als lokale *Extrema*. Hat die Funktion im offenen Intervall I ein *globales Maximum (Minimum)*⁹, so ist dieses insbesondere ein *lokales Maximum bzw. Minimum*. Der Punkt, an dem eine Funktion ein Extremum hat, heißt (lokale/globale) *Extremalstelle*, und der Funktionswert an dieser Stelle heißt (lokaler/globaler) *Extremalwert*.

SATZ 4.15 (Notwendige Bedingung für lokale Extrema). Sei I ein offenes Intervall und $f : I \rightarrow \mathbb{R}$. Hat f in $x \in I$ ein lokales Extremum und ist f in x differenzierbar, so gilt $f'(x) = 0$.

BEWEIS. Wir behandeln nur den Fall eines lokalen Maximums, der andere Fall folgt analog oder durch Übergang zu $-f$.

Sei $\epsilon > 0$ so klein, daß $f(x) \geq f(y)$ für alle $y \in B_\epsilon(x)$. Sei $(x_n)_{n \in \mathbb{N}} \subset (x, x + \epsilon)$ eine (von rechts) gegen x konvergente Folge, dann gilt nach Annahme der Differenzierbarkeit

$$f'(x) = \lim_{n \rightarrow \infty} \frac{f(x) - f(x_n)}{x - x_n} \leq 0,$$

⁸Ein anderer Beweis der linearen Unabhängigkeit der Potenzfunktionen verwendet die Regularität der VANDERMONDE-Matrix, die Ihnen vielleicht in der Linearen Algebra oder der Numerik begegnen wird.

⁹das bedeutet: $f(x) \geq f(y)$ bzw. $f(x) \leq f(y)$ für alle $y \in I$, nicht nur solche in einer Umgebung von x .

da $f(x) - f(x_n) \geq 0$, aber $x - x_n < 0$.

Ist andererseits $(x'_n)_{n \in \mathbb{N}} \subset (x - \epsilon, x)$ eine weitere, nun von links gegen x konvergente Folge, so gilt in ähnlicher Weise

$$f'(x) = \lim_{n \rightarrow \infty} \frac{f(x) - f(x'_n)}{x - x'_n} \geq 0,$$

da wieder $f(x) - f(x'_n) \geq 0$, aber nun $x - x'_n > 0$. Insgesamt folgt wie behauptet $f'(x) = 0$. $\square$

BEMERKUNG 4.16. Die Bedingung $f'(x) = 0$ ist notwendig, aber nicht hinreichend für das Vorliegen eines lokalen Extremums: Die Funktion $f: x \mapsto x^3$ etwa hat $f'(0) = 0$, obwohl bei 0 kein Maximum oder Minimum vorliegt (man bezeichnet in diesem Beispiel 0 als einen *Sattelpunkt* der Funktion).

Man kann diese Bedingung aber verwenden, um alle ‚Kandidaten‘ für eine globale Extremalstelle zu finden: Ist die Funktion, deren Extremum gesucht wird, in einem offenen Intervall differenzierbar, so findet man mit der Gleichung $f'(x) = 0$ alle möglichen lokalen Extremalstellen, und kann dann unter allen diesen Möglichkeiten die globale Minimal- bzw. Maximalstelle durch Vergleich der zugehörigen Funktionswerte ermitteln.

Äußerste Vorsicht ist allerdings auf nicht-offenen Intervallen geboten: Auf $[0, 1]$ hat etwa die Identität $x \mapsto x$ die (globale) Maximalstelle $x = 1$ und die Minimalstelle $x = 0$, aber die Ableitung ist niemals null (sondern stets 1). Bei der Suche nach globalen Extrema muß man also ggf. auch die Randpunkte des Definitionsbereichs miteinbeziehen.

Auch die Annahme der Differenzierbarkeit ist entscheidend: Die Funktion $x \mapsto |x|$ hat in $\mathbb{R}$ ein (globales, also auch lokales) Minimum in $x = 0$, aber dort ist sie nicht differenzierbar. Auch hier würde also die Suche nach Nullstellen der Ableitung nicht das gewünschte Ergebnis liefern.

SATZ 4.17 (Mittelwertsatz der Differentialrechnung). Sei $I = [a, b] \subset \mathbb{R}$ ein kompaktes Intervall mit $a < b$. Sei $f: I \rightarrow \mathbb{R}$ stetig und in (a, b) differenzierbar. Dann existiert $\xi \in (a, b)$, sodaß

$$f'(\xi) = \frac{f(b) - f(a)}{b - a}.$$

BEWEIS. Wir nehmen zunächst an $f(a) = f(b) = 0$, d.h. wir müssen zeigen, daß ein $\xi \in (a, b)$ existiert mit $f'(\xi) = 0$. Da f stetig auf dem kompakten Intervall I ist, nimmt es nach Korollar 3.17 sein Maximum und sein Minimum an. Werden Maximum und Minimum beide am Intervallrand angenommen, so ist f identisch null, und $f'(\xi) = 0$ für beliebiges $\xi \in (a, b)$. Wird dagegen das Maximum oder das Minimum in einem Punkt $\xi \in (a, b)$ angenommen, so liegt dort insbesondere ein lokales Extremum vor, und nach Satz 4.15 gilt dort $f'(\xi) = 0$. Damit ist der Spezialfall $f(a) = f(b) = 0$ erledigt¹⁰.

Für den allgemeinen Fall betrachte die Funktion $g: I \rightarrow \mathbb{R}$ gegeben durch

$$g(x) := f(x) - f(a) - \frac{f(b) - f(a)}{b - a}(x - a).$$

Offenbar erfüllt auch g alle Voraussetzungen des zu beweisenden Satzes, und zusätzlich $g(a) = g(b) = 0$. Nach dem ersten Beweisteil existiert daher ein $\xi \in (a, b)$, für das $g'(\xi) = 0$. Nach Definition von g ist dies aber äquivalent zu

$$0 = f'(\xi) - \frac{f(b) - f(a)}{b - a},$$

und dies ist genau die Behauptung. $\square$

¹⁰Aus historischen Gründen hat dieser Spezialfall des Mittelwertsatzes einen eigenen Namen: Er heißt *Satz von ROLLE*.

Wenn wir uns an die mechanische Interpretation der Ableitung erinnern, erschließt sich die Bezeichnung ‚Mittelwertsatz‘: Ist nämlich $s = s(t)$ der zurückgelegte Weg und $v = \dot{s}$, so ist $v(t)$ die Momentangeschwindigkeit zur Zeit t und

$$\frac{s(t_2) - s(t_1)}{t_2 - t_1}$$

die Durchschnittsgeschwindigkeit im Zeitintervall $[t_1, t_2]$; der Mittelwertsatz besagt dann, dann die mittlere Geschwindigkeit zu irgendeinem Zeitpunkt gleich der Momentangeschwindigkeit ist.

In geometrischer Interpretation besagt der Satz, daß es eine Tangente an den Graphen von f gibt, die parallel zur Sekante durch die Randpunkte verläuft.

KOROLLAR 4.18. Sei $I = [a, b] \subset \mathbb{R}$ ein kompaktes Intervall mit $a < b$. Sei $f : I \rightarrow \mathbb{R}$ stetig und in (a, b) differenzierbar mit $f'(x) = 0$ für alle $x \in (a, b)$. Dann ist f konstant.

BEWEIS. Seien $x_1, x_2 \in [a, b]$ mit $x_1 < x_2$ beliebig. Dann erfüllt f auch auf dem Intervall $[x_1, x_2]$ die Voraussetzungen des Mittelwertsatzes, und gemäß diesem gibt es $\xi \in (x_1, x_2)$ mit

$$f'(\xi) = \frac{f(x_2) - f(x_1)}{x_2 - x_1}.$$

Da aber nach Voraussetzung $f'(\xi) = 0$, folgt $f(x_1) = f(x_2)$, und da x_1, x_2 beliebig waren, ist f konstant. $\square$

4.2.1.2. Monotonie.

SATZ 4.19. Sei $I \subset \mathbb{R}$ ein Intervall und f auf I differenzierbar. Dann ist f auf I monoton wachsend genau dann, wenn $f'(x) \geq 0$ für alle $x \in I$.

Ebenso ist f monoton fallend genau dann, wenn $f'(x) \leq 0$ für alle $x \in I$.

Ist schließlich $f'(x) > 0$ für alle $x \in I$, so ist f sogar streng monoton wachsend in I . Analog ist f auf I streng monoton fallend, falls $f'(x) < 0$ für alle $x \in I$.

BEMERKUNG 4.20. Beachte, daß die Umkehrung der letzten Aussage nicht gilt: Die Funktion $x \mapsto x^3$ ist in ganz $\mathbb{R}$ differenzierbar und streng monoton steigend, aber $f'(0) = 0$.

BEWEIS. Sei zunächst f in I monoton wachsend und $x \in I$, dann ist $f(x+h) \geq f(x)$ für jedes $h > 0$, und daher

$$f'(x) = \lim_{h \searrow 0} \frac{f(x+h) - f(x)}{h} \geq 0.$$

(Falls x der rechte Intervallrand ist, argumentiere analog mit $h < 0$.)

Sei umgekehrt $f'(x) \geq 0$ für alle $x \in I$. Angenommen, f wäre nicht monoton wachsend, dann gäbe es $x, x' \in I$ mit $x < x'$, aber $f(x) > f(x')$. Nach Mittelwertsatz gäbe es dann ein $\xi \in (x, x')$, sodaß

$$f'(\xi) = \frac{f(x') - f(x)}{x' - x} < 0,$$

im Widerspruch zur Voraussetzung. Ist sogar $f'(x) > 0$ für alle $x \in I$, so erhielte man aus der Annahme, f sei nicht streng monoton wachsend, $x, x' \in I$ mit $x < x'$, aber $f(x) \geq f(x')$, und der Mittelwertsatz lieferte ein $\xi \in (x, x')$ mit $f'(\xi) \leq 0$, Widerspruch.

Die übrigen Aussagen folgen durch Übergang zu $-f$. $\square$

4.2.2. Die zweite Ableitung.

4.2.2.1. *Hinreichende Bedingung für lokale Extrema.* In Satz 4.15 haben wir eine *notwendige* Bedingung für das Vorliegen eines lokalen Extremums einer differenzierbaren Funktion etabliert: Wenn f in x ein lokales Extremum besitzt, so ist dort $f'(x) = 0$. Die Frage nach der Umkehrung haben wir aber offengelassen: Angenommen, wir haben einen Nullstelle der Ableitung gefunden, wie können wir wissen, ob dort auch tatsächlich ein lokales Extremum vorliegt? Und handelt es sich ggf. um ein lokales Maximum oder Minimum? Die Untersuchung der zweiten Ableitung kann hier Abhilfe schaffen.

SATZ 4.21 (hinreichende Bedingung für lokale Extrema). Sei $I \subset \mathbb{R}$ ein offenes Intervall und $f \in C^2(I)$ ¹¹. Ist $f'(x) = 0$ und $f''(x) > 0$, so hat f in x ein lokales Minimum. Ist $f'(x) = 0$ und $f''(x) < 0$, so hat f in x ein lokales Maximum.

BEMERKUNG 4.22. Im Falle $f'(x) = f''(x) = 0$ ist keine allgemeine Aussage möglich: Die Funktionen $x \mapsto x^3$, $x \mapsto |x|^3$, $x \mapsto -|x|^3$ haben jeweils verschwindende erste und zweite Ableitungen in $x = 0$, haben dort aber kein lokales Extremum bzw. ein lokales Minimum bzw. ein lokales Maximum.

BEWEIS. Wir zeigen nur die erste Aussage, da die zweite durch Übergang zu $-f$ folgt. Da f'' nach Voraussetzung in x stetig ist und $f''(x) > 0$, so ist sogar $f'' > 0$ in einer Umgebung von x . Sei $h \neq 0$ so gewählt, daß $x + h$ Element dieser Umgebung ist.

Anwendung des Mittelwertsatzes auf f im Intervall $[x, x + h]$ ¹² ergibt ein $\xi_1 \in (x, x + h)$ mit

$$f(x + h) = f(x) + hf'(\xi_1). \quad (4.4)$$

Erneute Anwendung des Mittelwertsatzes auf f' im Intervall $[x, \xi_1]$ ergibt nun ein $\xi_2 \in (x, \xi_1)$ mit

$$f'(\xi_1) = f'(x) + (\xi_1 - x)f''(\xi_2) = (\xi_1 - x)f''(\xi_2) \quad (4.5)$$

da ja $f'(x) = 0$. Wir bemerken $h(\xi_1 - x) > 0$, da h und $(\xi_1 - x)$ das gleiche Vorzeichen haben, und nach Wahl von h ist außerdem $f''(\xi_2) > 0$. Daher erhalten wir durch Einsetzen von (4.5) in (4.4):

$$f(x + h) = f(x) + hf'(\xi_1) = f(x) + h(\xi_1 - x)f''(\xi_2) > f(x),$$

und somit ist x eine lokale Minimalstelle von f . □

Eine Art Umkehrung dieses Satzes lautet:

KOROLLAR 4.23. Sei $I \subset \mathbb{R}$ ein offenes Intervall und $f \in C^2(I)$. Hat f in x ein lokales Minimum, so gilt dort $f'(x) = 0$ und $f''(x) \geq 0$. Hat f dagegen in x ein lokales Maximum, so ist $f'(x) = 0$ und $f''(x) \leq 0$.

BEWEIS. Wir zeigen wieder nur den ersten Teil. Die Aussage über die erste Ableitung ist genau Satz 4.15. Wäre nun $f''(x) < 0$, so hätte f nach Satz 4.21 ein lokales Maximum. Da bei x also sowohl eine lokale Maximal- wie auch Minimalstelle vorliegt, ist f in einer Umgebung von x konstant; dann aber sind alle Ableitungen von f in x gleich null, insbesondere $f''(x) = 0$, im Widerspruch zur Annahme $f''(x) < 0$. □

4.2.2.2. Konvexität.

DEFINITION 4.24 (Konvexität). Sei $I \subset \mathbb{R}$ ein Intervall. Eine Funktion $f : I \rightarrow \mathbb{R}$ heißt *konvex*, falls für alle $x, y \in I$ und $\lambda \in [0, 1]$ gilt:

$$f(\lambda x + (1 - \lambda)y) \leq \lambda f(x) + (1 - \lambda)f(y).$$

Eine Funktion f heißt *konkav*, wenn $-f$ konvex ist.

¹¹Wir erinnern uns: $C^2(I)$ ist der Raum der zweimal stetig differenzierbaren Funktionen in I .

¹²Für den Fall $h < 0$ beachte unsere Konvention aus Abschnitt 3.2.2.1.

In geometrischer Interpretation durchläuft der Graph einer konvexen Funktion eine ‚Linkskurve‘, der Graph einer konkaven Funktion dagegen eine ‚Rechtskurve‘. Ist eine Funktion auf einem Intervall gleichzeitig konvex und konkav, so ist sie dort affin (d.h. ihr Graph ist eine gerade Strecke).

Konvexität läßt sich mit der zweiten Ableitung charakterisieren. Wir benötigen ein vorbereitendes Resultat:

LEMMA 4.25. Seien $a, b, c, d \in \mathbb{R}$ mit $a < b < d$ und $a < c < d$, und sei $f : [a, d] \rightarrow \mathbb{R}$ konvex. Dann gilt

$$\frac{f(b) - f(a)}{b - a} \leq \frac{f(d) - f(a)}{d - a} \leq \frac{f(d) - f(c)}{d - c}.$$

BEWEIS. Wähle zuerst $\lambda = \frac{d-b}{d-a} \in (0, 1)$ und wende die Definition der Konvexität mit diesem λ und mit $x = a$, $y = d$ an. Unter Beachtung von $\lambda a + (1 - \lambda)d = b$ ergibt sich

$$f(b) \leq \lambda f(a) + (1 - \lambda)f(d),$$

und Einsetzen des Werts für λ liefert nach kurzer Umformung

$$\frac{f(b) - f(a)}{b - a} \leq \frac{f(d) - f(a)}{d - a}. \quad (4.6)$$

In ähnlicher Weise erhält man mit $\lambda = \frac{d-c}{d-a}$ aus der Konvexität

$$f(c) \leq \lambda f(a) + (1 - \lambda)f(d),$$

und nach Umformung

$$\frac{f(d) - f(c)}{d - c} \geq \frac{f(d) - f(a)}{d - a}. \quad (4.7)$$

Aus (4.6) und (4.7) folgt die Behauptung. $\square$

SATZ 4.26. Sei $I \subset \mathbb{R}$ ein Intervall und $f \in C^2(I)$. Dann ist f auf I konvex genau dann, wenn $f''(x) \geq 0$ für alle $x \in I$.

BEWEIS. Sei f konvex und seien x, y Punkte im Inneren von I mit $x < y$. Sei $h > 0$ so gewählt, daß $x + h, y + h \in I$. Nach Lemma 4.25 mit $a = x$, $b = x + h$, $c = y$, $d = y + h$ gilt

$$\frac{f(x + h) - f(x)}{h} \leq \frac{f(y + h) - f(y)}{h}$$

und daher, nach Übergang zum Limes $h \rightarrow 0$, auch $f'(x) \leq f'(y)$. Im Inneren von I ist also f' monoton wachsend, und nach Satz 4.19 gilt daher $f''(x) \geq 0$ für alle x im Inneren von I . Da aber nach Voraussetzung f'' ggf. an den Intervallrändern stetig ist, gilt auch dort $f'' \geq 0$.

Sei nun umgekehrt $f''(x) \geq 0$ für alle $x \in I$. Dann ist wiederum nach Satz 4.19 f' monoton wachsend. Seien $x, y \in I$ und $\lambda \in [0, 1]$, und schreibe $\bar{x} := \lambda x + (1 - \lambda)y$. Nach Mittelwertsatz existieren $\xi_1 \in (x, \bar{x})$ und $\xi_2 \in (\bar{x}, y)$ sodaß

$$\frac{f(\bar{x}) - f(x)}{\bar{x} - x} = f'(\xi_1) \leq f'(\xi_2) = \frac{f(y) - f(\bar{x})}{y - \bar{x}},$$

wobei wir die Monotonie von f' verwendet haben. Daraus folgt aber nach elementaren Umformungen (unter Beachtung der Wahl von $\bar{x}$), daß

$$f(\bar{x}) \leq \lambda f(x) + (1 - \lambda)f(y),$$

also die gewünschte Konvexität. $\square$

4.3. Das Integral stetiger Funktionen

4.3.1. Definition und Eigenschaften.

4.3.1.1. *Definition.* Sei bis auf Widerruf $I = [a, b] \subset \mathbb{R}$ ein kompaktes Intervall mit $a < b$. Eine Zerlegung

$$a = x_0 < x_1 < \dots < x_{N-1} < x_N = b$$

von I in N Teilintervalle hat die *Feinheit* $\max_{n=1, \dots, N} (x_n - x_{n-1})$.

SATZ 4.27 (Konvergenz der RIEMANN-Summen). Sei $f \in C(I)$. Dann existiert eine reelle Zahl $\mathcal{I}(f)$, sodaß für jedes $\epsilon > 0$ ein $\delta > 0$ existiert mit der folgenden Eigenschaft:

Ist $a = x_0 < x_1 < \dots < x_{N-1} < x_N = b$ eine Zerlegung mit Feinheit kleiner δ , so gilt

$$\left| \mathcal{I}(f) - \sum_{n=1}^N f(x_n)(x_n - x_{n-1}) \right| < \epsilon.$$

Man nennt Ausdrücke der Gestalt $\sum_{n=1}^N f(x_n)(x_n - x_{n-1})$ RIEMANN-Summen zur Funktion f .

BEWEIS. Seien $\epsilon > 0$ und $a = x_0 < x_1 < \dots < x_{N-1} < x_N = b$ sowie $a = y_0 < y_1 < \dots < y_{M-1} < y_M = b$ zwei unterschiedliche Zerlegungen von I . Wir zeigen zunächst: Es existiert ein $\delta > 0$, sodaß

$$\left| \sum_{m=1}^M f(y_m)(y_m - y_{m-1}) - \sum_{n=1}^N f(x_n)(x_n - x_{n-1}) \right| < \frac{\epsilon}{2}, \quad (4.8)$$

sofern die Feinheiten beider Zerlegungen kleiner δ sind.

Nach Satz 3.20 ist f als stetige Funktion auf einem kompakten Intervall sogar gleichmäßig stetig, d.h. zu unserem gegebenen $\epsilon > 0$ gibt es ein $\delta > 0$, sodaß für alle $x, y \in [a, b]$ mit $|x - y| < \delta$ gilt $|f(x) - f(y)| < \frac{\epsilon}{4(b-a)}$. Wir zeigen, daß mit diesem δ Eigenschaft (4.8) erfüllt ist.

Sei dazu $a = z_0 < z_1 < \dots < z_{R-1} < z_R = b$ die *gemeinsame Verfeinerung* der beiden Zerlegungen $(x_n)_{n=0}^N$ und $(y_m)_{m=0}^M$, das heißt,

$$\bigcup_{r=0}^R \{z_r\} = \bigcup_{n=0}^N \{x_n\} \cup \bigcup_{m=0}^M \{y_m\},$$

und die z_r sind aufsteigend angeordnet.

Bezeichne für $n = 1, \dots, N$ mit $J_n \subset \{1, \dots, R\}$ die Menge derjenigen Indizes r , für die $[z_r - z_{r-1}] \subset [x_n, x_{n+1}]$. Dann ist $\bigcup_{n=1}^N J_n = \{1, \dots, R\}$, $\sum_{r \in J_n} (z_r - z_{r-1}) = x_n - x_{n-1}$ und, wegen der gleichmäßigen Stetigkeit,

$$|f(z_r) - f(x_n)| < \frac{\epsilon}{4(b-a)} \quad \text{für alle } n = 1, \dots, N \text{ und alle } r \in J_n,$$

sofern die Feinheit der x -Zerlegung kleiner δ ist. Es folgt

$$\begin{aligned} & \left| \sum_{r=1}^R f(z_r)(z_r - z_{r-1}) - \sum_{n=1}^N f(x_n)(x_n - x_{n-1}) \right| \\ &= \left| \sum_{r=1}^R f(z_r)(z_r - z_{r-1}) - \sum_{n=1}^N \sum_{r \in J_n} f(x_n)(z_r - z_{r-1}) \right| \\ &= \left| \sum_{n=1}^N \sum_{r \in J_n} (f(z_r) - f(x_n))(z_r - z_{r-1}) \right| \\ &\leq \sum_{n=1}^N \sum_{r \in J_n} |f(z_r) - f(x_n)|(z_r - z_{r-1}) \\ &< \frac{\epsilon}{4(b-a)} \sum_{r=1}^R (z_r - z_{r-1}) = \frac{\epsilon}{4}. \end{aligned}$$

Analog haben wir

$$\left| \sum_{r=1}^R f(z_r)(z_r - z_{r-1}) - \sum_{m=1}^M f(y_m)(y_m - y_{m-1}) \right| < \frac{\epsilon}{4},$$

und zusammen folgt (4.8).

Der Abschluß des Beweises ist nun nicht mehr schwer: Gemäß (4.8) bildet jede Folge von RIEMANN-Summen zu f , deren Feinheiten gegen null konvergieren, eine Cauchyfolge und konvergiert somit. Wir wählen eine solche Folge aus, bezeichnen ihren Grenzwert mit $\mathcal{I}(f)$, und betrachten eine RIEMANN-Summe $\sum_{m=1}^M f(y_m)(y_m - y_{m-1})$ aus dieser Folge mit Feinheit kleiner δ und so, daß

$$\left| \mathcal{I}(f) - \sum_{m=1}^M f(y_m)(y_m - y_{m-1}) \right| < \frac{\epsilon}{2}.$$

Ist dann $\sum_{n=1}^N f(x_n)(x_n - x_{n-1})$ eine beliebige RIEMANN-Summe der Feinheit kleiner δ , so folgt aus (4.8)

$$\begin{aligned} & \left| \mathcal{I}(f) - \sum_{n=1}^N f(x_n)(x_n - x_{n-1}) \right| \\ & \leq \left| \mathcal{I}(f) - \sum_{m=1}^M f(y_m)(y_m - y_{m-1}) \right| + \left| \sum_{m=1}^M f(y_m)(y_m - y_{m-1}) - \sum_{n=1}^N f(x_n)(x_n - x_{n-1}) \right| \\ & < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon. \end{aligned}$$

□

Dieser Grenzwert ist dann das Integral von f :

DEFINITION 4.28 (Integral einer stetigen Funktion). Die Zahl

$$\int_a^b f(x)dx := \mathcal{I}(f)$$

aus Satz 4.27 heißt (bestimmtes) *Integral* von f im Intervall $[a, b]$, und f heißt *Integrand* in diesem Integral.

Geometrisch interpretiert man $\int_a^b f(x)dx$ als Fläche unter dem Graphen von f , also die Fläche, die der Graph von f mit den Vertikalen $x = a$ und $x = b$ und der x -Achse einschließt, wobei Flächenstücke, die unterhalb der x -Achse verlaufen, negativ gewichtet werden. Die Idee bei der Berechnung bzw. Definition dieser Fläche ist die Approximation durch N sehr dünne Rechtecke mit Grundlinie $x_n - x_{n-1}$ und Höhe $f(x_n)$ (also dem Funktionswert am rechten Randpunkt der Grundlinie des Rechtecks).

BEMERKUNG 4.29. Satz 4.27 bleibt auch für eine viel größere Klasse von Funktionen gültig, die Stetigkeit des Integranden ist also keinesfalls notwendig für die Konvergenz der RIEMANN-Summen. Funktionen, deren RIEMANN-Summen konvergieren, heißen RIEMANN-integrierbar. Wir verzichten hier auf die volle Allgemeinheit der RIEMANN-Integration, weil wir zunächst ohnehin nur an stetigen Integranden interessiert sind und sich herausstellen wird, daß das RIEMANN-Integral für weiterführende Zwecke in der Funktionalanalysis, der Theorie der partiellen Differentialgleichungen und der stochastischen Analysis und Finanzmathematik ungeeignet ist. Wir führen deshalb im dritten Semester in der Maßtheorie das leistungsfähigere LEBESGUE-Integral ein, das glücklicherweise für stetige Funktionen auf kompakten Intervallen mit der hier gegebenen Definition des (RIEMANN)-Integrals übereinstimmt.

4.3.1.2. *Einfache Beispiele.*

- (1) Die konstante Funktion $x \mapsto c$ für ein $c \in \mathbb{R}$ hat Integral

$$\int_a^b c dx = c(b-a),$$

da jede RIEMANN-Summe die Form $\sum_{n=1}^N c(x_n - x_{n-1}) = c \sum_{n=1}^N (x_n - x_{n-1}) = c(b-a)$ hat.

- (2) Für die Identität $x \mapsto x$ erhalten wir als RIEMANN-Summe mit der speziellen (sog. äquidistanten) Intervallzerlegung $x_n = a + (b-a) \frac{n}{N}$:

$$\sum_{n=1}^N \left(a + (b-a) \frac{n}{N} \right) \frac{b-a}{N} = a(b-a) + \frac{(b-a)^2}{N^2} \frac{N(N+1)}{2} \rightarrow \frac{b^2}{2} - \frac{a^2}{2}$$

für $N \rightarrow \infty$, wobei wir die GAUSSsche Summenformel (Beispiel 1.28) verwendet haben.

- (3) Für $x \mapsto x^2$ setzen wir der Einfachheit halber $I = [0, 1]$ (allgemeine Intervalle können ebenso behandelt werden, nur mit etwas mehr Schreibaufwand). Die RIEMANN-Summe zu einer äquidistanten Zerlegung lautet dann

$$\sum_{n=1}^N \left(\frac{n}{N} \right)^2 \frac{1}{N} = \frac{1}{N^3} \sum_{n=1}^N n^2 = \frac{N(N+1)(2N+1)}{6N^3} \rightarrow \frac{1}{3}$$

für $N \rightarrow \infty$, wobei wir die Summenformel aus Übungsblatt 3, Aufgabe 1 benutzt haben.

4.3.1.3. *Elementare Eigenschaften.*

PROPOSITION 4.30 (Linearität und Monotonie). *Seien $f, g \in C(I)$ und $\lambda \in \mathbb{R}$. Dann gilt*

- (1) $\int_a^b (f+g)(x) dx = \int_a^b f(x) dx + \int_a^b g(x) dx$;
- (2) $\int_a^b (\lambda f)(x) dx = \lambda \int_a^b f(x) dx$.

Ist außerdem $f \leq g$ (das bedeutet $f(x) \leq g(x)$ für alle $x \in I$), so ist

- (3) $\int_a^b f(x) dx \leq \int_a^b g(x) dx$.

BEWEIS. Dies folgt sofort aus den entsprechenden Eigenschaften für Summen und Grenzwerte. Zum Beispiel gilt für (3) nach Voraussetzung für die RIEMANN-Summen

$$\sum_{n=1}^N f(x_n)(x_n - x_{n-1}) \leq \sum_{n=1}^N g(x_n)(x_n - x_{n-1}),$$

und im Limes kleiner Feinheiten folgt die Behauptung über die Integrale. □

SATZ 4.31 (Dreiecksungleichung). Sei $f \in C(\mathbb{R})$ so gilt

$$\left| \int_a^b f(x) dx \right| \leq \int_a^b |f(x)| dx.$$

BEWEIS. Auch dies folgt sofort aus der entsprechenden Dreiecksungleichung für die RIEMANN-Summen:

$$\left| \sum_{n=1}^N f(x_n)(x_n - x_{n-1}) \right| \leq \sum_{n=1}^N |f(x_n)|(x_n - x_{n-1})$$

□

Aufgrund der Interpretation des Integrals als Fläche ist auch die folgende Aussage nicht überraschend:

PROPOSITION 4.32. Seien $a < c < b$ reelle Zahlen. Dann gilt für jedes $f \in C([a, b])$:

$$\int_a^b f(x)dx = \int_a^c f(x)dx + \int_c^b f(x)dx.$$

BEWEIS. Betrachte dazu RIEMANN-Summen bezüglich Zerlegungen $a = x_0 < x_1 < \dots < x_N = b$, für die $x_n = c$ für irgendein $n \in \{0, \dots, N\}$. Die entsprechende RIEMANN-Summe lautet dann

$$\sum_{k=1}^N f(x_k)(x_k - x_{k-1}) = \sum_{k=1}^n f(x_k)(x_k - x_{k-1}) + \sum_{k=n+1}^N f(x_k)(x_k - x_{k-1}),$$

und auf der rechten Seite stehen RIEMANN-Summen für die Intervalle $[a, c]$ und $[c, b]$. Die Behauptung ergibt sich wieder im Limes kleiner Feinheiten. $\square$

Wir etablieren noch folgende Konvention: Ist $a = b$, so setzen wir $\int_a^b f(x)dx := 0$, und für $a > b$ setzen wir

$$\int_a^b f(x)dx := -\int_b^a f(x)dx.$$

Mit dieser Konvention gilt Proposition 4.32 für beliebige reelle Zahlen a, b, c , ungeachtet ihrer Anordnung.

Damit läßt sich das Integral auf *stückweise stetige* Funktionen erweitern: Sind $a = \xi_0 < \xi_1 < \dots < \xi_l = b$ endlich viele Punkte in $[a, b]$, und ist $f : [a, b] \rightarrow \mathbb{R}$ stetig in jedem Teilintervall (ξ_j, ξ_{j+1}) (aber womöglich unstetig in den Punkten ξ_k), und existieren schließlich in jedem ξ_j die links- und rechtsseitigen Grenzwerte $\lim_{x \searrow \xi_j} f(x)$ und $\lim_{x \nearrow \xi_j} f(x)$ als reelle Zahlen, so definiert man

$$\int_a^b f(x)dx := \sum_{j=1}^l \int_{\xi_{j-1}}^{\xi_j} f(x)dx.$$

Man sieht leicht, daß alle bisher gezeigten Eigenschaften des Integrals auch für stückweise stetige Integranden weiterhin gelten.

PROPOSITION 4.33. Sei $f \geq 0$ stetig auf $[a, b]$. Dann ist

$$\int_a^b f(x)dx \geq 0 \tag{4.9}$$

mit Gleichheit genau dann, wenn f identisch null ist.

BEWEIS. Die Eigenschaft (4.9) folgt sofort aus $f \geq 0$ und der Monotonie des Integrals. Ist f identisch null, so ist sein Integral ebenfalls null. Sei schließlich f nicht identisch null, dann gibt es $x_0 \in [a, b]$ mit $M := f(x_0) > 0$ und wegen der Stetigkeit von f ein $\delta > 0$, sodaß $f(x) \geq \frac{M}{2}$ für alle $x \in (x_0 - \delta, x_0 + \delta) \cap [a, b]$. Mit anderen Worten: f ist größer oder gleich der stückweise stetigen Funktion, die in $x \in (x_0 - \delta, x_0 + \delta) \cap [a, b]$ den Wert $\frac{M}{2}$ und sonst den Wert null annimmt. Es folgt, wieder mithilfe der Monotonie,

$$\int_a^b f(x)dx \geq \frac{M}{2} \min\{\delta, b - a\} > 0.$$

$\square$

4.3.2. Der Mittelwertsatz der Integralrechnung.

SATZ 4.34 (Mittelwertsatz der Integralrechnung). Seien $f, g \in C([a, b])$ und $g \geq 0$ in $[a, b]$ (oder $g \leq 0$ in $[a, b]$). Dann existiert ein $\xi \in [a, b]$, sodaß

$$\int_a^b f(x)g(x)dx = f(\xi) \int_a^b g(x)dx.$$

BEWEIS. Wir behandeln nur den Fall $g \geq 0$, da $g \leq 0$ nach Übergang zu $-g$ folgt. Seien M bzw. m das Maximum bzw. Minimum der stetigen Funktion f auf dem kompakten Intervall $[a, b]$. Da $g \geq 0$, gilt mit Proposition 4.30 die Ungleichungskette

$$m \int_a^b g(x) dx \leq \int_a^b f(x)g(x) dx \leq M \int_a^b g(x) dx.$$

Daher existiert $\mu \in [m, M]$, sodaß $\int_a^b f(x)g(x) dx = \mu \int_a^b g(x) dx$: In der Tat, ist g identisch null, so sind alle Integrale null und die Wahl von $\mu \in [m, M]$ ist beliebig; ist dagegen g nicht identisch null, so ist nach Proposition 4.33 $\int_a^b g(x) dx > 0$, und wir können somit

$$\mu = \frac{\int_a^b f(x)g(x) dx}{\int_a^b g(x) dx}$$

wählen.

Nach dem Zwischenwertsatz (hier geht die Stetigkeit von f ein) existiert nun ein $\xi \in [a, b]$ mit $f(\xi) = \mu$. $\square$

Ein wichtiger Spezialfall ergibt sich für $g \equiv 1$: Für eine stetige Funktion $f : [a, b] \rightarrow \mathbb{R}$ existiert $\xi \in [a, b]$, sodaß

$$\int_a^b f(x) dx = f(\xi)(b - a).$$

Dies erklärt auch den Namen des Satzes: Der *Mittelwert* der Funktion f auf dem Intervall $[a, b]$ ist nämlich gegeben durch

$$\frac{1}{b - a} \int_a^b f(x) dx,$$

und der Mittelwertsatz der Integralrechnung besagt, dass dieser Mittelwert mindestens einmal angenommen wird (sofern f stetig ist).

Die Beziehung zwischen dem Mittelwertsatz der Integralrechnung und dem der Differenzialrechnung wird durch den Hauptsatz der Differential- und Integralrechnung, den wir im nächsten Abschnitt besprechen, eluzidiert.

4.4. Der Hauptsatz der Differential- und Integralrechnung

4.4.1. Stammfunktionen. Sei $I \subset \mathbb{R}$ ein Intervall. Eine differenzierbare Funktion $F : I \rightarrow \mathbb{R}$ heißt *Stammfunktion* von $f : I \rightarrow \mathbb{R}$, wenn $F' = f$ in I . Für gegebenes f ist die Stammfunktion bis auf eine additive Konstante eindeutig bestimmt: In der Tat, ist $F' = G' = f$ in I , so ist $(F - G)' = 0$, und mit Korollar 4.18 erhalten wir $F = G + C$ für eine Konstante C .

Auf der Menge der Abbildungen $I \rightarrow \mathbb{R}$ definieren wir deshalb die Äquivalenzrelation

$$F \sim G \quad \text{genau dann, wenn } F - G \text{ konstant.} \quad (4.10)$$

DEFINITION 4.35 (Unbestimmtes Integral). Sei F eine Stammfunktion einer Funktion $f : I \rightarrow \mathbb{R}$. Das *unbestimmte Integral*

$$\int f(x) dx$$

ist definiert als die Äquivalenzklasse bezüglich der Äquivalenzrelation (4.10), die F enthält.

Man schreibt häufig $\int f(x) dx = F + C$, um zu verdeutlichen, daß die Stammfunktion nur modulo einer Konstanten bestimmt ist. Das zuvor eingeführte Integral $\int_a^b f(x) dx$ wird zur Unterscheidung manchmal als *bestimmtes Integral* bezeichnet. Die Beziehung zwischen diesen beiden Begriffen wird durch den Hauptsatz der Differential- und Integralrechnung erklärt:

4.4.2. Der Hauptsatz.

SATZ 4.36 (Hauptsatz der Differential- und Integralrechnung, Version 1). Seien $a, b \in \mathbb{R}$ mit $a < b$ und $f : [a, b] \rightarrow \mathbb{R}$ stetig, dann ist die Abbildung $F : [a, b] \rightarrow \mathbb{R}$,

$$x \mapsto F(x) := \int_a^x f(t) dt,$$

eine Stammfunktion von f .

BEMERKUNG 4.37. Nach Proposition 4.32 kann die untere Integrationsgrenze a durch eine beliebige andere Zahl in $[a, b]$ ersetzt werden.

BEWEIS. Nach Proposition 4.32 gilt für $x \in [a, b]$ und $h \in \mathbb{R}$ mit $x + h \in [a, b]$:

$$\frac{F(x+h) - F(x)}{h} = \frac{1}{h} \int_x^{x+h} f(t) dt.$$

Nach dem Mittelwertsatz der Integralrechnung (Satz 4.34) existiert ein $\xi = \xi_h \in [x, x+h]$ mit

$$f(\xi_h) = \frac{1}{h} \int_x^{x+h} f(t) dt.$$

Für $h \rightarrow 0$ gilt (wegen $\xi_h \in [x, x+h]$) $\xi_h \rightarrow x$ und, da f stetig ist, auch $f(\xi_h) \rightarrow f(x)$. Es folgt die Differenzierbarkeit von F an der Stelle x und $F'(x) = f(x)$, wie behauptet. $\square$

SATZ 4.38 (Hauptsatz der Differential- und Integralrechnung, Version 2). Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und F eine Stammfunktion von f . Dann gilt

$$\int_a^b f(x) dx = F(b) - F(a).$$

BEWEIS. Nach Satz 4.36 ist die durch $G(x) := \int_a^x f(t) dt$ definierte Abbildung eine Stammfunktion von f auf $[a, b]$, und es gilt

$$\int_a^b f(x) dx = G(b). \quad (4.11)$$

Da Stammfunktionen derselben Funktion sich nur um eine Konstante unterscheiden, gibt es für eine beliebige andere Stammfunktion F ein $C \in \mathbb{R}$ mit $F = G + C$. Da $G(a) = 0$, haben wir $C = F(a) - G(a) = F(a)$, und es folgt mit (4.11)

$$\int_a^b f(x) dx = G(b) = F(b) - C = F(b) - F(a).$$

$\square$

Für die Differenz $F(b) - F(a)$ verwendet man oft die Schreibweise $F|_a^b$, sodaß der Hauptsatz folgendermaßen in Gestalt der Beziehung zwischen bestimmtem und unbestimmtem Integral formuliert werden kann:

$$\int_a^b f(x) dx = \int f(x) dx \Big|_a^b.$$

Nun können wir auch den Bezug zwischen den Mittelwertsätzen der Differential- bzw. Integralrechnung herstellen: Ist nämlich $f : [a, b] \rightarrow \mathbb{R}$ stetig und F eine Stammfunktion, dann besagt der Mittelwertsatz der Differentialrechnung, angewendet auf F , daß es ein $\xi \in [a, b]$ (sogar in (a, b)) gibt mit

$$F(b) - F(a) = (b - a)F'(\xi) = (b - a)f(\xi).$$

Dies ist, nach Maßgabe des Hauptsatzes, äquivalent zu

$$\int_a^b f(x)dx = (b-a)f(\xi),$$

also erfüllt dieses ξ die Aussage des Mittelwertsatzes der *Integralrechnung*, angewendet auf f ! Wir sehen an diesem Argument übrigens auch, daß das ξ aus dem Mittelwertsatz der Integralrechnung (mit $g \equiv 1$) sogar im offenen Intervall (a, b) gewählt werden kann.

4.4.3. Integrationstechniken. Der Hauptsatz ermöglicht es in manchen Fällen, Integrale umzuformen oder gar explizit ihren Wert zu ermitteln. Zwei Techniken sind dabei fundamental: Die *Substitution* und die *partielle Integration*.

SATZ 4.39 (Substitution). Sei $I \subset \mathbb{R}$ ein kompaktes Intervall, $f : I \rightarrow \mathbb{R}$ stetig und $\phi : [a, b] \rightarrow I$ stetig differenzierbar. Dann gilt

$$\int_a^b f(\phi(z))\phi'(z)dz = \int_{\phi(a)}^{\phi(b)} f(x)dx. \quad (4.12)$$

BEWEIS. Sei F eine Stammfunktion von f (eine solche existiert nach Satz 4.36). Nach Kettenregel ist $(F \circ \phi)' = (f \circ \phi)\phi'$, sodaß nach Satz 4.38

$$\int_a^b f(\phi(z))\phi'(z)dz = \int_a^b (F(\phi(z)))'dz = F(\phi(b)) - F(\phi(a)) = \int_{\phi(a)}^{\phi(b)} f(x)dx.$$

□

Die Formel (4.12) kann ‚von links nach rechts‘ gelesen werden oder umgekehrt: Im ersteren Falle möchte man eine ‚komplizierte‘ Funktion vereinfachen, indem man die innere Funktion ϕ durch eine neue Variable substituiert; im zweiten Fall stellt man die Integrationsvariable als geschickt gewählte Funktion einer anderen Variable dar, um zum Ziel zu kommen. Wir geben je ein Beispiel und greifen dabei auf Schulwissen über Exponential- und Winkelfunktionen zurück:

BEISPIEL 4.40. (1) Wir wollen das Integral $\int_0^1 xe^{-x^2}dx$ bestimmen. Da wir für die Exponentialfunktion eine Stammfunktion kennen (nämlich sie selbst), liegt es nahe, $\phi(x) = -x^2$ zu wählen, dann gilt nämlich nach (4.12)

$$\begin{aligned} \int_0^1 xe^{-x^2}dx &= \int_0^1 \left(-\frac{1}{2}\phi'(x)\right)e^{\phi(x)}dx \\ &= -\frac{1}{2} \int_0^{-1} e^u du = -\frac{1}{2}(e^{-1} - e^0) = \frac{1}{2}(1 - e^{-1}). \end{aligned}$$

(2) Das Integral $2 \int_{-1}^1 \sqrt{1-x^2}dx$ gibt den Flächeninhalt der Einheitskreisscheibe an. Wir setzen $x = \phi(z) = \sin z$, dann ist $\phi' = \cos$, und es gilt gemäß der Substitutionsregel

$$\begin{aligned} 2 \int_{-1}^1 \sqrt{1-x^2}dx &= 2 \int_{\sin(-\pi/2)}^{\sin(\pi/2)} \sqrt{1-x^2}dx = 2 \int_{-\pi/2}^{\pi/2} \sqrt{1-(\sin z)^2} \cos z dz \\ &= 2 \int_{-\pi/2}^{\pi/2} (\cos z)^2 dz, \end{aligned}$$

da $\cos \geq 0$ auf dem Intervall $[-\frac{\pi}{2}, \frac{\pi}{2}]$. Als Spezialfall des Additionstheorems für den Kosinus¹³ erhalten wir

$$(\cos z)^2 = \frac{1}{2}(1 + \cos(2z)),$$

¹³ $\cos(x+y) = \cos x \cos y - \sin x \sin y$.

sodaß, wieder unter Verwendung des Hauptsatzes,

$$2 \int_{-1}^1 \sqrt{1-x^2} dx = \int_{-\pi/2}^{\pi/2} [1 + \cos(2z)] dz = \pi + \frac{1}{2} [\sin \pi - \sin(-\pi)] = \pi,$$

wobei man $\sin \pi = \sin(-\pi) = 0$ beachte.

BEMERKUNG 4.41. Landläufig wird die Zahl π genau als die Fläche der Einheitskreisscheibe *definiert*; wir werden stattdessen (wie in der akademischen Analysis üblich) $\frac{\pi}{2}$ als kleinste positive reelle Zahl definieren, für die der Sinus (den wir seinerseits natürlich ordentlich definieren werden) den Wert 1 annimmt, woraus dann mit obiger Rechnung die Charakterisierung von π als Fläche der Einheitskreisscheibe folgt.

SATZ 4.42 (partielle Integration). Seien $f, g : [a, b] \rightarrow \mathbb{R}$ stetig differenzierbar, so gilt

$$\int_a^b f'(x)g(x)dx = fg|_a^b - \int_a^b f(x)g'(x)dx. \quad (4.13)$$

BEWEIS. Nach LEIBNIZ-Regel gilt $(fg)' = f'g + fg'$. Integration beider Seiten über $[a, b]$ und Anwendung von Satz 4.38 auf $(fg)'$ ergibt wie gewünscht

$$f(b)g(b) - f(a)g(a) = \int_a^b f'(x)g(x)dx + \int_a^b f(x)g'(x)dx.$$

□

BEISPIEL 4.43. Sei mit $\log$ der natürliche Logarithmus, also die Umkehrfunktion der Exponentialfunktion, bezeichnet (Sie kennen ihn hoffentlich aus der Schule und wissen, daß $(\log x)' = \frac{1}{x}$). Für $x \in (0, \infty)$ schreiben wir $\int_1^x \log t dt = \int_1^x 1 \cdot \log t dt$ und setzen in Formel (4.13) $f(t) = t$ und $g(t) = \log t$. Da $\log'(t) = \frac{1}{t}$ und $f' = 1$, erhalten wir

$$\int_1^x 1 \cdot \log t dt = [t \log t]_1^x - \int_1^x t \cdot \frac{1}{t} dt = x \log x - (x - 1).$$

Insbesondere ist $x \mapsto x(\log x - 1)$ eine Stammfunktion von $\log$.

4.5. Uneigentliche Integrale

Wir relaxieren nun die Annahme, das Intervall $[a, b]$ sei kompakt, und betrachten nunmehr ein *offenes* Intervall (a, b) , wobei auch die Werte $a = -\infty$ oder $b = +\infty$ zulässig sind.

DEFINITION 4.44 (Uneigentliches Integral). Sei $f : (a, b) \rightarrow \mathbb{R}$ stetig. Dann heißt f auf (a, b) *uneigentlich integrierbar*, falls es ein $c \in (a, b)$ gibt, sodaß die Limites

$$\lim_{\alpha \searrow a} \int_{\alpha}^c f(x)dx \quad \text{und} \quad \lim_{\beta \nearrow b} \int_c^{\beta} f(x)dx \quad (4.14)$$

existieren. In diesem Falle heißt

$$\int_a^b f(x)dx := \lim_{\alpha \searrow a} \int_{\alpha}^c f(x)dx + \lim_{\beta \nearrow b} \int_c^{\beta} f(x)dx \quad (4.15)$$

das *uneigentliche Integral* von f von a bis b .

Uneigentliche Integrierbarkeit auf halboffenen Intervallen $[a, b)$ bedarf keiner Stützstelle c : Eine auf $[a, b)$ stetige Funktion ist demnach uneigentlich integrierbar, wenn $\lim_{\beta \nearrow b} \int_a^{\beta} f(x)dx$ existiert.

Man sagt im Falle der uneigentlichen Integrierbarkeit auch, das Integral $\int_a^b f(x)dx$ *konvergiere*.

Natürlich müssen wir uns davon überzeugen, daß das uneigentliche Integral unabhängig von der Wahl von c ist. In der Tat existieren aber nach Proposition 4.32 die Limites (4.14)

automatisch für jedes $c \in (a, b)$, sofern dies für eines der Fall ist, und die Summe (4.15) ist für alle solche c gleich.

Ist f nicht nur in (a, b) , sondern sogar auf $[a, b]$ stetig, so stimmt selbstverständlich das uneigentliche Integral von f über (a, b) mit dem Integral von f über $[a, b]$ überein: Dies folgt aus der Abschätzung

$$\left| \int_{\beta}^b f(x) dx \right| \leq (b - \beta) \sup_{x \in [a, b]} |f(x)| \rightarrow 0$$

für $\beta \nearrow b$, und analog für $\int_a^{\alpha} f(x) dx$. (Hier haben wir die Dreiecksungleichung verwendet.)

BEISPIEL 4.45. Sei $s \in \mathbb{Q}^{14}$. Das Integral $\int_1^{\infty} \frac{1}{x^s} dx$ konvergiert genau dann, wenn $s > 1$. In der Tat erhalten wir mithilfe des Hauptsatzes für jedes $1 < \beta < \infty$ und für $s \neq 1$:

$$\int_1^{\beta} x^{-s} dx = \frac{1}{1-s} x^{1-s} \Big|_1^{\beta} = \frac{1}{1-s} [\beta^{1-s} - 1],$$

was für $\beta \nearrow \infty$ im Falle $s > 1$ konvergiert (nämlich gegen $\frac{1}{s-1}$) und im Falle $s < 1$ divergiert. Im Falle $s = 1$ dagegen gilt

$$\int_1^{\beta} x^{-1} dx = \log x \Big|_1^{\beta} = \log \beta,$$

was für $\beta \nearrow \infty$ wiederum divergiert.

BEISPIEL 4.46. Das Integral $\int_0^1 \frac{1}{x^s} dx$ konvergiert genau dann, wenn $s < 1$: Wieder folgt mit dem Hauptsatz für jedes $0 < \alpha < 1$ und für $s \neq 1$:

$$\int_{\alpha}^1 x^{-s} dx = \frac{1}{1-s} x^{1-s} \Big|_{\alpha}^1 = \frac{1}{1-s} [1 - \alpha^{1-s}],$$

was für $\alpha \searrow 0$ im Falle $s < 1$ konvergiert (nämlich gegen $\frac{1}{1-s}$) und im Falle $s > 1$ divergiert. Im Falle $s = 1$ gilt

$$\int_{\alpha}^1 x^{-1} dx = \log x \Big|_{\alpha}^1 = -\log \alpha,$$

was für $\alpha \searrow 0$ divergiert.

Die beiden Beispiele in Kombination zeigen insbesondere, daß $\int_0^{\infty} \frac{1}{x^s} dx$ niemals konvergiert.

Ein warnendes Beispiel ist die Funktion $x \mapsto x$, die auf $\mathbb{R}$ nicht integrierbar ist: Zwar existiert $\lim_{\beta \nearrow \infty} \int_{-\beta}^{\beta} x dx = 0$, aber für jedes $c \in \mathbb{R}$ ist $\int_c^{\beta} x dx$ divergent, ebenso wie $\int_{-\beta}^c x dx$. Beide Intervallgrenzen müssen also unabhängig voneinander approximiert werden.

4.6. Der Satz von TAYLOR

Wir haben gesehen, daß eine differenzierbare Funktion um einen fest gewählten Punkt linear approximiert werden kann, d.h. man findet ein Polynom ersten Grades, das in der Nähe dieses Punktes eine gute Annäherung an die gegebene Funktion darstellt (z.B. $\sin x \sim x$ für kleine x). Kann man, wenn man eine höhere Approximationsgüte erreichen will, auch quadratisch, kubisch oder allgemeinen mit einem Polynom beliebigen Grades approximieren? Falls die Funktion genügend oft differenzierbar ist, lautet die Antwort „ja“:

SATZ 4.47 (TAYLOR). Sei I ein Intervall, $a \in I$, und $f \in C^{m+1}(I)$ für ein $m \in \mathbb{N}$. Dann gilt für jedes $x \in I$:

$$f(x) = \sum_{k=0}^m \frac{f^{(k)}(a)}{k!} (x-a)^k + \frac{1}{m!} \int_a^x f^{(m+1)}(t) (x-t)^m dt. \quad (4.16)$$

¹⁴Das Ergebnis dieses Beispiels bleibt auch für $s \in \mathbb{R}$ gültig, aber reelle Exponenten haben wir noch nicht eingeführt.

BEMERKUNG 4.48. Die Summe auf der rechten Seite ist ein Polynom m -ten Grades in x und heißt *TAYLOR-Polynom m -ter Ordnung*; das Integral heißt *Restglied* (in Integraldarstellung).

BEWEIS. Wir führen Induktion nach m . Für $m = 0$ lautet die Behauptung

$$f(x) = f(a) + \int_a^x f'(t) dt,$$

und dies ist genau der Hauptsatz der Differential- und Integralrechnung. Angenommen also, (4.16) gilt für irgendein $m \in \mathbb{N}$. Ist f auf I sogar $(m+2)$ -mal stetig differenzierbar, so erhalten wir mit partieller Integration für $R_m(x) := \frac{1}{m!} \int_a^x f^{(m+1)}(t)(x-t)^m dt$:

$$\begin{aligned} R_m(x) &= \frac{1}{m!} \int_a^x f^{(m+1)}(t)(x-t)^m dt \\ &= \frac{1}{(m+1)!} \int_a^x f^{(m+2)}(t)(x-t)^{m+1} dt + \frac{1}{(m+1)!} f^{(m+1)}(a)(x-a)^{m+1}, \end{aligned}$$

und der Satz ist bewiesen. $\square$

KOROLLAR 4.49 (Restglied in LAGRANGE-Form). Sei I ein Intervall, $a \in I$, und $f \in C^{m+1}(I)$ für ein $m \in \mathbb{N}$. Dann existiert für jedes $x \in I$ ein $\xi \in [a, x]$, sodaß

$$f(x) = \sum_{k=0}^m \frac{f^{(k)}(a)}{k!} (x-a)^k + \frac{f^{(m+1)}(\xi)}{(m+1)!} (x-a)^{m+1}.$$

BEWEIS. Wir wenden auf das Restglied den Mittelwertsatz der Integralrechnung, Satz 4.34, mit $g(t) = (x-t)^m$ an. Beachte, daß diese Wahl von g zulässig ist, da g auf dem Integrationsbereich $[a, x]$ immer dasselbe Vorzeichen hat. Also existiert $\xi \in [a, x]$, sodaß

$$\frac{1}{m!} \int_a^x f^{(m+1)}(t)(x-t)^m dt = \frac{f^{(m+1)}(\xi)}{m!} \int_a^x (x-t)^m dt = \frac{f^{(m+1)}(\xi)}{(m+1)!} (x-a)^{m+1}.$$

$\square$

KOROLLAR 4.50. Sei I ein Intervall, $a \in I$, und $f \in C^m(I)$ für ein $m \in \mathbb{N}$. Dann gilt

$$f(x) = \sum_{k=0}^m \frac{f^{(k)}(a)}{k!} (x-a)^k + \eta(x)(x-a)^m,$$

wobei $\eta: I \rightarrow \mathbb{R}$ eine Funktion ist mit $\lim_{x \rightarrow a} \eta(x) = 0$.

BEWEIS. Unter Verwendung des LAGRANGE-Restglieds für das TAYLOR-Polynom $m-1$ -ter Ordnung gilt für ein $\xi \in [a, x]$

$$\begin{aligned} f(x) - \sum_{k=0}^{m-1} \frac{f^{(k)}(a)}{k!} (x-a)^k &= \frac{f^{(m)}(\xi)}{m!} (x-a)^m \\ &= \frac{f^{(m)}(a)}{m!} (x-a)^m + \left[\frac{f^{(m)}(\xi)}{m!} - \frac{f^{(m)}(a)}{m!} \right] (x-a)^m. \end{aligned}$$

Mit der Wahl $\eta(x) = \frac{f^{(m)}(\xi)}{m!} - \frac{f^{(m)}(a)}{m!}$ folgt die Behauptung, denn da $\xi = \xi(x)$ im Intervall $[a, x]$ liegt und $f^{(m)}$ stetig ist, folgt $\lim_{x \rightarrow a} \eta(x) = 0$. $\square$

Dieses Korollar besagt also, daß das TAYLOR-Polynom m -ten Grades unter allen Polynomen höchstens m -ten Grades die Funktion in der Nähe von a am besten approximiert.

BEISPIEL 4.51. Betrachte $f: [-1, +\infty) \rightarrow \mathbb{R}$, $x \mapsto \sqrt{1+x}$, und setze $a = 0$. Das TAYLOR-Polynom erster Ordnung lautet

$$f(0) + f'(0)x = 1 + \frac{1}{2}x.$$

Entwicklung bis zu zweiter Ordnung ergibt

$$f(0) + f'(0)x + f''(0)x^2 = 1 + \frac{1}{2}x - \frac{1}{4}x^2.$$

Ist (wie in diesem Beispiel) die Funktion am Entwicklungspunkt beliebig oft differenzierbar, so kann man TAYLOR-Polynome beliebiger Ordnung bilden. Falls die Restglieder mit wachsender Ordnung immer kleiner werden, könnte man also

$$f(x) = \sum_{k=0}^{\infty} \frac{f^{(k)}(a)}{k!} (x-a)^k$$

schreiben (dies ist die sogenannte TAYLOR-Reihe). Man hätte damit die Funktion f gewissermaßen als ‚Polynom vom Grad unendlich‘ dargestellt. Doch was bedeutet es, unendlich viele Terme aufzusummieren, und wann gilt tatsächlich $R_m(x) \rightarrow 0$ für $m \rightarrow \infty$? Diese Fragen diskutieren wir im folgenden Kapitel.

KAPITEL 5

Potenzreihen

Eine besonders einfach zu handhabende Klasse von Funktionen stellen die Polynomfunktionen dar. Man kann sie zum Beispiel mühelos differenzieren und integrieren. Eine Erweiterung des Polynombegriffs erhält man, wie am Ende des vorigen Kapitels angedeutet, indem man Polynome ‚unendlichen Grades‘ heranzieht, also Funktionen der Form

$$f(x) = \sum_{k=0}^{\infty} c_k (x - x_0)^k,$$

für geeignete Koeffizienten c_k . Während eine Polynomfunktion stets wohldefiniert ist, stellt sich für solche *Potenzreihen* die Frage nach der Konvergenz; so ist etwa die unendliche Summe (wir definieren das noch genauer) $\sum_{k=0}^{\infty} x^k$ für $x = 1$ nicht endlich.

Als Vorbereitung diskutieren wir deshalb allgemein die Konvergenz von *Reihen*, also unendlicher Summen. Als Kollateraleffekt erhalten wir die Dezimaldarstellung der reellen Zahlen, mit deren Hilfe wir unser Verständnis von $\mathbb{R}$ noch vertiefen können.

Schließlich stellen wir spezielle Potenzreihen vor, durch die man die Exponentialfunktion und ihre Verwandten (Logarithmen und Winkelfunktionen) definieren kann. Die Beziehung zwischen Exponential- und Winkelfunktionen wird nur in den *komplexen Zahlen* klar; überhaupt ist ein gutes Verständnis *analytischer Funktionen*, also solcher, die als Potenzreihe darstellbar sind, nur im Komplexen möglich. Wir beginnen dieses Kapitel deshalb mit einer Einführung in die komplexen Zahlen.

5.1. Komplexe Zahlen

Beim Aufbau des Zahlensystems sind wir stets von Gleichungen ausgegangen, die keine Lösung besaßen, und haben den Zahlenraum entsprechend erweitert. So hatte die Gleichung $x + 2 = 1$ keine Lösung in $\mathbb{N}$, weswegen wir die ganzen Zahlen $\mathbb{Z}$ eingeführt haben; die Gleichung $2x = 3$ hatte wiederum in $\mathbb{Z}$ keine Lösung, was zur Erweiterung auf die rationalen Zahlen $\mathbb{Q}$ führte. Schließlich erwuchs die Motivation für die Einführung der reellen Zahlen $\mathbb{R}$ unter anderem aus dem Wunsch, Gleichungen wie $x^2 = 2$ zu lösen. Doch auch in $\mathbb{R}$ können unlösbare Gleichungen gefunden werden – die einfachste unter ihnen ist $x^2 = -1$. In der Tat, nach Korollar 1.41 sind Quadrate reeller Zahlen stets nichtnegativ. Dies zeigt bereits, daß Körper, in denen $x^2 = -1$ lösbar ist, nicht angeordnet sein können.

Die komplexen Zahlen $\mathbb{C}$ erlauben die Lösung von $x^2 = -1$. Kann man nun ad infinitum weitere unlösbare Gleichungen finden und somit Zahlenräume erweitern? Zum Glück ist dies nicht mehr nötig, da die komplexen Zahlen *algebraisch abgeschlossen* sind, das heißt: Jede algebraische Gleichung, also eine Gleichung der Form

$$x^n + c_{n-1}x^{n-1} + \dots + c_1x + c_0 = 0$$

vom Grad $n \geq 1$ mit Koeffizienten $c_k \in \mathbb{C}$, besitzt mindestens eine Lösung in $\mathbb{C}$. Dies ist die Aussage des *Fundamentalsatzes der Algebra*, dessen Beweis Sie wahrscheinlich in der Funktionentheorie¹ sehen werden.

¹Die Funktionentheorie (engl. complex analysis) behandelt komplex differenzierbare Funktionen von (einer Teilmenge von) $\mathbb{C}$ nach $\mathbb{C}$ und sollte daher besser, wie im Englischen, komplexe Analysis heißen.

5.1.1. Definition, Körpereigenschaften und Betrag. Wir skizzieren zunächst den naiven Zugang zu komplexen Zahlen. Da es keine reelle Lösung von $x^2 = -1$ gibt, *adjungiert* man einfach ein Objekt zu $\mathbb{R}$, das diese Gleichung löst. Man schreibt dafür i , sodaß also gilt $i^2 = -1$. Eine *komplexe Zahl* ist dann eine reelle Linearkombination von 1 und i , d.h. jede komplexe Zahl läßt sich eindeutig schreiben als $z = x + iy$ mit $x, y \in \mathbb{R}$ (man bezeichnet dann x als *Realteil* und y als *Imaginärteil* von z). Die Grundrechenarten sind dann so erklärt, als sei i lediglich eine Variable, die ggf. gemäß $i^2 = -1$ vereinfacht werden kann; das bedeutet: Sind $z_1 = x_1 + iy_1$ und $z_2 = x_2 + iy_2$ komplexe Zahlen, so setzt man

$$\begin{aligned} z_1 \pm z_2 &:= (x_1 \pm x_2) + i(y_1 \pm y_2), \\ z_1 z_2 &:= (x_1 x_2 - y_1 y_2) + i(x_1 y_2 + x_2 y_1), \\ \frac{z_1}{z_2} &:= \frac{x_1 x_2 + y_1 y_2}{x_2^2 + y_2^2} + i \frac{x_2 y_1 - x_1 y_2}{x_2^2 + y_2^2} \quad \text{falls } x_2^2 + y_2^2 > 0. \end{aligned}$$

(Für die Division haben wir den Bruch $\frac{x_1 + iy_1}{y_1 + iy_2}$ mit der Zahl $y_1 - iy_2$ erweitert.)

Dies ist nun alles mathematisch nicht ganz sauber, da die imaginäre Einheit i irgendwie ‚vom Himmel gefallen‘ ist. Formal definiert man daher wie folgt:

DEFINITION UND SATZ 5.1 (Der Körper der komplexen Zahlen). Eine *komplexe Zahl* ist ein Paar $(x, y) \in \mathbb{R}^2$. Sind $z_1 = (x_1, y_1)$ und $z_2 = (x_2, y_2)$ komplexe Zahlen, so definieren wir die Grundrechenarten durch

$$\begin{aligned} z_1 \pm z_2 &:= (x_1 \pm x_2, y_1 \pm y_2), \\ z_1 z_2 &:= (x_1 x_2 - y_1 y_2, x_1 y_2 + x_2 y_1), \\ \frac{z_1}{z_2} &:= \left(\frac{x_1 x_2 + y_1 y_2}{x_2^2 + y_2^2}, \frac{x_2 y_1 - x_1 y_2}{x_2^2 + y_2^2} \right) \quad \text{falls } x_2^2 + y_2^2 > 0. \end{aligned}$$

Mit diesen Operationen wird die Menge $\mathbb{C}$ ein Körper (erfüllt also die üblichen Rechengesetze der vier Grundrechenarten), wobei $0 := (0, 0)$ und $1 := (1, 0)$ die neutralen Elemente der Addition bzw. Multiplikation sind.

BEWEIS. Die Rechengesetze folgen aus denen für $\mathbb{R}$; zum Beispiel das Kommutativgesetz der Multiplikation:

$$z_1 z_2 = (x_1 x_2 - y_1 y_2, x_1 y_2 + x_2 y_1) = (x_2 x_1 - y_2 y_1, x_2 y_1 + x_1 y_2) = z_2 z_1.$$

Aus den Rechengesetzen sieht man außerdem sofort $z + 0 = z$ sowie $z \cdot 1 = z$ für alle $z \in \mathbb{C}$.

Ist $z = (x, y)$, so ist $z = 0$ äquivalent zu $x = 0$ und $y = 0$, also zu $x^2 + y^2 = 0$, und damit ist die Division für jeden Nenner ungleich Null definiert. Damit ist jede komplexe Zahl ungleich null invertierbar mit Inverser $\frac{1}{z}$, denn

$$z \cdot \frac{1}{z} = (x, y) \cdot \left(\frac{x}{x^2 + y^2}, \frac{-y}{x^2 + y^2} \right) = \left(\frac{x^2 + y^2}{x^2 + y^2}, \frac{-xy + yx}{x^2 + y^2} \right) = (1, 0) = 1.$$

Noch einfacher ist es zu prüfen, daß das additive Inverse von $z = (x, y)$ durch $-z := (-x, -y)$ gegeben ist. $\square$

Betrachtet man $\mathbb{R}$ als Teilmenge von $\mathbb{C}$, indem man $x \in \mathbb{R}$ mit $(x, 0) \in \mathbb{C}$ identifiziert, so prüft man leicht, daß die Definition der Grundrechenarten in $\mathbb{C}$ konsistent mit denen in $\mathbb{R}$ sind.

Wir schreiben ab sofort wieder $x + iy$ statt (x, y) . Für $z = x + iy \in \mathbb{C}$ definiert man die *komplex Konjugierte* durch $\bar{z} := x - iy$. Offenbar ist die komplexe Konjugation $z \mapsto \bar{z}$ ein Körperautomorphismus auf $\mathbb{C}$, d.h. sie ist eine Bijektion mit $\overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2$ und $\overline{z_1 z_2} = \bar{z}_1 \bar{z}_2$. Außerdem ist die Konjugation eine *Involution*, d.h. $\bar{\bar{z}} = z$.

Für eine komplexe Zahl $z = x + iy$ sind *Realteil* und *Imaginärteil* gegeben durch $\Re z := x$ und $\Im z := y$.

DEFINITION 5.2 (Betrag). Der Betrag einer komplexen Zahl $z = x + iy$ ist definiert als

$$|z| := \sqrt{z\bar{z}} = \sqrt{x^2 + y^2}.$$

Interpretiert man eine komplexe Zahl $x + iy$ als Punkt in der Ebene mit Koordinaten (x, y) (‘GAUSSsche Zahlenebene’), so ist ihr Betrag nach dem Satz des PYTHAGORAS gerade ihr Abstand vom Koordinatenursprung, oder mit anderen Worten der Betrag des Vektors $(x, y) \in \mathbb{R}^2$. Offenbar sind die Definitionen des Betrags in $\mathbb{R}$ und $\mathbb{C}$ konsistent. Der Betrag hat im Komplexen ähnliche Eigenschaften wie im Reellen (vgl. Proposition 1.43):

PROPOSITION 5.3. Seien $z_1, z_2 \in \mathbb{C}$. Dann gilt

- (1) $|z_1| \geq 0$, und $|z_1| = 0$ genau dann, wenn $z_1 = 0$;
- (2) $|z_1 z_2| = |z_1| |z_2|$;
- (3) $|z_1 + z_2| \leq |z_1| + |z_2|$.

BEWEIS. i) ist genau die Äquivalenz von $z = 0$ mit $x^2 + y^2 = 0$; ii) folgt mit

$$|z_1 z_2|^2 = z_1 z_2 \overline{z_1 z_2} = z_1 \bar{z}_1 z_2 \bar{z}_2 = |z_1|^2 |z_2|^2;$$

und für iii) beachten wir

$$\begin{aligned} |z_1 + z_2|^2 &= (z_1 + z_2)(\bar{z}_1 + \bar{z}_2) = z_1 \bar{z}_1 + z_2 \bar{z}_2 + z_1 \bar{z}_2 + \bar{z}_1 z_2 \\ &= |z_1|^2 + |z_2|^2 + z_1 \bar{z}_2 + \overline{z_1 \bar{z}_2} \\ &= |z_1|^2 + |z_2|^2 + 2\Re(z_1 \bar{z}_2) \\ &\leq |z_1|^2 + |z_2|^2 + 2|z_1||z_2| = (|z_1| + |z_2|)^2, \end{aligned}$$

wobei wir im vorletzten Schritt verwendet haben $\Re z \leq |z|$ und im Schritt davor $z + \bar{z} = 2\Re z$. $\square$

5.1.2. Konvergenz und Stetigkeit. Da also $\mathbb{C}$ sowohl die Körpereigenschaften als auch die wichtigsten Eigenschaften des Betrags mit $\mathbb{R}$ gemein hat, ist es nicht verwunderlich, daß alle Eigenschaften, die die Konvergenz betreffen, sich ebenso übertragen. Die Definition der konvergenten Folge in $\mathbb{C}$ ist wortgleich mit der reellen Version:

DEFINITION 5.4 (Konvergenz einer Folge). Eine Folge $(z_n)_{n \in \mathbb{N}}$ komplexer Zahlen heißt *konvergent* gegen ein $z \in \mathbb{C}$, wenn für alle $\epsilon > 0$ ein $N \in \mathbb{N}$ existiert, sodaß $|z_n - z| < \epsilon$ für alle $n \geq N$.

Man sieht sofort, daß die Rechenregeln für Grenzwerte (Satz 2.9) sich ebenfalls wortgleich auf $\mathbb{C}$ übertragen.

SATZ 5.5. Eine Folge $(z_n)_{n \in \mathbb{N}}$ konvergiert genau dann, wenn $(\Re z_n)_{n \in \mathbb{N}}$ und $(\Im z_n)_{n \in \mathbb{N}}$ konvergieren, und in diesem Falle ist

$$\lim_{n \rightarrow \infty} z_n = \lim_{n \rightarrow \infty} \Re z_n + i \lim_{n \rightarrow \infty} \Im z_n.$$

BEWEIS. Konvergiere $(z_n)_{n \in \mathbb{N}}$ gegen $z \in \mathbb{C}$, und sei $\epsilon > 0$. Wähle $N \in \mathbb{N}$ so groß, daß $|z_n - z| < \epsilon$ für $n \geq N$. Nach Definition des Betrags sieht man sofort $|\Re(z_n - z)|, |\Im(z_n - z)| \leq |z_n - z| < \epsilon$ für solche n , und somit konvergieren Real- und Imaginärteil von z_n gegen $\Re z$ bzw. $\Im z$, wie behauptet.

Seien nun umgekehrt $x_n := \Re z_n$ und $y_n := \Im z_n$ konvergent gegen x bzw. y , und sei wieder $\epsilon > 0$. Wähle $N \in \mathbb{N}$ so groß, daß $|x_n - x|, |y_n - y| < \frac{\epsilon}{2}$ für $n \geq N$. Dann gilt für solche n mit $z := x + iy$:

$$|z_n - z| = [(x_n - x)^2 + (y_n - y)^2]^{1/2} < \left[\frac{\epsilon^2}{4} + \frac{\epsilon^2}{4} \right]^{1/2} = \frac{\epsilon}{\sqrt{2}} < \epsilon.$$

$\square$

Die Definition der Cauchyfolge ist wortgleich dem reellen Fall, und da man nach obigem Satz Real- und Imaginärteil in Fragen der Konvergenz separat behandeln kann, ist jede komplexe Cauchyfolge konvergent. Ähnlich verhält es sich mit der Definition der Stetigkeit:

DEFINITION 5.6. Sei $U \subset \mathbb{C}$ und $f : U \rightarrow \mathbb{C}$ eine Funktion. Dann heißt f stetig in $z \in U$, wenn gilt:

$$\forall \epsilon > 0 \quad \exists \delta > 0 \quad \forall w \in U : \quad |z - w| < \delta \implies |f(z) - f(w)| < \epsilon.$$

Da $\mathbb{R}$ als Teilmenge von $\mathbb{C}$ aufgefaßt werden kann, ist damit auch die Stetigkeit von Abbildungen $\mathbb{C} \rightarrow \mathbb{R}$ oder $\mathbb{R} \rightarrow \mathbb{C}$ erklärt. Eine Funktion ist stetig genau dann, wenn

$$\lim_{w \rightarrow z} f(w) = f(z),$$

vergleiche Satz 3.5 über die Folgenstetigkeit.

Als einfache Übung zeige man, daß die Betragsfunktion $\mathbb{C} \rightarrow \mathbb{R}$, $z \mapsto |z|$ auf ganz $\mathbb{C}$ stetig ist, und daß daher $\lim_{n \rightarrow \infty} z_n = 0$ genau dann, wenn $\lim_{n \rightarrow \infty} |z_n| = 0$.

5.2. Reihen

5.2.1. Definitionen und Beispiele. Reihen sollten besser *unendliche Summen* heißen, denn es handelt sich dabei um Ausdrücke der Gestalt $\sum_{k=1}^{\infty} c_k$. Genauer definieren wir:

DEFINITION 5.7 (Reihen). Sei $(c_k)_{k \in \mathbb{N}}$ eine Folge komplexer Zahlen. Dann heißt die Folge der Partialsummen

$$(s_n)_{n \in \mathbb{N}}, \quad s_n := \sum_{k=1}^n c_k$$

die *Reihe* über $(c_k)_{k \in \mathbb{N}}$. Im Falle der Konvergenz der Partialsummen schreibt man

$$\sum_{k=1}^{\infty} c_k := \lim_{n \rightarrow \infty} s_n = \lim_{n \rightarrow \infty} \sum_{k=1}^n c_k$$

und sagt, die Reihe konvergiere. Ist sogar $\sum_{k=1}^{\infty} |c_k|$ konvergent, so sagt man, die Reihe über $(c_k)_{k \in \mathbb{N}}$ konvergiere *absolut*. Eine konvergente, aber nicht absolut konvergente Reihe heißt *bedingt konvergent*.

Da $\mathbb{C}$ vollständig ist, ist eine Reihe genau dann konvergent, wenn die Folge der Partialsummen Cauchy ist, wenn also zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$ existiert, sodaß für $N \leq n \leq m$ gilt:

$$|s_n - s_m| = \left| \sum_{k=n+1}^m c_k \right| < \epsilon.$$

Daraus folgt sofort: Jede absolut konvergente Reihe ist auch im gewöhnlichen Sinne konvergent, denn für $n \leq m \in \mathbb{N}$ gilt

$$|s_n - s_m| = \left| \sum_{k=n+1}^m c_k \right| \leq \sum_{k=n+1}^m |c_k|,$$

und letzere Summe wird aufgrund der Konvergenz von $\sum_{k=1}^{\infty} |c_k|$ beliebig klein, wenn nur n, m hinreichend groß sind.

Bevor wir einige Beispiele betrachten, halten wir ein triviales Konvergenzkriterium fest: Ist $\sum_{k=1}^{\infty} c_k$ konvergent, so gilt notwendigerweise $\lim_{k \rightarrow \infty} c_k = 0$, denn $c_k = s_k - s_{k-1}$, und letzteres konvergiert gegen null, da $(s_k)_{k \in \mathbb{N}}$ Cauchy ist. Die Umkehrung gilt freilich nicht: Wir werden gleich divergente Reihen kennenlernen, für die $\lim_{k \rightarrow \infty} c_k = 0$.

BEISPIEL 5.8. (1) Sei $q \in \mathbb{C}$. Die *geometrische Reihe* $\sum_{k=0}^{\infty} q^k$ konvergiert genau dann, wenn $|q| < 1$. Um dies einzusehen, betrachte die Partialsummen und schreibe sie in folgender Weise:

$$\begin{aligned}s_n &= 1 + q + q^2 + \dots + q^n \\ qs_n &= q + q^2 + \dots + q^n + q^{n+1}.\end{aligned}$$

Subtraktion auf beiden Seiten ergibt eine Kürzung aller Terme außer 1 und q^{n+1} , also

$$s_n(1 - q) = 1 - q^{n+1}$$

und daher, für $q \neq 1$ und $n \in \mathbb{N}$,

$$s_n = \frac{1 - q^{n+1}}{1 - q}.$$

Falls $|q| < 1$, so konvergiert dies gegen

$$\sum_{k=0}^{\infty} q^k = \frac{1}{1 - q}.$$

Ist dagegen $|q| \geq 1$, so auch $|q^k| = |q|^k \geq 1$ für alle k , und daher ist die Reihe in diesem Falle divergent, weil ja dann $\lim_{k \rightarrow \infty} q^k \neq 0$.

- (2) Die *harmonische Reihe* $\sum_{k=1}^{\infty} \frac{1}{k}$ ist divergent: Angenommen, sie konvergierte gegen einen Wert $s \in \mathbb{R}$, so gälte

$$\begin{aligned}s &= 1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \frac{1}{5} + \frac{1}{6} + \dots \\ &> \frac{1}{2} + \frac{1}{2} + \frac{1}{4} + \frac{1}{4} + \frac{1}{6} + \frac{1}{6} + \dots \\ &= 1 + \frac{1}{2} + \frac{1}{3} + \dots = s,\end{aligned}$$

Widerspruch. Die Ausarbeitung dieser Idee in Form eines strengen Beweises, der über Partialsummen argumentiert, ist zur Übung empfohlen.

- (3) Dagegen ist die *alternierende harmonische Reihe* $\sum_{k=1}^{\infty} (-1)^{k+1} \frac{1}{k}$ konvergent, wie wir bald zeigen werden. Da die Reihe ihrer Beträge, wie eben gezeigt, divergiert, ist sie jedoch nur bedingt konvergent.
- (4) Die Reihe $\sum_{k=1}^{\infty} (-1)^k$ ist divergent, da die Partialsummen zwischen 0 und -1 alternieren und somit nicht konvergent sind.

5.2.2. Konvergenzkriterien. Es gibt zahlreiche Kriterien dafür, daß eine Reihe konvergiert. Wir stellen die wichtigsten vor.

SATZ 5.9 (Integralvergleichskriterium). Sei f auf $[1, \infty)$ stetig, nichtnegativ und monoton fallend. Dann konvergiert die Reihe $\sum_{n=1}^{\infty} f(n)$ genau dann, wenn $\int_1^{\infty} f(x) dx$ konvergiert.

BEWEIS. Wegen der Monotonie gilt für alle $n \geq 1$

$$f(n) \geq f(x) \geq f(n+1), \quad x \in [n, n+1).$$

Aufgrund der Monotonie des Integrals ist für jedes $N \in \mathbb{N}$ daher

$$\sum_{n=1}^{N-1} f(n) \geq \int_1^N f(x) dx \geq \sum_{n=1}^{N-1} f(n+1). \quad (5.1)$$

Konvergiert nun die Reihe $\sum f(n)$, so folgt (wegen $f \geq 0$) $\sum_{n=1}^{N-1} f(n) \leq \sum_{n=1}^{\infty} f(n) < \infty$, und daher ist nach (5.1) $N \mapsto \int_1^N f(x) dx$ eine monoton steigende, durch $\sum_{n=1}^{\infty} f(n)$ beschränkte

Folge, die somit nach Korollar 2.30 konvergiert. Da außerdem $\beta \mapsto \int_1^\beta f(x)dx$ monoton steigt, existiert sogar $\lim_{\beta \nearrow \infty} \int_1^\beta f(x)dx$.

Ist umgekehrt das Integral $\int_1^\infty f(x)dx$ konvergent, so ist $\int_1^N f(x)dx \leq \int_1^\infty f(x)dx$ für alle N , und nach (5.1) ist die monoton steigende Folge $N \mapsto \sum_{n=1}^N f(n)$ durch $f(1) + \int_1^\infty f(x)dx$ nach oben beschränkt, und damit wieder nach Korollar 2.30 konvergent. $\square$

Unter Verwendung von Beispiel 4.45 sieht man mit dem Integralvergleichskriterium, daß $\sum_{n=1}^\infty \frac{1}{n^s}$ genau für $s > 1$ konvergiert. (Die Divergenz der harmonischen Reihe, also für den Fall $s = 1$, haben wir bereits ‚zu Fuß‘ gezeigt.)

SATZ 5.10 (Majorantenkriterium). Sei $|c_k| \leq d_k$ für alle $k \in \mathbb{N}$, und sei $\sum_{k=1}^\infty d_k$ konvergent. Dann ist $\sum_{k=1}^\infty c_k$ absolut konvergent.

BEWEIS. Sei $\epsilon > 0$. Da $\mathbb{R} \ni d_k \geq 0$, gibt es nach Annahme $N \in \mathbb{N}$, sodaß für $N \leq n \leq m$ gilt $\sum_{k=n+1}^m d_k < \epsilon$. Dann aber gilt für solche n, m auch

$$\left| \sum_{k=n+1}^m c_k \right| \leq \sum_{k=n+1}^m |c_k| \leq \sum_{k=n+1}^m d_k < \epsilon.$$

Also ist $\sum c_k$ Cauchy und damit konvergent. $\square$

BEISPIEL 5.11. Wir kommen auf die bereits genannte alternierende harmonische Reihe $\sum_{k=1}^\infty (-1)^{k+1} \frac{1}{k}$ zurück und zeigen ihre Konvergenz. Wir sehen dies wie folgt:

Für gerade Indizes haben wir für die Partialsummen

$$s_{2n} = \sum_{k=1}^{2n} (-1)^{k+1} \frac{1}{k} = \sum_{j=1}^n \left(\frac{1}{2j-1} - \frac{1}{2j} \right) = \sum_{j=1}^n \frac{1}{4j^2 - 2j},$$

und dies konvergiert nach dem Majorantenkriterium, da $\frac{1}{4j^2 - 2j} \leq \frac{1}{j^2}$. Andererseits gilt für die ungeraden Partialsummen

$$s_{2n+1} = s_{2n} + \frac{1}{2n+1},$$

und da $\frac{1}{2n+1} \rightarrow 0$, konvergieren die ungeraden Partialsummen gegen denselben Grenzwert wie die geraden. Also ist die alternierende harmonische Reihe konvergent.

SATZ 5.12 (Quotientenkriterium). Es existiere $\theta < 1$, sodaß gilt:

$$\left| \frac{c_{k+1}}{c_k} \right| \leq \theta \quad \text{für alle } k \in \mathbb{N}.$$

Dann ist $\sum_{k=1}^\infty c_k$ absolut konvergent.

BEWEIS. Für jedes $k \in \mathbb{N}$ gilt nach Voraussetzung $|c_{k+1}| \leq \theta |c_k|$. Durch Induktion zeigt man daraus leicht die Abschätzung

$$|c_k| \leq \theta^{k-1} |c_1| \quad \text{für alle } k \in \mathbb{N}.$$

Da $\sum_{k=1}^\infty |c_1| \theta^{k-1} = |c_1| \sum_{k=0}^\infty \theta^k$ als geometrische Reihe konvergent ist, gilt dies nach dem Majorantenkriterium auch (sogar im absoluten Sinne) für $\sum_{k=1}^\infty c_k$. $\square$

BEMERKUNG 5.13. Für das Majoranten- wie auch das Quotientenkriterium gilt: Für die Konvergenz genügt es sogar, daß die Voraussetzung jeweils für fast alle (also für alle bis auf endlich viele) $k \in \mathbb{N}$ erfüllt ist. Denn das Konvergenzverhalten einer Reihe ändert sich nicht durch Abänderung endlich vieler Glieder (der Wert der Reihe allerdings im Allgemeinen schon!).

BEISPIEL 5.14. Betrachten wir die Reihe $\sum_{n=1}^{\infty} \frac{n^s}{2^n}$ für beliebiges $s \in \mathbb{Q}$. Es gilt

$$\left| \frac{2^n(n+1)^s}{2^{n+1}n^s} \right| = \frac{1}{2} \left(\frac{n+1}{n} \right)^s.$$

Da $\lim_{n \rightarrow \infty} \frac{1}{2} \left(\frac{n+1}{n} \right)^s = \frac{1}{2}$, ist das Quotientenkriterium für hinreichend große Indizes z.B. mit $\theta = \frac{3}{4}$ erfüllt. Nach Quotientenkriterium und der anschließenden Bemerkung ist die Reihe also konvergent.

BEMERKUNG 5.15. Achtung: Es ist für die Konvergenz *nicht* hinreichend, daß $\left| \frac{c_{k+1}}{c_k} \right| < 1$ für alle k . Betrachte dazu erneut die harmonische Reihe, die divergiert, obwohl $\frac{c_{k+1}}{c_k} = \frac{k}{k+1} < 1$.

SATZ 5.16 (Wurzelkriterium). Es existiere $\theta < 1$, sodaß $\sqrt[k]{|c_k|} < \theta$ für (fast) alle $k \in \mathbb{N}$. Dann ist $\sum_{k=1}^{\infty} c_k$ absolut konvergent.

BEWEIS. Wegen $|c_k| < \theta^k$ ist die geometrische Reihe $\sum_{k=1}^{\infty} \theta^k$ eine konvergente Majorante. $\square$

BEMERKUNG 5.17. Das Wurzelkriterium kann äquivalent so formuliert werden: Ist

$$\limsup_{k \rightarrow \infty} \sqrt[k]{|c_k|} < 1,$$

so ist $\sum_{k=1}^{\infty} c_k$ absolut konvergent.

Ebenso für das Quotientenkriterium: Ist

$$\limsup_{k \rightarrow \infty} \left| \frac{c_{k+1}}{c_k} \right| < 1,$$

so konvergiert die Reihe absolut.

5.2.3. Dezimaldarstellung reeller Zahlen. Wir betrachten spezielle Reihen der Form $\pm \sum_{k=z}^{\infty} d_k 10^{-k}$, wobei $z \in \mathbb{Z}$ und $(d_k)_{k=z, \dots, \infty}$ eine Folge ganzer Zahlen zwischen 0 und 9 ist. Auf diese Weise gewinnen wir eine Dezimalzahl

$$d_z d_{z-1} \dots d_0, d_1 d_2 \dots,$$

zum Beispiel $1 = 1,0000\dots$ oder $\pi = 3,1415\dots$. Besteht die Dezimalentwicklung nicht nur aus Nullen, dann können wir ohne Beschränkung der Allgemeinheit $d_z \neq 0$ wählen (sonst würde die Dezimaldarstellung mit einer Null beginnen, die wir einfach streichen könnten), und dann gibt 10^{-z} die *Größenordnung* der so dargestellten Zahl an. Die Kreiszahl π hat also Größenordnung $10^0 = 1$, die Lichtgeschwindigkeit im Vakuum (gemessen in m/s) die Größenordnung 10^8 , und der Durchmesser des Wasserstoffatoms (in Metern) 10^{-11} .

Zunächst ist klar, daß jede Dezimaldarstellung absolut konvergiert: Dies folgt aus dem Majorantenkriterium wegen

$$|d_k \cdot 10^{-k}| \leq 9 \cdot 10^{-k},$$

und letzteres sind die Glieder einer geometrischen Reihe mit Basis $q = \frac{1}{10}$. Jede Dezimalzahl ist also eine reelle Zahl. Die Umkehrung gilt ebenfalls:

SATZ 5.18 (Dezimalentwicklung einer reellen Zahl). Jede reelle Zahl kann in der Form $\pm \sum_{k=z}^{\infty} d_k \cdot 10^{-k}$ dargestellt werden, wobei $z \in \mathbb{Z}$ und $(d_k)_{k=z, \dots, \infty}$ eine Folge ganzer Zahlen zwischen 0 und 9 ist.

BEWEIS. Wir müssen nur den Fall einer positiven reellen Zahl betrachten, denn für eine negative Zahl drehen wir einfach das Vorzeichen um, und die Null hat offenbar die Dezimaldarstellung $0,000\dots$

Sei also $0 < x \in \mathbb{R}$ und $z \in \mathbb{Z}$ die größte ganze Zahl, für die $x < 10^{-z+1}$. Wähle nun d_z als die größte ganze Zahl, für die $x \geq d_z \cdot 10^{-z}$. Dann ist d_z zwischen 1 und 9, denn für

$d_z = 0$ würde folgen $x < 10^{-z}$ im Widerspruch zur Maximalität von z , und für $d_z > 9$ würde folgen $x \geq 10^{-z+1}$, wiederum im Widerspruch zur Wahl von z . Damit ist die erste Ziffer der Dezimaldarstellung von x gefunden. Beachte

$$0 \leq x - d_z \cdot 10^{-z} < 10^{-z} \quad (5.2)$$

(eine Verletzung der zweiten Ungleichung stünde nämlich im Widerspruch zur Maximalität von d_z).

Seien bereits $d_z, d_{z+1}, \dots, d_n$ definiert ($n \geq z$). Als Induktionsannahme dürfen wir wegen (5.2) verwenden

$$0 \leq x - \sum_{k=z}^n d_k \cdot 10^{-k} < 10^{-n}.$$

Dann wählen wir d_{n+1} als die größte ganze Zahl, für die $x - \sum_{k=z}^n d_k \cdot 10^{-k} \geq d_{n+1} \cdot 10^{-(n+1)}$. Dann ist nach Induktionsannahme einerseits $d_{n+1} \geq 0$ und andererseits $d_{n+1} \leq 9$, da sonst $x - \sum_{k=z}^n d_k \cdot 10^{-k} \geq 10^{-n}$.

Wir halten fest, daß damit gilt

$$0 \leq x - \sum_{k=z}^{n+1} d_k \cdot 10^{-k} < 10^{-(n+1)}. \quad (5.3)$$

Mit dieser rekursiven Definition der d_k liegt somit eine wohldefinierte Dezimalreihe vor, die wegen (5.3) gegen x konvergiert, d.h.

$$x = \sum_{k=z}^{\infty} d_k \cdot 10^{-k}.$$

□

BEMERKUNG 5.19. Die Dezimalentwicklung einer reellen Zahl braucht nicht eindeutig bestimmt zu sein: so ist etwa²

$$1 = 1,000\dots = 0,999\dots$$

Letzteres sieht man leicht aus der Summenformel für die geometrische Reihe, denn

$$0,999\dots := \sum_{k=1}^{\infty} 9 \cdot 10^{-k} = 9 \left(\frac{1}{1 - \frac{1}{10}} - 1 \right) = 1.$$

Eine lohnenswerte Übung ist es zu zeigen, daß die Dezimaldarstellung auch *nur* in diesem Falle nichteindeutig ist, genauer: Besitzt eine reelle Zahl zwei unterschiedliche Dezimaldarstellungen, so bricht die eine nach endlich vielen Stellen ab, und die andere enthält nach endlich vielen Stellen nur die Ziffer 9. In diesem Falle bezeichnen wir diejenige Darstellung, die nach endlich vielen Stellen abbricht, als *kanonische Dezimaldarstellung*. Im obigen Beispiel wäre also $1,000\dots$ die kanonische Darstellung der Zahl 1. Die kanonische Dezimaldarstellung ist stets eindeutig.

BEMERKUNG 5.20. Aus mathematischer Sicht gibt es keinen Grund für die besondere Rolle, die die Basis 10 in der Darstellung reeller Zahlen spielt. Genausogut kann man eine beliebige natürliche Basis $b > 1$ wählen, und entsprechend

$$x = \sum_{k=z}^{\infty} d_k b^k$$

²Den Lesenden wird empfohlen, sich zu überlegen, welche der beiden Darstellungen von 1 man aus dem Beweis von Satz 5.18 erhält.

schreiben, wobei nun d_k ganze Zahlen zwischen 0 und $b-1$ sind. Eine besondere Bedeutung in vielen Anwendungen hat die auf LEIBNIZ zurückgehende *Binärdarstellung* ($b = 2$). So hat man zum Beispiel

$$2019_{10} = 11111100011_2,$$

wobei die Indizes 10 bzw. 2 anzeigen, daß die jeweilige Ziffernfolge im Dezimal- bzw. Binärsystem zu interpretieren ist³.

5.2.4. Die Überabzählbarkeit von $\mathbb{R}$. Als Konsequenz aus der Dezimaldarstellung können wir nun noch mehr Erkenntnisse über die Beschaffenheit der Menge $\mathbb{R}$ gewinnen. Genauer gesagt können wir zeigen, daß es ‚mehr‘ reelle als natürliche Zahlen gibt. Da es sich bei beiden um unendliche Mengen handelt, müssen wir zunächst klären, was das heißt. Die Ideen, die diesem Abschnitt zugrundeliegen, stammen von CANTOR aus dem späten 19. Jahrhundert.

DEFINITION 5.21 (Abzählbarkeit). Eine Menge M heißt *abzählbar*, wenn es eine Surjektion $\mathbb{N} \rightarrow M$ gibt.

Die Idee hinter dieser Definition liegt darin, daß eine Surjektion von einer *endlichen* Menge N auf eine andere Menge M genau dann existiert, wenn N mindestens so viele Elemente besitzt wie M . Man kann die Abzählbarkeit einer Menge also dahingehend auffassen, daß sie nicht ‚größer‘ ist als $\mathbb{N}$.

Jede endliche Menge ist abzählbar, denn wenn wir die Elemente von M mit $x_1, x_2, \dots, x_N$ bezeichnen (damit haben wir dann bereits eine ‚Abzählung‘ etabliert), so ist die Abbildung $s : \mathbb{N} \rightarrow M$ definiert durch

$$s(n) = \begin{cases} x_n & \text{falls } n \in \{1, 2, \dots, N\}, \\ x_1 & \text{sonst} \end{cases}$$

surjektiv.⁴

Doch auch unendliche Mengen können abzählbar sein: Selbstverständlich ist $\mathbb{N}$ selbst abzählbar. Paradoxerweise können auch Mengen, die $\mathbb{N}$ als echte Teilmengen enthalten (und damit ‚größer‘ sind als $\mathbb{N}$), abzählbar sein: Definiere etwa eine Surjektion $s : \mathbb{N} \rightarrow \mathbb{Z}$ durch

$$s(n) = \begin{cases} 0 & \text{falls } n = 1, \\ k & \text{falls } n = 2k \text{ für ein } k \in \mathbb{N}, \\ -k & \text{falls } n = 2k + 1 \text{ für ein } k \in \mathbb{N}. \end{cases}$$

Dies zeigt die Abzählbarkeit von $\mathbb{Z}$. Sogar $\mathbb{Q}$ ist abzählbar. Dazu zeigen wir zunächst eine allgemeine Aussage über abzählbare Mengen:

SATZ 5.22. Abzählbare Vereinigungen abzählbarer Mengen sind wieder abzählbar.

BEWEIS. Sei $(M_n)_{n \in \mathbb{N}}$ eine Familie abzählbarer Mengen. Sei für jedes $n \in \mathbb{N}$ die Abbildung $s_n : \mathbb{N} \rightarrow M_n$ surjektiv (eine solche Abbildung existiert nach Voraussetzung der Abzählbarkeit von M_n). Dann gilt

$$\mathcal{M} := \bigcup_{n \in \mathbb{N}} M_n = \{s_n(m) : n, m \in \mathbb{N}\}.$$

Wir müssen eine Surjektion $\mathbb{N} \rightarrow \mathcal{M}$ finden. Zunächst ist die Abbildung $(n, m) \mapsto s_n(m)$ eine Surjektion von $\mathbb{N}^2$ nach $\mathcal{M}$, sodaß es genügt, eine Surjektion $\mathbb{N} \rightarrow \mathbb{N}^2$ zu finden – mit

³Man beachte, daß diese Indizes ihrerseits dezimal dargestellt werden.

⁴Auch die leere Menge ist abzählbar, denn die leere Abbildung $\mathbb{N} \rightarrow \emptyset$ (definiert als die leere Menge, die als Teilmenge von $\mathbb{N} \times \emptyset$ betrachtet wird) ist surjektiv – schließlich ist ihr Bild genau die leere Menge. Diese Argumentation ist zugegebenermaßen hart an der Grenze zur Esoterik, aber trotzdem korrekt.

anderen Worten: eine Folge in $\mathbb{N}^2$, die jedes Element von $\mathbb{N}^2$ enthält. Eine solche Folge sieht etwa wie folgt aus:

$$(1, 1), (1, 2), (2, 1), (1, 3), (2, 2), (3, 1), (1, 4), (2, 3), (3, 2), (4, 1), \dots$$

Man zählt also der Reihe nach diejenigen Paare natürlicher Zahlen ab, deren Summe 2, dann 3, dann 4 usw. ist.⁵

□

KOROLLAR 5.23. $\mathbb{Q}$ ist abzählbar.

BEWEIS. Da

$$\mathbb{Q} = \bigcup_{q \in \mathbb{N}} \left\{ \frac{p}{q} : p \in \mathbb{Z} \right\}$$

und jede der Mengen $\left\{ \frac{p}{q} : p \in \mathbb{Z} \right\}$ bijektiv auf die abzählbare Menge $\mathbb{Z}$ abgebildet werden kann, folgt die Behauptung aus dem vorigen Satz. □

SATZ 5.24. $\mathbb{R}$ ist nicht abzählbar.

BEWEIS. Der Beweis wird durch Widerspruch erbracht und folgt dem ‚CANTORSchen Diagonalargument‘. Wir zeigen, daß sogar die Teilmenge $[0, 1] \subset \mathbb{R}$ nicht abzählbar ist.

Angenommen, dies wäre der Fall, und es gäbe eine Abzählung $(x_n)_{n \in \mathbb{N}}$ von $[0, 1]$. Nach Satz 5.18 kann jedes x_n in kanonischer Dezimaldarstellung (vgl. Bemerkung 5.19) als

$$x_n = \sum_{k=1}^{\infty} d_{n,k} \cdot 10^{-k}$$

geschrieben werden. Dann ist auch

$$x^* = \sum_{k=1}^{\infty} d_k \cdot 10^{-k}$$

eine reelle Zahl in $[0, 1]$, wobei wir d_k so wählen, daß

$$d_k \neq d_{k,k} \quad \text{und} \quad d_k \neq 9 \quad \text{für alle } k \in \mathbb{N}.$$

Da diese Darstellung von x^* keine 9 enthält, handelt es sich dabei um die eindeutig bestimmte kanonische Dezimaldarstellung von x^* (vgl. Bemerkung 5.19), die nach Wahl der d_k an der n -ten Stelle von der Dezimaldarstellung von x_n verschieden ist. Da kanonische Darstellungen eindeutig sind, folgt $x^* \neq x_n$ für alle $n \in \mathbb{N}$ – im Widerspruch zur Annahme, daß die Familie $(x_n)_{n \in \mathbb{N}}$ alle Elemente von $[0, 1]$ enthält. □

BEMERKUNG 5.25. In vielen Texten wird die Terminologie etwas anders verwendet als hier: Dort heißen nur solche Mengen abzählbar, die unendlich und in unserem Sinne abzählbar sind; nach dieser Konvention wären endliche Mengen also *nicht* abzählbar. Was bei uns abzählbar heißt, wird dann als *höchstens abzählbar* bezeichnet. Am sichersten fährt man, wenn man nur die Begriffe *abzählbar unendlich* und *höchstens abzählbar* verwendet, weil dann immer klar ist, was gemeint ist.

Zur weiteren Lektüre über die Inhalte dieses (mengentheoretischen) Abschnitts sei das Buch [6] empfohlen, das für Studienanfänger*innen bestens geeignet ist.

5.3. Gleichmäßige Konvergenz

Sei $U \subset \mathbb{C}$. Wir betrachten Funktionen $f : U \rightarrow \mathbb{C}$. Da $\mathbb{R}$ als Teilmenge von $\mathbb{C}$ aufgefaßt werden kann, gilt alles im Folgenden Gesagte ebenso für Funktionen von (einer Teilmenge von) $\mathbb{R}$ nach $\mathbb{R}$.

⁵Man möge sich dies in einem zweidimensionalen Koordinatensystem veranschaulichen.

5.3.1. Punktweise vs. gleichmäßige Konvergenz. Betrachte eine *Funktionenfolge*, also eine Folge $(f_n)_{n \in \mathbb{N}}$ von Abbildungen $f_n : U \rightarrow \mathbb{C}$. Was bedeutet es, daß eine solche Folge gegen eine Abbildung $f : U \rightarrow \mathbb{C}$ *konvergiert*? Im Gegensatz zum Fall von Zahlenfolgen gibt es hierauf zahlreiche gleichermaßen valide Antworten. Zwei davon gibt die folgende Definition:

DEFINITION 5.26 (punktweise und gleichmäßige Konvergenz). Sei $(f_n)_{n \in \mathbb{N}}$ eine Funktionenfolge.

- (1) Die Folge (f_n) konvergiert *punktweise* gegen $f : U \rightarrow \mathbb{C}$, falls

$$\forall x \in U : f(x) = \lim_{n \rightarrow \infty} f_n(x).$$

- (2) Die Folge (f_n) konvergiert *gleichmäßig* gegen $f : U \rightarrow \mathbb{C}$, falls

$$\lim_{n \rightarrow \infty} \sup_{x \in U} |f(x) - f_n(x)| = 0.$$

Schreibt man die Definition des Limes und des Supremums aus, erhält man als äquivalente Charakterisierung der punktweisen Konvergenz:

$$\forall \epsilon > 0 \quad \forall x \in X \quad \exists N \in \mathbb{N} \quad \forall n \geq N : |f(x) - f_n(x)| < \epsilon,$$

und für die gleichmäßige Konvergenz

$$\forall \epsilon > 0 \quad \exists N \in \mathbb{N} \quad \forall x \in X \quad \forall n \geq N : |f(x) - f_n(x)| < \epsilon.$$

Der einzige Unterschied besteht in der Reihenfolge der Quantoren: Im ersten Falle gibt es *für alle x ein N* , im zweiten gibt es *ein N für alle x* . Im ersten Falle darf also das N von x abhängen, im zweiten nicht (vgl. die Definition der gleichmäßigen Stetigkeit). Es folgt insbesondere, daß gleichmäßige Konvergenz punktweise Konvergenz impliziert. Es ist außerdem klar, daß der punktweise bzw. gleichmäßige Limes eindeutig ist, sofern er existiert.

In den folgenden Beispielen sei stets $U = [0, 1] \subset \mathbb{R}$.

BEISPIEL 5.27. Betrachte die Folge

$$f_n(x) = \begin{cases} n & \text{falls } x \in (0, \frac{1}{n}), \\ 0 & \text{sonst.} \end{cases}$$

Diese Folge konvergiert punktweise gegen die Nullfunktion, denn $f_n(0) = 0$ für alle n , und für jedes $x \in (0, 1]$ existiert N mit $\frac{1}{N} < x$, sodaß für alle Folgenglieder mit Index mindestens N gilt $f_n(x) = 0$. Die Folge konvergiert jedoch nicht gleichmäßig: Täte sie dies, müßten gleichmäßiger und punktweiser Limes übereinstimmen (da gleichmäßige Konvergenz punktweise Konvergenz impliziert, und zwar gegen die selbe Funktion). Die Folge konvergiert aber nicht gleichmäßig gegen null, denn $\sup_{x \in [0, 1]} |f_n(x)| = n$. Man beachte, daß für jedes $n \in \mathbb{N}$ das Integral von f_n gleich 1 ist, das Integral der punktweisen Limesfunktion hingegen null.

BEISPIEL 5.28. Die durch $f_n(x) = x^n$ definierte Folge konvergiert punktweise, aber nicht gleichmäßig gegen die Funktion gegeben durch

$$f(x) = \begin{cases} 0 & \text{falls } x \in [0, 1), \\ 1 & \text{falls } x = 1. \end{cases}$$

Sei nämlich $\epsilon > 0$, dann gibt es für jedes $x \in [0, 1)$ ein N , sodaß für alle Indizes ab N gilt $x^n < \epsilon$. Zudem ist $1^n = 1$ für alle n . Es folgt die punktweise Konvergenz. Andererseits existiert für jedes $n \in \mathbb{N}$ ein $x_n \in (0, 1)$ mit $f_n(x_n) = \frac{1}{2}$, denn für festes n ist $0^n = 0$ und $\lim_{x \nearrow 1} x^n = 1$, und die Existenz des gewünschten x_n folgt somit aus dem Zwischenwertsatz. Daher ist für jedes $n \in \mathbb{N}$ $\sup_{x \in [0, 1]} |f(x) - f_n(x)| \geq \frac{1}{2}$, und die Folge konvergiert nicht gleichmäßig.

Beachte, daß der punktweise Limes unstetig ist, obwohl jedes f_n stetig ist.

BEISPIEL 5.29. $f_n(x) = \sin(nx)$ konvergiert weder punktweise noch gleichmäßig (zum Beispiel weil $\sin(n\pi/4) = 0$, falls n Vielfaches von 4 ist, und $\sin(n\pi/4) = 1$, falls n die Form $8k+2$ hat).

Die Folge $g_n(x) = \frac{1}{n} \sin(nx)$ dagegen konvergiert gleichmäßig (und damit auch punktweise) gegen null. Die Folge der Ableitungen $g'_n(x) = \cos(nx)$ konvergiert wiederum *nicht* gegen null!

Die Frage nach der Verträglichkeit von Limites von Funktionenfolgen mit Differentiation und Integration wird uns im nächsten Abschnitt beschäftigen.

Wir werden auch *Funktionenreihen* betrachten, also Funktionenfolgen der Form $n \mapsto \sum_{j=1}^n f_j$. Da eine Reihe die Folge ihrer Partialsummen ist, sind durch Definition 5.26 auch punktweise bzw. gleichmäßige Konvergenz solcher Reihen erklärt. Man sagt außerdem, eine Funktionenreihe konvergiere *absolut*, wenn die Reihe $\sum_{j=1}^{\infty} |f_j|$ punktweise konvergiert.

BEISPIEL 5.30. Betrachte die Reihe $\sum_{j=0}^{\infty} \frac{x^j}{j!}$. Für jedes $x \in \mathbb{C}$ ergibt das Quotientenkriterium

$$\left| \frac{x^{j+1}}{(j+1)!} \frac{j!}{x^j} \right| = \frac{|x|}{j+1},$$

was für fast alle j kleiner als $\theta = \frac{1}{2}$ ist. Daher konvergiert diese Funktionenreihe absolut. Man nennt $\exp: \mathbb{C} \rightarrow \mathbb{C}$,

$$\exp(x) := \sum_{j=0}^{\infty} \frac{x^j}{j!}$$

die *Exponentialfunktion*⁶.

5.3.2. Eigenschaften gleichmäßiger Konvergenz.

SATZ 5.31 (Gleichmäßige Limites stetiger Funktionen sind stetig). Sei $(f_n)_{n \in \mathbb{N}}$ eine Folge stetiger Funktionen $f_n: U \rightarrow \mathbb{C}$, die gleichmäßig gegen $f: U \rightarrow \mathbb{C}$ konvergiert. Dann ist f stetig.

BEWEIS. Seien $\epsilon > 0$ und $x \in U$. Nach Definition der gleichmäßigen Konvergenz existiert $N \in \mathbb{N}$, sodaß für alle $y \in U$ gilt $|f(y) - f_N(y)| < \frac{\epsilon}{3}$. Da f_N stetig ist, existiert außerdem ein $\delta > 0$, sodaß aus $|x - y| < \delta$ folgt $|f_N(x) - f_N(y)| < \frac{\epsilon}{3}$. Ist also $y \in U$ mit $|x - y| < \delta$, so folgt nach Dreiecksungleichung

$$|f(x) - f(y)| \leq |f(x) - f_N(x)| + |f_N(x) - f_N(y)| + |f_N(y) - f(y)| < \frac{\epsilon}{3} + \frac{\epsilon}{3} + \frac{\epsilon}{3} = \epsilon,$$

d.h. f ist stetig. □

Beachte, daß punktweise Konvergenz für die Aussage des Satzes nicht ausreicht, wie in Beispiel 5.28 demonstriert wird.

SATZ 5.32 (Vertauschbarkeit von gleichmäßigem Limes und Integral). Sei $U = [a, b] \subset \mathbb{R}$ ein kompaktes Intervall mit $a < b$. Ist $(f_n)_{n \in \mathbb{N}}$ eine Folge stetiger Funktionen $f_n: [a, b] \rightarrow \mathbb{R}$, die gleichmäßig gegen $f: [a, b] \rightarrow \mathbb{R}$ konvergiert, so gilt

$$\int_a^b f(x) dx = \lim_{n \rightarrow \infty} \int_a^b f_n(x) dx.$$

⁶Der Zusammenhang zwischen diesem Konzept und dem, was Sie in der Schule als Exponentialfunktion kennengelernt haben, wird im nächsten Abschnitt klar.

BEWEIS. Nach dem vorigen Satz ist f stetig, also ist das Integral wohldefiniert. Es gilt die Abschätzung

$$\left| \int_a^b (f - f_n)(x) dx \right| \leq \int_a^b |f - f_n|(x) dx \leq (b - a) \sup_{x \in [a, b]} |f - f_n|(x) \rightarrow 0,$$

da nach Definition der gleichmäßigen Konvergenz gilt $\lim_{n \rightarrow \infty} \sup_{x \in [a, b]} |f - f_n|(x) = 0$. $\square$

Auch diese Aussage ist im allgemeinen falsch, wenn man gleichmäßige durch punktweise Konvergenz ersetzt (siehe Beispiel 5.27). Die Aussage ist für *uneigentliche* Integrale im allgemeinen ebenfalls falsch: Betrachte dazu auf $U = (0, \infty)$ die Funktionenfolge

$$f_n(x) = \begin{cases} \frac{1}{n} & \text{falls } x \in (0, n), \\ 0 & \text{sonst.} \end{cases}$$

Dann ist $\int_0^\infty f_n(x) dx = 1$ für jedes $n \in \mathbb{N}$, aber die Folge konvergiert gleichmäßig gegen null. Gewissermaßen hat die Folge unendlich viel Platz, um die Gesamtmasse ‚auszuschmieren‘.

SATZ 5.33 (Vertauschbarkeit von Limes und Ableitung). Sei $U \subset \mathbb{R}$ ein (nicht notwendig kompaktes) Intervall. Sei weiter $(f_n)_{n \in \mathbb{N}}$ eine Folge stetig differenzierbarer Funktionen $f_n : U \rightarrow \mathbb{R}$, die punktweise gegen $f : U \rightarrow \mathbb{R}$ konvergiert. Konvergiert die Folge der Ableitungen $(f'_n)_{n \in \mathbb{N}}$ gleichmäßig, so ist f stetig differenzierbar, und der (gleichmäßige) Limes der f'_n ist gleich f' .

BEWEIS. Sei $f^* : U \rightarrow \mathbb{R}$ der gleichmäßige Limes der Folge (f'_n) . Nach dem Hauptsatz gilt für jedes $n \in \mathbb{N}$ und $x \in U$

$$f_n(x) = f_n(a) + \int_a^x f'_n(t) dt, \quad (5.4)$$

wobei $a \in U$ beliebig gewählt ist.

Satz 5.32 impliziert $\lim_{n \rightarrow \infty} \int_a^x f'_n(t) dt = \int_a^x f^*(t) dt$, sodaß aus (5.4) und der punktweisen Konvergenz der f_n folgt

$$f(x) = f(a) + \int_a^x f^*(t) dt.$$

Nach Satz 5.31 ist f^* stetig, sodaß Ableiten auf beiden Seiten unter Verwendung von Satz 4.36 wie gewünscht $f' = f^*$ ergibt. $\square$

Im Folgenden wird es nützlich sein, für ein $f : U \rightarrow \mathbb{C}$ die Schreibweise

$$\|f\|_\infty := \sup_{x \in U} |f(x)|$$

zu verwenden. Dann ist $\|f\|_\infty < \infty$ genau dann, wenn f auf U beschränkt ist.

SATZ 5.34 (Konvergenzkriterium von WEIERSTRASS). Seien $f_n : U \rightarrow \mathbb{C}$ beschränkt für jedes $n \in \mathbb{N}$. Falls die Reihe

$$\sum_{n=1}^{\infty} \|f_n\|_\infty \quad (5.5)$$

konvergiert, so konvergiert die Reihe $\sum_{n=1}^{\infty} f_n$ absolut und gleichmäßig.

BEMERKUNG 5.35. Man vergleiche dies mit dem Majorantenkriterium.

BEWEIS. Sei $x \in U$. Da für alle $n \in \mathbb{N}$ gilt $\|f_n\|_\infty \geq |f_n(x)|$, ist nach dem Majorantenkriterium die Reihe $\sum_{n=1}^{\infty} |f_n(x)|$ konvergent, und es folgt die absolute Konvergenz der Funktionenreihe $\sum_{n=1}^{\infty} f_n$. Deshalb ist durch

$$F(x) := \sum_{n=1}^{\infty} f_n(x)$$

eine Funktion auf U wohldefiniert. Wir zeigen, daß die Reihe sogar *gleichmäßig* gegen F konvergiert: Ist nämlich $\epsilon > 0$, so existiert nach Voraussetzung (5.5) ein $N \in \mathbb{N}$, sodaß

$$\sum_{n=k+1}^{\infty} \|f_n\|_{\infty} < \epsilon \quad \text{für alle } k \geq N.$$

Daraus erhalten wir für $k \geq N$ und jedes $x \in X$ die Abschätzung

$$\left| F(x) - \sum_{n=0}^k f_n(x) \right| = \left| \sum_{n=k+1}^{\infty} f_n(x) \right| \leq \sum_{n=k+1}^{\infty} |f_n(x)| \leq \sum_{n=k+1}^{\infty} \|f_n\|_{\infty} < \epsilon,$$

und da dies für alle $x \in U$ gilt (mit einem von x unabhängigen N), folgt die gleichmäßige Konvergenz. $\square$

5.4. Potenzreihen

In diesem Abschnitt sei $U = \mathbb{C}$. In Beispiel 5.30 haben wir bereits gesehen, daß die Funktionenreihe

$$\sum_{n=0}^{\infty} \frac{z^n}{n!}$$

für jedes $z \in \mathbb{K}$ absolut konvergiert und damit auf ganz $\mathbb{C}$ eine Funktion definiert, die als Exponentialfunktion bezeichnet wird. Eine *Potenzreihe* ist allgemeiner eine Funktionenreihe der Form

$$\sum_{n=0}^{\infty} c_n (z - z_0)^n, \tag{5.6}$$

wobei $z_0, c_n \in \mathbb{C}$. Ein kleiner Hinweis zur Terminologie bzw. Notation: Der Ausdruck $\sum_{n=0}^{\infty} a_n$ hat genaugenommen zwei Bedeutungen – einerseits bezeichnet er die Folge der Partialsummen, also die Folge $(\sum_{n=0}^N a_n)_{N \in \mathbb{N}}$, und andererseits bezeichnet er den *Wert* der Reihe, d.h. die Zahl $\lim_{N \rightarrow \infty} \sum_{n=0}^N a_n$. In der ersten Bedeutung ist eine Potenzreihe $\sum_{n=0}^{\infty} c_n (z - z_0)^n$ also stets wohldefiniert, nämlich als Funktionenfolge $N \mapsto F_N : \mathbb{C} \rightarrow \mathbb{C}$, $F_N(z) = \sum_{n=0}^N c_n (z - z_0)^n$.

5.4.1. Konvergenzradius. Nicht immer konvergiert eine Potenzreihe auf ganz $\mathbb{C}$: Die Reihe $\sum_{n=1}^{\infty} \frac{z^n}{n}$ divergiert etwa an der Stelle $z = 1$, denn dort erhält man die harmonische Reihe. Die Reihe $\sum_{n=0}^{\infty} n! z^n$ divergiert sogar für jedes $z \neq 0$, wie man leicht mithilfe des Quotientenkriteriums sieht (es ist klar, daß jede Potenzreihe der Form (5.6) bei $z = z_0$ konvergiert, nämlich gegen c_0).

Wir schreiben wieder $B_{z_0}(r) := \{z \in \mathbb{K} : |z - z_0| < r\}$ für den offenen Kreis um z_0 mit Radius r , und $\overline{B_{z_0}}(r) := \{z \in \mathbb{K} : |z - z_0| \leq r\}$ für seinen Abschluß.

SATZ 5.36. Die Potenzreihe (5.6) konvergiere für ein $z_1 \in \mathbb{C}$. Ist $0 < r \in \mathbb{R}$ dergestalt, daß $r < |z_1 - z_0|$, so konvergiert die Potenzreihe auf $\overline{B_{z_0}}(r)$ absolut und gleichmäßig.

BEWEIS. Setze $f_n(z) := c_n (z - z_0)^n$. Da die Glieder einer konvergenten Reihe beschränkt sind, gibt es ein $M > 0$, sodaß $|f_n(z_1)| \leq M$ für alle $n \in \mathbb{N}$. Ist nun $z \in \overline{B_{z_0}}(r)$, so gilt deshalb

$$|f_n(z)| = |c_n (z - z_0)^n| = |c_n (z_1 - z_0)^n| \frac{|z - z_0|^n}{|z_1 - z_0|^n} \leq M q^n$$

mit

$$q = \frac{|z - z_0|}{|z_1 - z_0|} \leq \frac{r}{|z_1 - z_0|} < 1.$$

Auf $\overline{B_{z_0}}(r)$ gilt also

$$\sum_{n=0}^{\infty} \|f_n\|_{\infty} \leq M \sum_{n=0}^{\infty} q^n < \infty,$$

und das Kriterium von WEIERSTRASS (Satz 5.34) impliziert die absolute und gleichmäßige Konvergenz der Potenzreihe $\sum_{n=0}^{\infty} f_n$ auf $\overline{B_{z_0}}(r)$. $\square$

DEFINITION 5.37 (Konvergenzradius). Der *Konvergenzradius* der Potenzreihe (5.6) ist definiert durch

$$R = \sup\{|z - z_0| : \sum_{n=0}^{\infty} c_n(z - z_0)^n \text{ ist konvergent}\}.$$

Der Konvergenzradius kann auch null oder unendlich sein. Nach Satz 5.36 konvergiert eine Potenzreihe also absolut und gleichmäßig auf jeder Kreisscheibe $\overline{B_{z_0}}(r)$ mit $r < R$ und divergiert für jedes z mit $|z - z_0| > R$. Auf dem Rand $\{|z - z_0| = R\}$ sind dagegen verschiedene Szenarien möglich: Betrachte etwa die Reihen

$$\sum_{n=1}^{\infty} \frac{z^n}{n^2}, \quad \sum_{n=1}^{\infty} \frac{z^n}{n}, \quad \sum_{n=1}^{\infty} z^n.$$

Man erkennt leicht mithilfe des Quotientenkriteriums, daß der Konvergenzradius jeweils 1 ist. Die erste Reihe konvergiert auf $\{|z| = 1\}$ absolut, denn die Reihe $\sum_{n=1}^{\infty} \frac{1}{n^2}$ ist eine konvergente Majorante; die zweite Reihe divergiert für $z = 1$ (harmonische Reihe) und konvergiert für $z = -1$ (alternierende harmonische Reihe); und die dritte Reihe divergiert für jedes z mit $|z| = 1$, da die Glieder einer konvergenten Reihe stets gegen null konvergieren.

Wir geben schließlich eine Charakterisierung des Konvergenzradius an:

SATZ 5.38 (Formel von HADAMARD). Der Konvergenzradius der Potenzreihe (5.6) ist

$$R = \left(\limsup_{n \rightarrow \infty} \sqrt[n]{|c_n|} \right)^{-1}$$

mit der Konvention $0^{-1} = \infty$ und $\infty^{-1} = 0$.

BEWEIS. Setze $R^* := \left(\limsup_{n \rightarrow \infty} \sqrt[n]{|c_n|} \right)^{-1}$. Ist $|z - z_0| < R^*$, so ist

$$\limsup_{n \rightarrow \infty} \sqrt[n]{|c_n(z - z_0)^n|} = |z - z_0| \limsup_{n \rightarrow \infty} \sqrt[n]{|c_n|} < R^* \limsup_{n \rightarrow \infty} \sqrt[n]{|c_n|} = 1,$$

sodaß $\sum_{n=0}^{\infty} c_n(z - z_0)^n$ nach dem Wurzelkriterium konvergiert. Ist dagegen $|z - z_0| > R^*$, so divergiert die Reihe analog nach dem Wurzelkriterium. Insgesamt erhalten wir für den Konvergenzradius also $R = R^*$. $\square$

5.4.2. Differentiation und Integration von Potenzreihen. Nun ist es ein leichtes, zu zeigen, daß Potenzreihen innerhalb ihres Konvergenzradius *gliedweise* integriert und differenziert werden können:

SATZ 5.39 (gliedweise Integration). Seien $z_0 \in \mathbb{R}$, $c_n \in \mathbb{R}$ für $n \in \mathbb{N} \cup \{0\}$. Die Potenzreihe $\sum_{n=0}^{\infty} c_n(z - z_0)^n$ habe Konvergenzradius R . Dann hat auch die Reihe

$$\sum_{n=0}^{\infty} c_n \frac{(z - z_0)^{n+1}}{n+1} \tag{5.7}$$

den Konvergenzradius R , und für alle $[a, b] \in (z_0 - R, z_0 + R) \subset \mathbb{R}$ gilt

$$\int_a^b \left(\sum_{n=0}^{\infty} c_n(z - z_0)^n \right) dz = \sum_{n=0}^{\infty} c_n \frac{(z - z_0)^{n+1}}{n+1} \Big|_a^b. \tag{5.8}$$

BEWEIS. Nach Satz 5.38 bestimmt sich der Kehrwert des Konvergenzradius der Reihe (5.7) zu

$$\limsup_{n \rightarrow \infty} \left(\frac{|c_n|}{n+1} \right)^{\frac{1}{n}} = \limsup_{n \rightarrow \infty} \sqrt[n]{|c_n|} = \frac{1}{R},$$

denn $\lim_{n \rightarrow \infty} (n+1)^{1/n} = 1$ (Übungsblatt 8 Aufgabe 5). Daher stimmen die Konvergenzradien der beiden Potenzreihen überein.

Nach Satz 5.36 konvergiert $\sum_{n=0}^{\infty} c_n(z-z_0)^n$ gleichmäßig auf $[a, b]$. Da für jedes $N \in \mathbb{N}$

$$\int_a^b \left(\sum_{n=0}^N c_n(z-z_0)^n \right) dz = \sum_{n=0}^N c_n \frac{(z-z_0)^{n+1}}{n+1} \Big|_a^b,$$

gilt nach Satz 5.32 sogar (5.8). $\square$

BEISPIEL 5.40. Betrachte die Exponentialfunktion $\exp(x) := \sum_{n=0}^{\infty} \frac{x^n}{n!}$. Dann ist nach Satz 5.39 eine Stammfunktion gegeben durch

$$\int \exp(x) dx = \sum_{n=0}^{\infty} \frac{x^{n+1}}{(n+1)!} + C = \sum_{n=0}^{\infty} \frac{x^n}{n!} + (C-1) = \exp(x) + C^*.$$

Nach Hauptsatz folgt daraus $\exp' = \exp$. Dies könnte man aber auch aus dem folgenden Satz herleiten:

SATZ 5.41 (gliedweise Differentiation). Seien $z_0 \in \mathbb{R}$ und $c_n \in \mathbb{R}$ für alle $n \in \mathbb{N} \cup \{0\}$. Die Potenzreihe

$$F(z) := \sum_{n=0}^{\infty} c_n(z-z_0)^n$$

habe Konvergenzradius R . Dann hat auch die Reihe

$$\sum_{n=1}^{\infty} c_n n(z-z_0)^{n-1}$$

den Konvergenzradius R , und F ist in $(z_0 - R, z_0 + R)$ differenzierbar mit

$$F'(z) = \sum_{n=1}^{\infty} c_n n(z-z_0)^{n-1}. \quad (5.9)$$

BEWEIS. Nach Satz 5.38 bestimmt sich der Kehrwert des Konvergenzradius der Reihe (5.7) zu

$$\limsup_{n \rightarrow \infty} (n|c_n|)^{\frac{1}{n}} = \limsup_{n \rightarrow \infty} \sqrt[n]{|c_n|} = \frac{1}{R},$$

denn $\lim_{n \rightarrow \infty} n^{1/n} = 1$. Daher stimmen die Konvergenzradien der beiden Potenzreihen überein. Da außerdem nach Satz 5.36 beide Reihen innerhalb ihres Konvergenzradius gleichmäßig konvergieren, folgt aus Satz 5.33 die Differenzierbarkeit von F und die Gleichheit (5.9) in $(z_0 - R, z_0 + R)$. $\square$

Durch wiederholte Anwendung dieses Satzes erhält man sogar

KOROLLAR 5.42. Eine reelle Potenzreihe definiert innerhalb ihres Konvergenzradius eine beliebig oft differenzierbare Funktion.

KOROLLAR 5.43 (TAYLOR-Reihe). Seien wieder $z_0 \in \mathbb{R}$ und $c_n \in \mathbb{R}$ für alle $n \in \mathbb{N} \cup \{0\}$. Ist $F(z) = \sum_{n=0}^{\infty} c_n(z-z_0)^n$ eine Potenzreihe mit Konvergenzradius $R > 0$, so gilt für die Koeffizienten

$$c_n = \frac{1}{n!} F^{(n)}(z_0),$$

wobei $F^{(n)}$ die n -te Ableitung bezeichnet. Man kann F in $(z_0 - R, z_0 + R)$ also darstellen durch seine TAYLOR-Reihe

$$F(z) = \sum_{n=0}^{\infty} \frac{1}{n!} F^{(n)}(z_0)(z-z_0)^n.$$

BEWEIS. Nach Satz 5.41 kann die F definierende Potenzreihe gliedweise differenziert werden, und man sieht leicht $F(z_0) = c_0$, $F'(z_0) = c_1$, ..., $F^{(n)}(z_0) = n!c_n$. $\square$

Doch Vorsicht ist geboten: Das Korollar besagt lediglich, daß F seiner TAYLOR-Reihe gleicht, *wenn F überhaupt als Potenzreihe darstellbar ist*. Um festzustellen, ob dies der Fall ist, muß man die Restglieder der TAYLOR-Polynome geeignet abschätzen (s. Satz 4.47 und dessen Korollar). Eine notwendige Bedingung dafür ist nach Korollar 5.42 jedenfalls die beliebig häufige Differenzierbarkeit. Die Betragsfunktion etwa ist bei null nicht differenzierbar und somit in einer Umgebung von null auch nicht als Potenzreihe darstellbar. Doch die Bedingung der beliebig häufigen Differenzierbarkeit ist nicht hinreichend: In Übungsblatt 12, Aufgabe 6 wurde gezeigt, daß die Funktion

$$f(x) = \begin{cases} \exp\left(-\frac{1}{x}\right) & \text{für } x > 0 \\ 0 & \text{für } x \leq 0 \end{cases}$$

in $\mathbb{R}$ beliebig oft differenzierbar ist mit $f^{(n)}(0) = 0$ für alle $n \in \mathbb{N}$. Wäre f als Potenzreihe mit Entwicklungspunkt $z_0 = 0$ darstellbar, so wären also nach Korollar 5.43 alle Koeffizienten null, und es würde folgen $f \equiv 0$ in einer Umgebung von null, was offensichtlich falsch ist.

Abschließend sei erwähnt, daß alle Resultate dieses Unterabschnitts auch in $\mathbb{C}$ Bestand haben. Funktionen, die (in $\mathbb{R}$ oder $\mathbb{C}$) als Potenzreihen dargestellt werden können, heißen *analytisch*; das Studium analytischer Funktionen $\mathbb{C} \rightarrow \mathbb{C}$ ist Gegenstand der Funktionentheorie.

5.5. Spezielle Funktionen

Wir behandeln hier die Exponentialfunktion, den Logarithmus sowie Sinus und Kosinus. Es stellt sich heraus, daß diese Funktionen eng miteinander verwandt sind. Aus unserer Diskussion der Winkelfunktionen erhalten wir außerdem die Zahl π .

5.5.1. Logarithmus und allgemeine Potenz.

5.5.1.1. *Der natürliche Logarithmus.* Die Exponentialfunktion $\exp : \mathbb{R} \rightarrow \mathbb{R}$ ist durch die absolut und gleichmäßig konvergente Reihe

$$\exp(x) = \sum_{n=0}^{\infty} \frac{x^n}{n!}$$

definiert (siehe Beispiel 5.30). In Analysis I, Übungsblatt 14, Aufgabe 7, wurde gezeigt, daß $\exp$ streng monoton wachsend ist mit $\exp(\mathbb{R}) = (0, \infty)$. Daher besitzt die Exponentialfunktion eine Umkehrfunktion:

DEFINITION 5.44 (natürlicher Logarithmus). Die Umkehrfunktion von $\exp$ heißt *natürlicher Logarithmus*,

$$\log : (0, \infty) \rightarrow \mathbb{R}.$$

- PROPOSITION 5.45. a) *Der Logarithmus ist streng monoton steigend mit $\lim_{x \searrow 0} \log(x) = -\infty$ and $\lim_{x \nearrow \infty} \log(x) = +\infty$. Es gilt $\log(1) = 0$.*
 b) *Es gilt für $x, y > 0$ die Funktionalgleichung $\log(xy) = \log(x) + \log(y)$.*
 c) *Der Logarithmus ist auf ganz $(0, \infty)$ differenzierbar mit $\log'(x) = \frac{1}{x}$.*

BEWEIS. a) Dies folgt sofort aus den entsprechenden Eigenschaften der Exponentialfunktion: $\exp : \mathbb{R} \rightarrow (0, \infty)$ ist surjektiv und streng monoton steigend, außerdem haben wir $\lim_{x \rightarrow -\infty} \exp(x) = 0$ und $\lim_{x \rightarrow +\infty} \exp(x) = +\infty$, sowie schließlich $\exp(0) = 1$ (Analysis I Blatt 14 Aufgabe 7).

b) Seien $x, y > 0$, so setze $\xi := \log(x)$ und $\eta := \log(y)$. Dann folgt aus der Funktionalgleichung der Exponentialfunktion (Analysis I Blatt 14 Aufgabe 7b):

$$\exp(\xi + \eta) = \exp(\xi) \exp(\eta).$$

Nach Definition von $\log$ ist die rechte Seite gleich xy und linke Seite gleich $\exp(\log(x) + \log(y))$; Logarithmieren auf beiden Seiten liefert die gewünschte Identität.

c) Nach dem Satz über die Ableitung der Umkehrfunktion (Satz 4.10) ist $\log$ überall differenzierbar, und es gilt für jedes $x > 0$ wegen $\exp' = \exp$:

$$\log'(x) = \frac{1}{\exp'(\log(x))} = \frac{1}{\exp(\log(x))} = \frac{1}{x}.$$

□

5.5.1.2. Allgemeine Potenzen.

DEFINITION 5.46. Sei $a > 0$, so definiert man für $x \in \mathbb{R}$:

$$\exp_a(x) := \exp(x \log(a)).$$

PROPOSITION 5.47. Die Funktion $\exp_a : \mathbb{R} \rightarrow \mathbb{R}$ ist stetig und es gilt

- a) $\exp_a(x + y) = \exp_a(x) \exp_a(y)$ für alle $x, y \in \mathbb{R}$;
- b) $\exp_a(x) = a^x$ für alle $x \in \mathbb{Q}$.

BEWEIS. Wie in Analysis I, Übungsblatt 14, Aufgabe 7. □

Die Proposition besagt insbesondere, daß $\exp_a$ die (eindeutige) stetige Fortsetzung der Abbildung $\mathbb{Q} \rightarrow \mathbb{R}$, $x \mapsto a^x$ nach ganz $\mathbb{R}$ ist. Man schreibt deshalb auch für $x \in \mathbb{R}$ meistens $a^x := \exp_a(x)$.

5.5.1.3. Allgemeine Logarithmen.

DEFINITION 5.48. Sei $a > 0$, so definiert man für $x > 0$:

$$\log_a(x) := \frac{\log(x)}{\log(a)}.$$

PROPOSITION 5.49. $\log_a : (0, \infty) \rightarrow \mathbb{R}$ ist die Umkehrfunktion von $\exp_a : \mathbb{R} \rightarrow (0, \infty)$.

BEWEIS. Sei $x \in \mathbb{R}$, so gilt, da $\log$ die Umkehrfunktion von $\exp$ ist,

$$\log_a(\exp_a x) = \frac{\log(\exp(x \log(a)))}{\log(a)} = \frac{x \log(a)}{\log(a)} = x,$$

und außerdem gilt für $x > 0$

$$\exp_a(\log_a(x)) = \exp(\log_a(x) \log(a)) = \exp\left(\frac{\log(x)}{\log(a)} \log(a)\right) = \exp(\log(x)) = x.$$

□

5.5.2. Die komplexe Exponentialfunktion. Wir erinnern uns erneut an die Definition der Exponentialfunktion (Beispiel 5.30), die auch in $\mathbb{C}$ gültig ist:

DEFINITION 5.50. Die Funktion $\exp : \mathbb{C} \rightarrow \mathbb{C}$ ist gegeben durch

$$\exp(z) = \sum_{n=0}^{\infty} \frac{z^n}{n!}.$$

Man schreibt für $\exp(z)$ oft auch e^z (vgl. Analysis I, Übungsblatt 14, Aufgabe 7c).

Wir haben bereits gesehen, daß diese Potenzreihe auf ganz $\mathbb{C}$ absolut konvergiert. Außerdem ist die Exponentialfunktion stetig, denn als Potenzreihe ist sie innerhalb ihres Konvergenzradius (also auf ganz $\mathbb{C}$) lokal gleichmäßig konvergent (Satz 5.36), aber der gleichmäßige Limes einer Folge stetiger Funktionen ist wieder stetig (Satz 5.31).

Eine Reihe weiterer Eigenschaften der Exponentialfunktion ist Ihnen aus Analysis I (Übungsblatt 14, Aufgabe 7) bekannt, allerdings nur im Reellen. Die wichtigste dieser Eigenschaften ist die *Funktionalgleichung* $e^{w+z} = e^w e^z$. Im Reellen haben wir dies über die gliedweise Differenzierbarkeit von Potenzreihen bewiesen, für die wir wiederum den

Hauptsatz der Differential- und Integralrechnung verwendet haben. Eine solche Argumentation kann man in $\mathbb{C}$ in ähnlicher Weise verfolgen, aber wir wollen hier nicht im Komplexen differenzieren und integrieren (das passiert dann in der Vorlesung Funktionentheorie). Stattdessen leiten wir die Funktionalgleichung über das CAUCHY-Produkt her:

LEMMA 5.51 (CAUCHY-Produkt). *Seien $\sum_{j=0}^{\infty} a_j$ und $\sum_{k=0}^{\infty} b_k$ konvergente Reihen, von denen eine sogar absolut konvergiert. Dann konvergiert auch die Reihe*

$$\sum_{n=0}^{\infty} c_n, \quad \text{wobei} \quad c_n := \sum_{l=0}^n a_l b_{n-l},$$

und es gilt

$$\sum_{n=0}^{\infty} c_n = \left(\sum_{j=0}^{\infty} a_j \right) \left(\sum_{k=0}^{\infty} b_k \right).$$

BEWEIS. Sei o.B.d.A. die Reihe $\sum_{j=0}^{\infty} a_j$ absolut konvergent. Wir schreiben für die Partialsummen

$$A_n := \sum_{j=0}^n a_j, \quad B_n := \sum_{k=0}^n b_k, \quad C_n := \sum_{l=0}^n c_l,$$

und für die Limites $A := \sum_{j=0}^{\infty} a_j$ und $B := \sum_{k=0}^{\infty} b_k$.

Da endliche Summen beliebig umgeordnet werden dürfen (Kommutativität der Addition), erhalten wir

$$C_n = \sum_{l=0}^n \sum_{j=0}^l a_j b_{l-j} = \sum_{l=0}^n \sum_{j=0}^l a_{n-l} b_j = \sum_{l=0}^n a_{n-l} B_l = \sum_{l=0}^n a_{n-l} (B_l - B) + A_n B.$$

Daher ist

$$\begin{aligned} |C_n - AB| &= \left| \sum_{l=0}^n a_{n-l} (B_l - B) + (A_n - A) B \right| \\ &\leq \sum_{l=0}^n |a_{n-l}| |B_l - B| + |A_n - A| |B|. \end{aligned} \tag{5.10}$$

Sei nun $\epsilon > 0$. Da $A_n \rightarrow A$ nach Voraussetzung, gibt es ein N so groß, daß $|A_n - A| < \frac{\epsilon}{3|B|}$ für alle $n \geq N$.

Da $B_l \rightarrow B$ und – wegen der absoluten Konvergenz – $\sum_{n=0}^{\infty} |a_n| < \infty$, gibt es ein L so groß, daß für $l \geq L$ gilt

$$|B_l - B| < \frac{\epsilon}{3(1 + \sum_{k=0}^{\infty} |a_k|)}.$$

Schließlich gibt es (wieder wegen der absoluten Konvergenz) ein M so groß, daß

$$|a_m| < \frac{\epsilon}{3L(1 + \max_{l=0, \dots, L-1} |B_l - B|)}$$

für alle $m \geq M$.

Sei $n \geq \max\{L + M, N\}$, dann können wir (5.10) also aufspalten als

$$\begin{aligned} |C_n - AB| &\leq \sum_{l=0}^n |a_{n-l}| |B_l - B| + |A_n - A| |B| \\ &\leq \sum_{l=0}^{L-1} |a_{n-l}| |B_l - B| + \sum_{l=L}^n |a_{n-l}| |B_l - B| + |A_n - A| |B| < \epsilon, \end{aligned}$$

und es folgt wie behauptet $\lim_{n \rightarrow \infty} C_n = AB$. $\square$

KOROLLAR 5.52 (Funktionalgleichung der Exponentialfunktion). Für alle $w, z \in \mathbb{C}$ gilt $e^{w+z} = e^w e^z$.

BEWEIS. Die Zahlen e^w und e^z sind durch absolut konvergente Reihen dargestellt, also gilt nach vorigem Lemma

$$e^w e^z = \left(\sum_{j=0}^{\infty} \frac{w^j}{j!} \right) \left(\sum_{l=0}^{\infty} \frac{z^l}{l!} \right) = \sum_{n=0}^{\infty} c_n,$$

wobei

$$c_n = \sum_{k=0}^n \frac{w^k}{k!} \frac{z^{n-k}}{(n-k)!} = \frac{1}{n!} \sum_{k=0}^n \binom{n}{k} w^k z^{n-k} = \frac{(w+z)^n}{n!},$$

womit bereits $e^w e^z = e^{w+z}$ gezeigt ist. $\square$

KOROLLAR 5.53. Für jedes $z \in \mathbb{C}$ ist $e^z \neq 0$.

BEWEIS. Nach Funktionalgleichung ist $1 = e^0 = e^z e^{-z}$. Wäre $e^z = 0$, würde sich somit ein Widerspruch ergeben. $\square$

Wir erinnern uns: Ist $z = x + iy \in \mathbb{C}$, so heißt $\bar{z} = x - iy$ die zu z *komplex Konjugierte*.

PROPOSITION 5.54. Für alle $z \in \mathbb{C}$ ist $\overline{e^z} = e^{\bar{z}}$.

BEWEIS. Für die n -te Partialsumme der Exponentialreihe gilt, da die Konjugation mit Summen und Produkten verträglich ist,

$$\overline{\sum_{j=0}^n \frac{z^j}{j!}} = \sum_{j=0}^n \frac{\bar{z}^j}{j!}. \quad (5.11)$$

Die rechte Seite konvergiert mit $n \rightarrow \infty$ gegen $e^{\bar{z}}$. Mit Satz 5.5 sieht man leicht, daß eine Folge komplexer Zahlen genau dann gegen z konvergiert, wenn die Folge der Konjugierten gegen $\bar{z}$ konvergiert; daraus folgt, daß die linke Seite von (5.11) mit $n \rightarrow \infty$ gegen $\overline{e^z}$ konvergiert, und die Behauptung folgt. $\square$

PROPOSITION 5.55. Für jedes $x \in \mathbb{R}$ gilt $|e^{ix}| = 1$.

BEWEIS. Nach Definition des Betrags einer komplexen Zahl und mit der vorigen Proposition sowie der Funktionalgleichung ist

$$|e^{ix}|^2 = e^{ix} \overline{e^{ix}} = e^{ix} e^{-ix} = e^0 = 1.$$

$\square$

Die Zahlen e^{ix} liegen also in der komplexen Ebene auf dem Einheitskreis. Wir werden bald sehen, daß dabei x den Winkel (im Bogenmaß) angibt, den e^{ix} mit der reellen Achse einschließt.

5.5.3. Winkelfunktionen. Mit *Winkelfunktionen* oder *trigonometrischen* Funktionen meint man hauptsächlich Sinus, Kosinus und Tangens; in staubigen Büchern finden Sie außerdem Sekans, Kosekans und Kotangens. Diese Funktionen wurden seit der Spätantike zur Berechnung ebener Dreiecke verwendet (mit Anwendungen in der Navigation und Astronomie), in der Neuzeit dann auch zur Beschreibung von Schwingungsphänomenen. Letzteres werden Sie vermutlich in der Vorlesung über gewöhnliche Differentialgleichungen sehen.

5.5.3.1. *Definition und Reihendarstellung.* Die folgende Definition ist allerdings nicht unmittelbar geometrisch:

DEFINITION 5.56 (Sinus, Kosinus). Die Funktionen $\sin, \cos : \mathbb{R} \rightarrow \mathbb{R}$ sind definiert durch

$$\cos x = \Re e^{ix}, \quad \sin x = \Im e^{ix}.$$

Insbesondere gilt $e^{ix} = \cos x + i \sin x$.

Daraus ergibt sich mit Proposition 5.55 sofort, daß $|\sin x|, |\cos x| \leq 1$ für alle $x \in \mathbb{R}$, und zudem $\sin^2 x + \cos^2 x = |e^{ix}| = 1$.

SATZ 5.57 (Reihenentwicklung). Sinus und Kosinus sind durch die auf ganz $\mathbb{R}$ konvergenten Potenzreihen

$$\cos x = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n}}{(2n)!}, \quad \sin x = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{(2n+1)!}$$

gegeben⁷.

BEWEIS. Bemerke zunächst

$$i^n = \begin{cases} (-1)^k, & \text{falls } n = 2k \text{ für ein } k \in \mathbb{N} \cup \{0\} \\ (-1)^k i, & \text{falls } n = 2k+1 \text{ für ein } k \in \mathbb{N} \cup \{0\}. \end{cases}$$

Da die Exponentialreihe auf ganz $\mathbb{R}$ absolut konvergiert, dürfen wir gerade und ungerade Indizes getrennt summieren (Analysis I, Übungsblatt 13, Aufgabe 9a). Also erhalten wir

$$\begin{aligned} e^{ix} &= \sum_{n=0}^{\infty} i^n \frac{x^n}{n!} \\ &= \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!} + \sum_{k=0}^{\infty} (-1)^k i \frac{x^{2k+1}}{(2k+1)!} \\ &= \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k}}{(2k)!} + i \sum_{k=0}^{\infty} (-1)^k \frac{x^{2k+1}}{(2k+1)!}. \end{aligned}$$

Vergleich von Real- und Imaginärteil ergibt die Behauptung. $\square$

Es folgt sofort, daß der Kosinus eine *gerade* Funktion ist (d.h. $\cos(-x) = \cos x$ für alle x) und der Sinus *ungerade* (d.h. $\sin(-x) = -\sin x$ für alle x). Der Graph von $\cos$ ist also achsensymmetrisch um die vertikale Achse, der Graph von $\sin$ ist dagegen punktsymmetrisch um den Ursprung.

KOROLLAR 5.58 (Ableitungen). Sinus und Kosinus sind auf ganz $\mathbb{R}$ beliebig oft differenzierbar, und es gilt

$$\sin' = \cos, \quad \cos' = -\sin.$$

BEWEIS. Die Differenzierbarkeit ergibt sich aus Korollar 5.42. Die Formeln für die Ableitung erhält man durch gliedweise Differentiation (Satz 5.41), denn

$$\sin' x = \left(\sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{(2n+1)!} \right)' = \sum_{n=0}^{\infty} (-1)^n (2n+1) \frac{x^{2n}}{(2n+1)!} = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n}}{(2n)!} = \cos x$$

und ähnlich für $\cos'$. $\square$

Insonderheit erfüllen Sinus und Kosinus die Differentialgleichung $f'' + f = 0$, die die Schwingung eines reibungsfreien Federpendels (oder eines widerstandsfreien elektrischen Schwingkreises etc.) beschreibt, sofern Masse und Federhärte gleich 1 sind.

KOROLLAR 5.59 (Restgliedabschätzungen). Es gilt für jedes $x \in \mathbb{R}$ und $n \in \mathbb{N}$

$$\begin{aligned} \left| \cos x - \sum_{k=0}^n (-1)^k \frac{x^{2k}}{(2k)!} \right| &\leq \frac{|x|^{2n+2}}{(2n+2)!}, \\ \left| \sin x - \sum_{k=0}^n (-1)^k \frac{x^{2k+1}}{(2k+1)!} \right| &\leq \frac{|x|^{2n+3}}{(2n+3)!}. \end{aligned}$$

⁷Durch diese Reihenentwicklungen kann man $\sin$ und $\cos$ sogar als Funktionen $\mathbb{C} \rightarrow \mathbb{C}$ definieren.

BEWEIS. Aus der Restgliedabschätzung für die TAYLOR-Reihe in LAGRANGE-Darstellung (Korollar 4.49), hier für $m = 2n + 1$, erhalten wir für ein $\xi \in [0, x]$:

$$\left| \cos x - \sum_{k=0}^n (-1)^k \frac{x^{2k}}{(2k)!} \right| = |\cos(\xi)| \frac{|x|^{2n+2}}{(2n+2)!},$$

und die Behauptung folgt mit der Beobachtung $|\cos| \leq 1$. Die Abschätzung für $\sin$ geht analog. $\square$

5.5.3.2. Die Zahl π .

SATZ 5.60. Der Kosinus hat im Intervall $[0, 2]$ genau eine Nullstelle.

BEWEIS. Aus der Reihenentwicklung sieht man $\cos 0 = 1$. Andererseits gilt gemäß der eben gezeigten Restgliedabschätzung

$$\left| \cos 2 - 1 + \frac{2^2}{2} \right| \leq \frac{2^4}{4!},$$

mithin $|\cos 2 + 1| \leq \frac{2}{3}$, und somit $\cos 2 \leq \frac{2}{3} - 1 < 0$. Da $\cos$ stetig ist, folgt aus dem Zwischenwertsatz die Existenz einer Nullstelle.

Um zu zeigen, daß es nur eine Nullstelle gibt, genügt es zu zeigen, daß $\cos$ im fraglichen Intervall streng monoton fällt. Dazu zeigen wir $\cos' = -\sin < 0$, also $\sin > 0$ auf $(0, 2]$.

Für jedes $x \in (0, 2]$ haben wir, wiederum nach Restgliedabschätzung für die Sinusreihe,

$$|\sin x - x| \leq \frac{x^3}{6},$$

also

$$\sin x \geq x - \frac{x^3}{6} = x \left(1 - \frac{x^2}{6} \right).$$

Da aber in $(0, 2]$ gilt $\frac{x^2}{6} \leq \frac{2}{3} < 1$, folgt $\sin x > 0$ wie gewünscht. $\square$

DEFINITION 5.61 (Die Zahl π). Die eindeutig bestimmte Nullstelle des Kosinus im Intervall $(0, 2)$ heißt $\frac{\pi}{2}$.

Daß dieses π wirklich ‚das‘ π ist – nämlich der Flächeninhalt der Einheitskreisscheibe –, folgt aus den im folgenden (und in den Übungen) bewiesenen Eigenschaften der Winkelfunktionen zusammen mit Beispiel 4.40.

5.5.3.3. Weitere Eigenschaften der Winkelfunktionen.

PROPOSITION 5.62 (Additionstheoreme). Für alle $x, y \in \mathbb{R}$ gilt

$$\sin(x + y) = \sin x \cos y + \cos x \sin y, \quad \cos(x + y) = \cos x \cos y - \sin x \sin y.$$

BEWEIS. Dies folgt einfach aus der Funktionalgleichung der Exponentialfunktion und einem Vergleich von Real- und Imaginärteil, denn

$$\begin{aligned} \cos(x + y) + i \sin(x + y) &= e^{i(x+y)} = e^{ix} e^{iy} \\ &= (\cos x + i \sin x)(\cos y + i \sin y) \\ &= (\cos x \cos y - \sin x \sin y) + i(\sin x \cos y + \cos x \sin y). \end{aligned}$$

$\square$

PROPOSITION 5.63. $e^{i\frac{\pi}{2}} = i$, $e^{i\pi} = -1$, $e^{i\frac{3\pi}{4}} = -i$, $e^{2\pi i} = 1$.

BEWEIS. Einerseits ist definitionsgemäß $\cos(\pi/2) = 0$, andererseits ist (nach Beweis von Satz 5.60) $\sin(\pi/2) > 0$, und schließlich gilt ja $\sin^2 = 1 - \cos^2$, woraus $\sin(\pi/2) = 1$ folgt. Daraus ergibt sich

$$e^{i\frac{\pi}{2}} = \cos\left(\frac{\pi}{2}\right) + i \sin\left(\frac{\pi}{2}\right) = i$$

wie behauptet.

Die übrigen Aussagen folgen aus der Funktionalgleichung der Exponentialfunktion und aus $i^2 = -1$, $i^3 = -i$, $i^4 = 1$. So ist etwa

$$e^{i\frac{3\pi}{4}} = \left(e^{i\frac{\pi}{2}}\right)^3 = i^3 = -i.$$

□

KOROLLAR 5.64 (Periodizität und Symmetrien). Für alle $x \in \mathbb{R}$ gilt

$$\begin{aligned}\cos(x + 2\pi) &= \cos x, & \sin(x + 2\pi) &= \sin x, \\ \cos(x + \pi) &= -\cos x, & \sin(x + \pi) &= -\sin x, \\ \cos x &= \sin\left(\frac{\pi}{2} - x\right), & \sin x &= \cos\left(\frac{\pi}{2} - x\right).\end{aligned}$$

BEWEIS. Dies folgt sofort aus den Additionstheoremen und den speziellen Werten von $\sin$ und $\cos$, die sich aus Proposition 5.63 und der Formel $\exp(ix) = \cos x + i \sin x$ ergeben.

□

KOROLLAR 5.65 (Nullstellen).

- i) $\{x \in \mathbb{R} : \cos x = 0\} = \{\frac{\pi}{2} + k\pi : k \in \mathbb{Z}\}$ und $\{x \in \mathbb{R} : \sin x = 0\} = \{k\pi : k \in \mathbb{Z}\}$.
- ii) $\{x \in \mathbb{R} : e^{ix} = 1\} = \{2\pi k : k \in \mathbb{Z}\}$.

BEWEIS. Übung.

□

DEFINITION 5.66. Die Tangensfunktion $\tan : \mathbb{R} \setminus \{\frac{\pi}{2} + k\pi : k \in \mathbb{Z}\} \rightarrow \mathbb{R}$ ist definiert durch

$$\tan x = \frac{\sin x}{\cos x}.$$

5.5.3.4. *Geometrische Interpretation der komplexen Exponentialfunktion.* Wie (zum Schluß des letzten Unterabschnitts) angekündigt wollen wir nun noch die Intuition plausibilisieren, daß e^{ix} derjenige Punkt auf dem komplexen Einheitskreis ist, der mit der reellen Achse den Winkel x einschließt. Wir messen hier wie üblich Winkel im *Bogenmaß*, d.h. ein Winkel wird mit der Länge des entsprechenden Einheitskreisbogens identifiziert. Den vulgären Winkel in Grad erhält man dann aus der Beziehung

$$\alpha := \frac{180}{\pi} x$$

zurück, wobei α in Grad und x im Bogenmaß gemessen wird.

Sei nun $x \in [0, 2\pi)$ gegeben und betrachte die $n+1$ Punkte auf dem Einheitskreis, die durch

$$z_{n,k} := e^{ix\frac{k}{n}}, \quad k = 0, \dots, n$$

definiert sind. Aus dem Monotonieverhalten von Sinus und Kosinus (Übung) wird klar, daß für festes n die Punkte $z_{n,k}$ mit aufsteigendem k auch aufsteigende Winkel einschließen.

Betrachte die Länge des Polygonzugs, der die $n+1$ Punkte $z_{n,k}$ verbindet; sie ist gegeben durch

$$L_n = \sum_{k=0}^{n-1} |z_{n,k+1} - z_{n,k}|.$$

Wir wollen zeigen $\lim_{n \rightarrow \infty} L_n = x$, d.h. die Länge des Kreisbogens von $e^0 = 1$ bis e^{ix} beträgt tatsächlich x .

Dazu rechnen wir

$$\begin{aligned} L_n &= \sum_{k=0}^{n-1} \left| e^{ix \frac{k+1}{n}} - e^{ix \frac{k}{n}} \right| = \sum_{k=0}^{n-1} \left| e^{ix \frac{k+\frac{1}{2}}{n}} \right| \left| e^{ix \frac{1}{2n}} - e^{-ix \frac{1}{2n}} \right| \\ &= \sum_{k=0}^{n-1} \left| e^{ix \frac{1}{2n}} - e^{-ix \frac{1}{2n}} \right| = n \left| e^{ix \frac{1}{2n}} - e^{-ix \frac{1}{2n}} \right| = 2n \sin \left(\frac{x}{2n} \right), \end{aligned}$$

wobei wir die Funktionalgleichung der Exponentialfunktion und späterhin die Identität $\sin x = \frac{e^{ix} - e^{-ix}}{2i}$ verwendet haben, die leicht aus $e^{ix} = \cos x + i \sin x$ und den Symmetrien von $\sin$ und $\cos$ folgt.

Nach Restgliedabschätzung für den Sinus haben wir nun⁸

$$\left| 2n \sin \left(\frac{x}{2n} \right) - x \right| \leq 2n \frac{|x|^3}{6(2n)^3},$$

und dies konvergiert mit $n \rightarrow \infty$ gegen null. Es folgt, wie behauptet, $\lim_{n \rightarrow \infty} L_n = 0$.

Zum Abschluß bemerken wir, daß aus dieser Interpretation von e^{ix} auch die übliche geometrische Interpretation von Sinus und Kosinus folgt: betrachte das rechtwinklige Dreieck mit Hypotenuse der Länge 1, die mit der horizontalen (reellen) Achse den Winkel x einschließt. Da Kosinus und Sinus die Projektionen dieser Hypotenuse auf die horizontale bzw. vertikale Achse sind (nichts anderes besagt ja die Definition $\cos x = \Re e^{ix}$, $\sin x = \Im e^{ix}$), beschreiben sie genau die Länge der Ankathete bzw. Gegenkathete.

⁸Ebenso kann man den wichtigen Grenzwert $\lim_{x \rightarrow 0} \frac{\sin x}{x} = 1$ herleiten.

Topologische Grundlagen

Wie mißt man den Abstand zweier Objekte? Sind diese Objekte reelle oder komplexe Zahlen, so tut man dies üblicherweise mit dem Betrag, d.h. der Abstand zweier Zahlen x, y beträgt $|x - y|$. In vielen Situationen gibt es allerdings mehrere Möglichkeiten, Abstände oder Entfernungen zu messen: So sind von Ulm aus gemessen Augsburg und Nördlingen Luftlinie ungefähr gleich weit entfernt, aber eine Bahnfahrt nach Nördlingen dauert mehr als doppelt so lange wie nach Augsburg. Zu zwei gegebenen Bahnhöfen kann also eine Distanz angegeben werden (nämlich die Fahrzeit mit dem Zug), die sich von der geographischen Entfernung, wie sie idealisiert durch den Betrag des Differenzvektors in $\mathbb{R}^2$ gemessen würde, unterscheidet, und darüber hinaus für Bahnfahrende auch deutlich nützlicher ist. Den allgemeinsten Rahmen, innerhalb dessen in der Mathematik Distanzen zwischen zwei Elementen einer Menge angegeben werden können, bilden die *metrischen Räume*.

Wir betrachten außerdem zwei weitere Konzepte: Die *topologischen Räume*, die allgemeiner sind als die metrischen, werden kurz angeschnitten, da man sie in einem Mathematikstudium unbedingt gesehen haben sollte, in Ulm aber leider keine regelmäßigen Lehrveranstaltungen zur Topologie angeboten werden.¹ Topologische Räume bieten eine Möglichkeit zu definieren, wann zwei Punkte einander ‘nahe’ sind, ohne einen quantifizierbaren Abstand angeben zu müssen.

Normierte Räume sind spezielle metrische Räume, die mit der Struktur eines Vektorraums verträglich sind. Ein wichtiges Beispiel für einen normierten Raum bildet $\mathbb{R}^n$.

6.1. Topologische Räume

Wir haben bereits von offenen, abgeschlossenen und kompakten Intervallen gesprochen. Die Definitionen dieser Konzepte in abstrakten topologischen Räumen werden Ihnen zunächst sehr fremd erscheinen², aber es wird sich herausstellen, daß sie im Falle der reellen Zahlen mit den bereits bekannten Begriffen übereinstimmen.

6.1.1. Grundbegriffe. Man erinnere sich an die Potenzmenge $\mathcal{P}(X)$ einer Menge X , also die Menge aller Teilmengen von X .

DEFINITION 6.1 (Topologischer Raum). Ein *topologischer Raum* ist eine nichtleere Menge X zusammen mit einer *Topologie* $\mathcal{T} \subset \mathcal{P}(X)$, die die folgenden Eigenschaften hat:

- i) $\emptyset \in \mathcal{T}$ und $X \in \mathcal{T}$,
- ii) die Vereinigung von Elementen aus $\mathcal{T}$ ist wieder in $\mathcal{T}$,
- iii) der *endliche* Durchschnitt von Elementen aus $\mathcal{T}$ ist wieder in $\mathcal{T}$.

Die Elemente der Topologie $\mathcal{T}$ heißen *offene Mengen*, die Komplemente offener Mengen heißen *abgeschlossen*.

Um an dieser Definition nicht irr zu werden, betrachten wir ein bereits bekanntes Beispiel:

¹Ein unterhaltsame Einführung in die Topologie bietet JÄNICH [12].

²In der Maßtheorie werden uns ähnliche Strukturen, nämlich σ -Algebren, begegnen. Es ist daher gut, wenn Sie sich schonmal an solche Konzepte gewöhnen.

BEISPIEL 6.2. Sei $X = \mathbb{R}$. Eine Teilmenge $V \subset \mathbb{R}$ heißt nach Definition 3.9 offen, wenn zu jedem $x \in V$ ein $\epsilon > 0$ existiert, sodaß das Intervall $(x - \epsilon, x + \epsilon) \subset V$. Wir zeigen, daß die so definierten offenen Teilmengen von $\mathbb{R}$ eine Topologie im Sinne von Definition 6.1 bilden:

i) Die leere Menge ist offen, da es kein Element von $\emptyset$ gibt, für das ein ϵ angegeben werden müßte. Die ganze Menge $\mathbb{R}$ ist ebenfalls offen, da für jedes $x \in \mathbb{R}$ und jedes $\epsilon > 0$ das Intervall $(x - \epsilon, x + \epsilon) \subset \mathbb{R}$ ist.

ii) Sei I irgendeine Indexmenge und $(V_i)_{i \in I}$ eine Familie offener Mengen. Wir wollen zeigen, daß $\mathcal{V} := \bigcup_{i \in I} V_i$ wieder offen ist. Sei dazu $x \in \mathcal{V}$, dann existiert ein $j \in I$ mit $x \in V_j$. Da V_j offen ist, existiert $\epsilon > 0$ mit $(x - \epsilon, x + \epsilon) \subset V_j$. Da aber $V_j \subset \mathcal{V}$, gilt auch $(x - \epsilon, x + \epsilon) \subset \mathcal{V}$.

iii) Sei $\mathcal{W} := \bigcap_{i=1}^N W_i$ ein endlicher Durchschnitt offener Mengen, und sei $x \in \mathcal{W}$. Dann ist $x \in W_i$ für alle $i = 1, \dots, N$, und da jedes W_i offen ist, existiert zu jedem i ein $\epsilon_i > 0$, sodaß $(x - \epsilon_i, x + \epsilon_i) \subset W_i$. Da es sich nur um endlich viele ϵ_i handelt, gilt $\epsilon := \min_{i=1, \dots, N} \epsilon_i > 0$, und außerdem

$$(x - \epsilon, x + \epsilon) \subset (x - \epsilon_i, x + \epsilon_i) \subset W_i$$

für jedes i . Daher ist sogar $(x - \epsilon, x + \epsilon) \subset \mathcal{W}$.

Dieses Beispiel zeigt auch, daß *beliebige* (unendliche) Durchschnitte offener Mengen nicht offen zu sein brauchen (vgl. Klausur zur Analysis I, Aufgabe 1f): Denn es ist $\{1\} = \bigcap_{n=1}^{\infty} (1 - \frac{1}{n}, 1 + \frac{1}{n})$ ein Schnitt offener Mengen, der selbst nicht offen ist.

Das Beispiel kann wörtlich auf $\mathbb{C}$ übertragen werden, sofern man $(x - \epsilon, x + \epsilon)$ durch $\{z \in \mathbb{C} : |x - z| < \epsilon\}$ ersetzt.

Ein verbreitetes Mißverständnis soll hier gleich zu Beginn ausgeräumt werden: Es gibt einerseits Mengen, die weder offen noch abgeschlossen sind, z.B. das Intervall $[0, 1)$. Denn dieses ist nicht offen, da es zu 0 keine in $[0, 1)$ enthaltene ϵ -Umgebung gibt; es ist aber auch nicht abgeschlossen, da sein Komplement

$$\mathbb{R} \setminus [0, 1) = (-\infty, 0) \cup [1, +\infty)$$

ebenfalls nicht offen ist.

Andererseits gibt es Mengen, die sowohl offen als auch abgeschlossen sind: In $\mathbb{R}$ nämlich die leere Menge und die ganze Menge $\mathbb{R}$ selbst. Allgemein sind $\emptyset$ und X stets offen und abgeschlossen³. Gibt es darüber hinaus weitere Teilmengen von $\mathbb{R}$ mit dieser Eigenschaft? Die Antwort lautet nein:

PROPOSITION 6.3. *Sei $\mathbb{R}$ mit der in Beispiel 6.2 angegebenen Topologie. Ist $V \subset \mathbb{R}$ offen und abgeschlossen, so ist $V = \mathbb{R}$ oder $V = \emptyset$.*

BEWEIS. Sei V offen und abgeschlossen, und wir nehmen an, daß $V \neq \emptyset$ und $V \neq \mathbb{R}$. Dann ist das Komplement V^c ebenfalls offen und abgeschlossen, und $V^c \neq \emptyset$ und $V^c \neq \mathbb{R}$. Wir zeigen zunächst, daß in diesem Falle ein $x \in \mathbb{R}$ existiert, sodaß für jedes $\epsilon > 0$ gilt: $(x - \epsilon, x + \epsilon) \cap V \neq \emptyset$ und $(x - \epsilon, x + \epsilon) \cap V^c \neq \emptyset$ (x ist also ein *Randpunkt* von V).

Seien dazu $y \in V$ und $z \in V^c$ beliebig gewählt; o.B.d.A. dürfen wir $y < z$ annehmen. Dann setze

$$x := \sup\{w \in [y, z] : w \in V\}.$$

Da es eine Zahl in $[y, z]$ gibt, die auch in V ist (nämlich y), wird das Supremum über eine nichtleere Menge genommen und existiert daher als reelle Zahl. Nach Definition des Supremums gilt außerdem für jedes $\epsilon > 0$, daß $(x - \epsilon, x + \epsilon) \cap V \neq \emptyset$, und es gilt ebenfalls $(x - \epsilon, x + \epsilon) \cap V^c \neq \emptyset$; denn falls $x = z$, so ist $x \in V^c$, und falls $x < z$, so ist nach Definition jede Zahl in $[y, z]$, die größer ist als x , in V^c .

³Im Englischen heißen Mengen, die offen und abgeschlossen (open and closed) sind, *clopen sets*. Eine deutsche Entsprechung wäre etwa *abgeschloffen*.

Es ist entweder $x \in V$ oder $x \in V^c$. Wir nehmen o.B.d.A. ersteres an. Da V offen ist, existiert $\delta > 0$, sodaß $(x - \delta, x + \delta) \subset V$. Andererseits haben wir soeben gezeigt $(x - \delta, x + \delta) \cap V^c \neq \emptyset$, also folgt $V \cap V^c \neq \emptyset$, und das ist der gewünschte Widerspruch. $\square$

Nicht jeder topologische Raum hat die Eigenschaft, daß es nur triviale offene und abgeschlossene Mengen gibt:

BEISPIEL 6.4. Sei $X = [0, 1] \cup [2, 3]$, und eine Teilmenge $V \subset X$ heiße offen, wenn es für jedes $x \in V$ ein $\epsilon > 0$ gibt, sodaß $(x - \epsilon, x + \epsilon) \cap X \subset V$ (vgl. Definition 3.9). Es ist leicht nachzuprüfen, daß diese offenen Mengen eine Topologie bilden⁴ (Übung).

Wir behaupten, daß die Teilmenge $[0, 1] \subset X$ offen und abgeschlossen ist: Sie ist offen, denn für jedes $x \in [0, 1]$ und $\epsilon = \frac{1}{2}$ gilt $(x - \epsilon, x + \epsilon) \cap X \subset [0, 1]$. Sie ist aber auch abgeschlossen, denn das Komplement $[2, 3]$ ist offen, wie man wieder für jedes $x \in [2, 3]$ mit der Wahl $\epsilon = \frac{1}{2}$ sieht.

DEFINITION 6.5. Ein topologischer Raum X heißt *zusammenhängend*, wenn gilt: Ist $V \subset X$ offen und abgeschlossen, so ist $V = \emptyset$ oder $V = X$.

In topologischen Räumen kann man Konvergenz und Stetigkeit definieren. Zunächst definieren wir den Begriff der *Umgebung*:

DEFINITION 6.6 (Umgebung). Sei X ein topologischer Raum und $x \in X$. Dann heißt $U \subset X$ eine *Umgebung* von x , wenn es eine offene Menge $V \subset X$ gibt mit $x \in V \subset U$.

Im Falle $X = \mathbb{R}$ mit der üblichen Topologie ist etwa $B_\epsilon(x) := (x - \epsilon, x + \epsilon)$ eine Umgebung von x , sofern $\epsilon > 0$.

DEFINITION 6.7 (Konvergenz). Sei X ein topologischer Raum und $(x_n)_{n \in \mathbb{N}} \subset X$ eine Folge. Man sagt, die Folge konvergiere gegen ein $x \in X$, wenn jede Umgebung von x alle bis auf endlich viele Glieder der Folge enthält.

Man überzeuge sich, daß in $\mathbb{R}$ diese Definition der Konvergenz mit der gewohnten übereinstimmt (Übung).

DEFINITION 6.8 (Stetigkeit). Seien X, Y topologische Räume und $f : X \rightarrow Y$ eine Abbildung.

a) f heißt *stetig* in $x \in X$, wenn das Urbild jeder offenen Umgebung von $f(x) \in Y$ unter f wieder offen ist.

b) f heißt *folgenstetig* in $x \in X$, falls für jede gegen x konvergente Folge $(x_n) \subset X$ die Folge $(f(x_n))_{n \in \mathbb{N}}$ gegen $f(x)$ konvergiert.

In Satz 3.10 haben wir bereits gesehen, daß in $\mathbb{R}$ Stetigkeit und Folgenstetigkeit äquivalent (und auch äquivalent zur ϵ - δ -Definition) sind. Man kann zeigen, daß dies in allgemeinen topologischen Räumen nicht mehr der Fall zu sein braucht.

6.1.2. Beispiele.

- (1) Sei $X \subset \mathbb{R}^n$ (oder $\mathbb{C}^n$), dann ist eine Topologie auf X folgendermaßen definiert: $V \subset X$ heißt offen, wenn es zu jedem $x \in V$ ein $\epsilon > 0$ gibt, sodaß

$$B_\epsilon(x) \cap X = \{y \in X : |x - y| < \epsilon\} \subset V.$$

Hier bezeichnet $|\cdot|$ die euklidische Norm eines Vektors, die Sie aus der Linearen Algebra kennen: Für $x \in \mathbb{R}^n$ ist $|x| := (\sum_{k=1}^n x_k^2)^{1/2}$.

So ist zum Beispiel $(0, 1) \times (3, 5) \subset \mathbb{R}^2$ offen in $\mathbb{R}^2$. Die Menge $[0, 1) \times [3, 5]$ ist offen im topologischen Raum $Y := [0, \infty) \times [3, 5]$, aber nicht in $Z := [0, \infty) \times [3, 6]$.

⁴Diese Topologie nennt man die von $\mathbb{R}$ auf die Teilmenge X *induzierte Topologie* oder *Relativtopologie*.

Ähnlich wie in $\mathbb{R}$ kann man zeigen (Übung), daß $\mathbb{R}^2$ mit dieser Topologie zusammenhängend ist. Dagegen ist $\mathbb{R}^2 \setminus \{(x, y) \in \mathbb{R}^2 : y = 0\}$ nicht zusammenhängend, da die Mengen $\{(x, y) \in \mathbb{R}^2 : y < 0\}$ und $\{(x, y) \in \mathbb{R}^2 : y > 0\}$ offen und abgeschlossen sind.

Eine Folge $(x_n)_{n \in \mathbb{N}} \subset \mathbb{R}^n$ konvergiert genau dann gegen $x \in \mathbb{R}^n$, wenn für jedes $\epsilon > 0$ ein $N \in \mathbb{N}$ existiert, sodaß $|x - x_n| < \epsilon$ für alle $n \geq N$ (wobei $|\cdot|$ wieder die euklidische Norm bezeichnet). Es ist nicht schwer zu sehen, daß dies äquivalent zur komponentenweisen Konvergenz ist, daß also $x_n^j \rightarrow x^j$ für jedes $j = 1, \dots, n$ im Sinne der Konvergenz in $\mathbb{R}$.

- (2) Sei X eine nichtleere Menge. Die *diskrete Topologie* ist einfach $\mathcal{P}(X)$, d.h. jede Teilmenge von X ist offen. Man prüft leicht nach, daß $\mathcal{P}(X)$ tatsächlich eine Topologie ist.

Sei $(x_n)_{n \in \mathbb{N}}$ eine Folge. Damit sie gegen $x \in X$ konvergiere, muß jede Umgebung von x fast alle x_n enthalten. Allerdings ist $\{x\}$ eine Umgebung von x . Das bedeutet: Eine Folge konvergiert genau dann in der diskreten Topologie gegen x , wenn fast alle Folgenglieder konstant gleich x sind.

Sei Y ein weiterer (beliebiger) topologischer Raum. Dann ist jede Abbildung $f: X \rightarrow Y$ stetig, denn jedes Urbild unter f ist offen.

Jede Menge X mit mindestens zwei Elementen ist in der diskreten Topologie nicht zusammenhängend, denn jede einelementige Menge $\{x\}$ ist offen und abgeschlossen.

Man nennt die diskrete Topologie auch die *feinste* Topologie, mit der man eine Menge versehen kann.

- (3) Der Gegensatz zur diskreten Topologie auf einer nichtleeren Menge X ist die *größte Topologie*, die nur aus der leeren Menge und dem ganzen Raum X besteht. Auch hier prüft man wieder nach, daß es sich in der Tat um eine Topologie handelt.

Ist $x \in X$, so ist X die einzige Umgebung. Daher konvergiert *jede Folge gegen jeden Grenzwert*. Daran sieht man, daß in allgemeinen topologischen Räumen Grenzwerte nicht eindeutig zu sein brauchen.

Sei $f: X \rightarrow \mathbb{R}$ eine Abbildung, dann ist sie genau dann stetig, wenn f konstant ist. Denn einerseits ist das Urbild von $y \in \mathbb{R}$ unter einer konstanten Funktion entweder $\emptyset$ oder X ; ist andererseits f nicht konstant, so nimmt es mindestens zwei Werte $y_1 \neq y_2$ an, und damit ist für $\epsilon := \frac{|y_1 - y_2|}{2}$: $f^{-1}(B_\epsilon(y_1)) \notin \{\emptyset, X\}$, also ist f nicht stetig.

Allerdings ist jede Abbildung von einem topologischen Raum Y nach X stetig, denn die Urbilder der (einzigen) offenen Mengen $\emptyset$ und X sind $\emptyset$ bzw. Y , und die sind stets offen.

Offensichtlich ist X mit der groben Topologie zusammenhängend.

- (4) ⁵ Wir wählen jetzt $X = \mathbb{Z}$ und definieren für $a, b \in \mathbb{Z}$, $b > 0$ Teilmengen

$$N_{a,b} := \{a + nb : n \in \mathbb{Z}\}.$$

(Man nennt eine solche Menge eine *arithmetische Folge*.) Eine Teilmenge $V \subset \mathbb{Z}$ heiße offen, wenn entweder $V = \emptyset$ oder zu jedem $a \in V$ ein $b \in \mathbb{N}$ existiert, sodaß $N_{a,b} \subset V$. Wir zeigen, daß dadurch eine Topologie auf $\mathbb{Z}$ definiert wird:

Nach Definition ist $\emptyset$ offen. Die Offenheit von $\mathbb{Z}$ folgt ebenfalls sofort aus der Definition (für jedes a kann dafür b beliebig gewählt werden).

Sei $a \in \bigcup_{i \in I} V_i$, wobei jedes V_i offen ist. dann gibt es ein $j \in I$ mit $a \in V_j$, und daher existiert $b \in \mathbb{N}$, sodaß $a \in N_{a,b} \subset V_j \subset \bigcup_{i \in I} V_i$. Also ist die Vereinigung offener Mengen wieder offen.

⁵Dieses bezaubernde Beispiel ist dem auch sonst sehr empfehlenswerten Buch [1] entnommen.

Seien schließlich V_1, V_2 offen und $a \in V_1 \cap V_2$. Dann gibt es nach Definition $b_1, b_2 \in \mathbb{N}$, sodaß $N_{a,b_1} \subset V_1$ und $N_{a,b_2} \subset V_2$. Da $N_{a,b_1 b_2} \subset N_{a,b_1}$ und ebenso $N_{a,b_1 b_2} \subset N_{a,b_2}$, ist $N_{a,b_1 b_2} \subset V_1 \cap V_2$, und $V_1 \cap V_2$ ist somit offen. Wenn aber der Durchschnitt zweier offener Mengen offen ist, so auch der Durchschnitt endlich vieler offener Mengen (denn dann ist ja auch $(V_1 \cap V_2) \cap V_3$ offen usw.). Damit ist gezeigt, daß die so definierten offenen Mengen eine Topologie bilden.

Da $N_{a,b} = \mathbb{Z} \setminus \bigcup_{i=1}^{b-1} N_{a+i,b}$, ist $N_{a,b}$ als Komplement einer offenen Menge abgeschlossen. Insbesondere ist $\mathbb{Z}$ mit dieser Topologie nicht zusammenhängend.

Als Anwendung zeigen wir, daß die Menge $\mathbb{P}$ der Primzahlen unendlich ist. (Eine natürliche Zahl heißt bekanntlich prim, wenn sie genau zwei Teiler hat; so sind etwa 2, 3, 5, 7, 11, etc. prim, da z.B. 2 genau die Teiler 1 und 2 hat.) Da jede ganze Zahl außer ± 1 mindestens einen Primteiler hat (warum?), gibt es zu jedem $z \in \mathbb{Z} \setminus \{\pm 1\}$ ein $p \in \mathbb{P}$ mit $z \in N_{0,p}$. Daher gilt

$$\mathbb{Z} \setminus \{-1, +1\} = \bigcup_{p \in \mathbb{P}} N_{0,p}. \quad (6.1)$$

Angenommen, $\mathbb{P}$ wäre endlich, so wäre $\bigcup_{p \in \mathbb{P}} N_{0,p}$ eine endliche Vereinigung abgeschlossener Mengen (die Abgeschlossenheit von $N_{a,b}$ haben wir ja eben gezeigt). Nach Maßgabe der mengentheoretischen Identität

$$\left(\bigcap_{i \in I} M_i \right)^c = \bigcup_{i \in I} M_i^c$$

ist aber die endliche Vereinigung abgeschlossener Mengen wieder abgeschlossen, und nach (6.1) wäre also $\{-1, 1\}$ als Komplement einer abgeschlossenen Menge offen. Unmittelbar aus der Definition der offenen Menge ist aber ersichtlich, daß jede nichtleere offene Menge unendlich ist, und dies ergibt den gewünschten Widerspruch zur Endlichkeit von $\mathbb{P}$.

Wir bemerken noch, daß es selbstverständlich einen vielen elementareren Beweis der Unendlichkeit der Primzahlen gibt, nämlich den von EUKLID: Angenommen, es gäbe nur endlich viele Primzahlen $p_1, \dots, p_n$. Betrachte die Zahl

$$P := \prod_{i=1}^n p_i + 1.$$

Dann ist P durch kein p_i teilbar, es enthält also einen von allen p_i verschiedenen Primfaktor, und wir erhalten so einen Widerspruch.

6.1.3. Abschluß und Rand. Hier nur ein paar kurze Definitionen:

DEFINITION 6.9 (Topologischer Abschluß). Sei X ein topologischer Raum und $\Omega \subset X$ eine beliebige Teilmenge. Dann heißt der Durchschnitt aller abgeschlossenen Mengen, die Ω enthalten, der (*topologische*) *Abschluß* von Ω . Man schreibt dafür $\bar{\Omega}$.

Man beachte, daß beliebige Durchschnitte abgeschlossener Mengen wieder abgeschlossen sind, denn das Komplement des Durchchnitts ist die Vereinigung der Komplemente, und da letztere offen sind und die Vereinigung offener Mengen wieder offen ist, ist der Durchschnitt abgeschlossener Mengen abgeschlossen. Der Abschluß einer beliebigen Menge Ω ist also abgeschlossen und enthält Ω . Ist Ω bereits abgeschlossen, so ist offenbar $\bar{\Omega} = \Omega$.

DEFINITION 6.10 (offener Kern). Sei X ein topologischer Raum und $\Omega \subset X$ eine beliebige Teilmenge. Dann heißt

$$\Omega^\circ := (\bar{\Omega^c})^c$$

der *offene Kern* von Ω .

Der offene Kern ist (als Komplement einer abgeschlossenen Menge) offen und in Ω enthalten. Ist Ω offen, so ist $\Omega = \Omega^\circ$.

DEFINITION 6.11 (Rand). Sei X ein topologischer Raum und $\Omega \subset X$ eine beliebige Teilmenge. Dann heißt

$$\partial\Omega := \bar{\Omega} \setminus \Omega^\circ$$

der (*topologische*) Rand von Ω .

Der Rand einer Menge ist stets abgeschlossen (warum?).

BEISPIEL 6.12. (1) Sei $X = \mathbb{R}$ mit der üblichen Topologie (s. Beispiel 6.2) und $\Omega = [0, 1)$. Dann ist $\bar{\Omega} = [0, 1]$, $\Omega^\circ = (0, 1)$, und $\partial\Omega = \{0, 1\}$.

(2) Sei $X = \mathbb{C}$ mit der Topologie aus Beispiel 6.1.2 (1), und $\Omega = B_1(0) = \{z \in \mathbb{C} : |z| < 1\}$. Dann ist $\bar{\Omega} = \{z \in \mathbb{C} : |z| \leq 1\}$, $\Omega^\circ = \Omega$, und $\partial\Omega = \{z \in \mathbb{C} : |z| = 1\}$. Für $\Omega' = \{z \in \mathbb{C} : |z| = 1\}$ gilt $\overline{\Omega'} = \Omega'$, $(\Omega')^\circ = \emptyset$ und $\partial\Omega' = \Omega'$.

6.1.4. Kompaktheit.

DEFINITION 6.13 (Kompaktheit). Ein topologischer Raum X heißt *kompakt*, wenn gilt: Ist $X = \bigcup_{i \in I} V_i$ eine offene Überdeckung von X (d.h. jedes V_i ist offen), so gibt es eine endliche Teilüberdeckung $X = \bigcup_{i \in I_0} V_i$ (d.h. I_0 ist endlich).

Diese zugegebenmaßen recht unhandlich anmutende Definition wollen wir uns an einem bekannten Beispiel klarmachen:

BEISPIEL 6.14. (1) Sei $X = [0, 1] \subset \mathbb{R}$ mit der üblichen Topologie (s. Beispiel 6.2). Dann ist X kompakt. Sei dazu $[0, 1] = \bigcup_{i \in I} V_i$ eine offene Überdeckung. Angenommen, $[0, 1]$ könnte nicht durch endlich viele V_i überdeckt werden. Unter dieser Annahme konstruieren wir nun eine Intervallschachtelung $(I_n)_{n \in \mathbb{N}}$, sodaß $|I_{n+1}| = \frac{1}{2}|I_n|$ für jedes $n \in \mathbb{N}$, und sodaß jedes I_n nicht durch endlich viele V_i überdeckt werden kann.

Wähle dazu $I_1 := [0, 1]$. Ist $I_n = [a_n, b_n]$ bereits konstruiert, so bilden wir die beiden Teilintervalle $I_n^- := [a_n, \frac{a_n+b_n}{2}]$ sowie $I_n^+ := [\frac{a_n+b_n}{2}, b_n]$. Da nach Induktionsannahme I_n nicht durch endlich viele V_i überdeckt wird, wird auch I_n^- oder I_n^+ nicht durch endlich viele V_i überdeckt (sonst könnte man die beiden endlichen Überdeckungen von I_n^- und I_n^+ einfach vereinigen, um eine endliche Überdeckung von I_n zu erhalten). Wähle von den beiden Intervallen I_n^- und I_n^+ eines aus, das nicht von endlich vielen V_i überdeckt wird, und setze dieses Intervall als I_{n+1} .

Nach Intervallschachtelungsprinzip (Satz 2.25) existiert nun ein $x \in [0, 1]$, sodaß $x \in I_n$ für alle $n \in \mathbb{N}$. Für dieses x existiert nach Voraussetzung ein $j \in I$ mit $x \in V_j$, und da V_j offen ist, gibt es ein $\epsilon > 0$ mit $(x - \epsilon, x + \epsilon) \cap [0, 1] \subset V_j$. Wähle $n \in \mathbb{N}$ so groß, daß $2^{-n} < \epsilon$. Da $|I_n| = 2^{-n} < \epsilon$, gilt $I_n \subset V_j$, also ist $I_n \subset \bigcup_{i \in \{j\}} V_i$ eine endliche Überdeckung von I_n , im Widerspruch zur Konstruktion von I_n . Daher ist unsere Annahme, $[0, 1]$ hätte keine Überdeckung aus endlich vielen V_i , falsch, und die Kompaktheit ist bewiesen.

Es ist klar, daß diese Argumentation ebenso für Intervalle $[a, b]$ mit $a < b$ funktioniert. Unsere frühere Bezeichnung solcher Intervalle als kompakt ist also mit der neuen Definition konsistent.

(2) Sei nun $X = (0, 1)$ mit der üblichen Topologie. Wir zeigen, daß X nicht kompakt ist. Dazu reicht es aus, eine offene Überdeckung von $(0, 1)$ anzugeben, die keine endliche Teilüberdeckung zuläßt. Sei dazu $V_0 := (\frac{1}{4}, 1)$ und $V_n := (2^{-n-2}, 2^{-n})$ für alle $n \in \mathbb{N}$. Da für jedes $x \in (0, 1)$ ein $n \in \mathbb{N} \cup \{0\}$ existiert, sodaß $x > 2^{-n-2}$ und $x < 2^{-n}$, ist $(0, 1) = \bigcup_{i=0}^{\infty} V_i$ eine offene Überdeckung.

Betrachte eine endliche Teilfamilie von Mengen $V_{i_1}, \dots, V_{i_N}$ mit $i_1 < i_2 < \dots < i_N$, so gilt für jedes $x \in (0, 2^{i_N-2}]$

$$x \notin \bigcup_{k=1}^N V_{i_k}.$$

Daher ist $\bigcup_{k=1}^N V_{i_k}$ keine Überdeckung von $(0, 1)$. Da die Auswahl der endlichen Teilfamilie beliebig war, zeigt dies, daß $(0, 1)$ nicht kompakt ist.

Ist X ein topologischer Raum und $\Omega \subset X$, so ist auf Ω durch die Mengen der Form $\{V \cap \Omega : V \subset X \text{ offen}\}$ eine Topologie definiert, die als die *von X induzierte Topologie* oder *Relativtopologie auf Ω bzgl. X* bezeichnet wird (vgl. Übung 2 Aufgabe 1 iii).

SATZ 6.15 (vgl. Satz 3.16). Seien X, Y topologische Räume und $f : X \rightarrow Y$ stetig. Ist X kompakt, so auch $f(X)$ bezüglich der von Y induzierten Topologie.

BEWEIS. Sei $f(X) \subset \bigcup_{i \in I} V_i \subset Y$ eine offene Überdeckung (d.h. $V_i \subset Y$ sind offen, daher sind auch $V_i \cap f(X)$ in $f(X)$ offen). Da f stetig ist, ist auch jedes $f^{-1}(V_i) \subset X$ offen, und wegen $f^{-1}(\bigcup V_i) = \bigcup f^{-1}(V_i)$ folgt

$$X = \bigcup_{i \in I} f^{-1}(V_i).$$

Da X kompakt ist, gibt es eine endliche Teilüberdeckung $X = \bigcup_{k=1}^N f^{-1}(V_{i_k})$. Ist nun $y \in f(X)$, so existiert $x \in X$ mit $f(x) = y$ und daher auch $k \in \{1, \dots, N\}$ mit $x \in f^{-1}(V_{i_k})$, also $y = f(x) \in V_{i_k}$. Damit ist gezeigt, daß

$$f(X) \subset \bigcup_{k=1}^N V_{i_k},$$

und dies ist die (bzw. eine) gesuchte endliche Teilüberdeckung von $f(X)$. $\square$

DEFINITION 6.16 (Folgenkompaktheit, vgl. Satz von BOLZANO-WEIERSTRASS). Ein topologischer Raum X heißt *folgenkompakt*, wenn jede Folge in X eine konvergente Teilfolge besitzt.

In allgemeinen topologischen Räumen sind Überdeckungskompaktheit und Folgenkompaktheit nicht äquivalent; wir studieren nun aber eine spezielle Klasse topologischer Räume, die *metrischen Räume*, in denen die beiden Konzepte, ebenso wie Stetigkeit und Folgenstetigkeit, gleichbedeutend sind.

6.2. Metrische Räume

6.2.1. Definition und Beispiele.

DEFINITION 6.17 (Metrischer Raum). Ein *metrischer Raum* ist eine nichtleere Menge X zusammen mit einer Abbildung (der *Metrik*) $d : X \times X \rightarrow \mathbb{R}$, sodaß gilt:

- i) *Positivität*: $d(x, y) \geq 0$ für alle $x, y \in X$, mit Gleichheit genau dann, wenn $x = y$;
- ii) *Symmetrie*: $d(x, y) = d(y, x)$ für alle $x, y \in X$;
- iii) *Dreiecksungleichung*: $d(x, z) \leq d(x, y) + d(y, z)$ für alle $x, y, z \in X$.

Man sollte sich $d(x, y)$ als den Abstand zwischen den Punkten x und y vorstellen.

BEISPIEL 6.18. (1) Die Standardmetrik auf $\mathbb{R}^n$ ist definiert als

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

wobei $|\cdot|$ wieder die euklidische Norm bezeichnet. Ebenso ist mit dieser Metrik jede Teilmenge von $\mathbb{R}^n$ ein metrischer Raum.

(2) Für jede nichtleere Menge X definiert

$$d(x, y) := \begin{cases} 1, & \text{falls } x \neq y, \\ 0, & \text{falls } x = y \end{cases}$$

eine Metrik.

- (3) Sei $P := \{x^1, x^2, \dots, x^N\}$ eine endliche Teilmenge des $\mathbb{R}^2$, und betrachte die Menge $\mathbb{R}^2/P := (\mathbb{R}^2 \setminus P) \cup \{P\}$

(d.h. man betrachtet $\mathbb{R}^2$ und identifiziert dabei die Punkte aus P). Dann ist durch

$$d(x, y) := \min \left\{ |x - y|; \min_{j=1, \dots, N} |x - x^j| + \min_{k=1, \dots, N} |y - x^k| \right\}$$

für $x, y \notin P$ und $d(x, P) := \min_{j=1, \dots, N} |x - x^j|$ die ‚U-Bahn-Metrik‘ auf $\mathbb{R}^2/P$ gegeben. Übung: Zeige, daß dies tatsächlich eine Metrik ist. Was hat sie mit der U-Bahn zu tun?

DEFINITION 6.19. Sei (X, d) ein metrischer Raum. Eine Menge $V \subset X$ heißt *offen*, wenn es zu jedem $x \in V$ ein $\epsilon > 0$ gibt, sodaß

$$B_\epsilon(x) := \{y \in X : d(x, y) < \epsilon\} \subset V.$$

Die offenen Mengen definieren eine Topologie auf X , die als die *von d induzierte Topologie* auf X bezeichnet wird.

Die letzte Aussage bedarf eines kurzen Beweises: $\emptyset$ und X sind nach dieser Definition gewiß offen; sind V_i offen für jedes $i \in I$ und ist $x \in \bigcup V_i$, so ist $x \in V_j$ für ein $j \in I$, und daher gibt es $\epsilon > 0$ sodaß $B_\epsilon(x) \subset V_j \subset \bigcup V_i$; und sind V_i offen für $i = 1, \dots, N$, so gibt es für jedes $i = 1, \dots, N$ ein $\epsilon_i > 0$ mit $B_{\epsilon_i}(x) \subset V_i$ – dann aber ist für $\epsilon := \min_{i=1, \dots, N} \epsilon_i > 0$ auch $B_\epsilon(x) \subset V_i$ für alle i , und somit auch $B_\epsilon(x) \subset \bigcap V_i$.

Damit sind alle in topologischen Räumen eingeführten Begriffe (Konvergenz, Stetigkeit, Kompaktheit, Rand, etc.) auch in metrischen Räumen definiert.

PROPOSITION 6.20. Sei (X, d) ein metrischer Raum. Eine Folge $(x_n)_{n \in \mathbb{N}} \subset X$ konvergiert genau dann gegen $x \in X$, wenn gilt:

$$\forall \epsilon > 0 \quad \exists N \in \mathbb{N} \quad \forall n \geq N : \quad d(x_n, x) < \epsilon. \quad (6.2)$$

Im Falle der Konvergenz ist der Grenzwert eindeutig bestimmt.

BEWEIS. Sei $(x_n)_{n \in \mathbb{N}}$ konvergent gegen x , d.h. jede Umgebung von x enthält fast alle Folgenglieder. Insbesondere ist jedes $B_\epsilon(x)$ Umgebung von x , und somit gibt es zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$, sodaß (6.2) erfüllt ist.

Gelte umgekehrt (6.2), und sei $U \supset x$ eine Umgebung von x . Das bedeutet, daß U eine offene Menge enthält, die ihrerseits x enthält. Daher existiert $\epsilon > 0$ mit $B_\epsilon(x) \subset U$, und nach (6.2) liegen fast alle Folgenglieder in $B_\epsilon(x) \subset U$.

Zur Eindeutigkeit des Grenzwerts: Angenommen $(x_n)_{n \in \mathbb{N}}$ konvergiere sowohl gegen x als auch gegen y . Dann existiert für jedes $\epsilon > 0$ ein $N \in \mathbb{N}$, sodaß $x_N \in B_\epsilon(x)$ und $x_N \in B_\epsilon(y)$. Nach Dreiecksungleichung ist aber $d(x, y) \leq d(x, x_N) + d(x_N, y) < 2\epsilon$, und da dies für jedes $\epsilon > 0$ der Fall ist, folgt $d(x, y) = 0$, und somit $x = y$. $\square$

6.2.2. Kompaktheit.

DEFINITION 6.21 (Totale Beschränktheit). Ein metrischer Raum (X, d) heißt *total beschränkt*, wenn es für jedes $\epsilon > 0$ eine Überdeckung von X aus endlich vielen Kugeln vom Radius ϵ gibt, d.h. $X = \bigcup_{i=1}^{N(\epsilon)} B_\epsilon(x_i)$.

PROPOSITION 6.22. Jeder folgenkompakte metrische Raum ist total beschränkt.

BEWEIS. Sei X folgenkompakt. Angenommen, er wäre nicht total beschränkt, so existierte $\epsilon > 0$ dergestalt, daß X nicht durch endlich viele Kugeln vom Radius ϵ überdeckt würde.

Wir definieren rekursiv eine Folge wie folgt. Sei $x_1 \in X$ beliebig gewählt. Da $X \neq B_\epsilon(x_1)$, existiert $x_2 \in X$ mit $d(x_1, x_2) \geq \epsilon$.

Seien für ein $n \in \mathbb{N}$ $(x_1, x_2, \dots, x_n)$ bereits so gewählt, daß $d(x_i, x_j) \geq \epsilon$ für $1 \leq i \neq j \leq n$. Da $\bigcup_{i=1}^n B_\epsilon(x_i) \neq X$, gibt es ein x_{n+1} mit $d(x_{n+1}, x_j) \geq \epsilon$ für alle $j \leq n$. Somit ist eine Folge in X definiert, deren beliebige zwei Glieder Abstand mindestens ϵ voneinander haben. Daraus folgt, daß diese Folge keine konvergente Teilfolge hat (sonst gäbe es nämlich einen Häufungspunkt $x \in X$ und Indizes $i \neq j$, sodaß $d(x_i, x_j) \leq d(x_i, x) + d(x, x_j) < \epsilon$). Dies ist der gewünschte Widerspruch zur Folgenkompaktheit. $\square$

LEMMA 6.23 (LEBESGUE-Zahl). *Sei X ein folgenkompakter metrischer Raum und $X = \bigcup_{i \in I} V_i$ eine offene Überdeckung. Dann existiert ein $\delta > 0$, sodaß gilt: Für jedes $x \in X$ existiert $i \in I$, sodaß $B_\delta(x) \subset V_i$.*

BEWEIS. Angenommen, dies wäre nicht der Fall. Dann gäbe es zu jedem $n \in \mathbb{N}$ ein $x_n \in X$, sodaß $B_{1/n}(x_n) \not\subset V_i$ für alle $i \in I$.

Aufgrund der Folgenkompaktheit hat die Folge $(x_n)_{n \in \mathbb{N}}$ eine Teilfolge $(x_{n_k})_{k \in \mathbb{N}}$, die gegen $x \in X$ konvergiert. Da $x \in V_i$ für ein $i \in I$, existiert $\epsilon > 0$ mit $B_\epsilon(x) \subset V_i$. Sei nun k so groß, daß einerseits $\frac{1}{n_k} < \frac{\epsilon}{2}$ und andererseits $d(x_{n_k}, x) < \frac{\epsilon}{2}$. Dann gilt für jedes $y \in B_{1/n_k}(x_{n_k})$:

$$d(y, x) \leq d(y, x_{n_k}) + d(x_{n_k}, x) < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon, \quad (6.3)$$

also ist $B_{1/n_k}(x_{n_k}) \subset B_\epsilon(x) \subset V_i$, im Widerspruch zur Konstruktion der Folge $(x_n)_{n \in \mathbb{N}}$. $\square$

SATZ 6.24. Ein metrischer Raum ist genau dann kompakt, wenn er folgenkompakt ist.

BEWEIS. Sei X kompakt und $(x_n)_{n \in \mathbb{N}} \subset X$ eine Folge. Angenommen, Die Folge besäße keine konvergente Teilfolge. Dann existiert zu jedem $y \in X$ eine offene Umgebung V_y , die nur endlich viele x_n enthält (Übung). Da X kompakt ist, existiert eine endliche Teilüberdeckung $X = \bigcup_{j=1}^N V_{y_j}$. Dann würde aber X nur endlich viele Folgenglieder enthalten, Widerspruch.

Sei umgekehrt X folgenkompakt und $X = \bigcup_{i \in I} V_i$ eine offene Überdeckung. Nach Lemma 6.23 existiert $\delta > 0$, sodaß für alle $x \in X$ ein $i_x \in I$ existiert mit $B_\delta(x) \subset V_{i_x}$. Da X nach Proposition 6.22 total beschränkt ist, existieren endlich viele x_j , $j = 1, \dots, N$, sodaß $X = \bigcup_{j=1}^N B_\delta(x_j)$. Daher ist

$$X = \bigcup_{j=1}^N V_{i_{x_j}}$$

die gesuchte endliche Teilüberdeckung. $\square$

6.2.3. Weitere Eigenschaften metrischer Räume. Wir stellen einige Eigenschaften metrischer Räume zusammen, die wir bereits von den reellen Zahlen kennen.

PROPOSITION 6.25. *Eine Teilmenge A eines metrischen Raums ist genau dann abgeschlossen, wenn gilt: Ist $(x_n)_{n \in \mathbb{N}} \subset A$ eine Folge, die gegen x konvergiert, so ist $x \in A$.*

BEWEIS. Sei A abgeschlossen und $(x_n)_{n \in \mathbb{N}}$ eine gegen x konvergente Folge. Wäre $x \notin A$, so wäre $x \in A^c$, was eine offene Menge ist. Das bedeutet, daß es ein $B_\epsilon(x) \subset A^c$ gäbe. Dann aber gälte $d(x, x_n) \geq \epsilon$ für alle $n \in \mathbb{N}$, und somit könnte $(x_n)_{n \in \mathbb{N}}$ nicht gegen x konvergieren.

Für die umgekehrte Implikation sei $x \in A^c$, dann müssen wir zeigen, daß es ein $B_\epsilon(x) \subset A^c$ gibt. Wäre dies nicht der Fall, so gäbe es zu jedem $n \in \mathbb{N}$ ein $x_n \in A$ mit $d(x, x_n) < \frac{1}{n}$, und die so gewonnene Folge $(x_n)_{n \in \mathbb{N}}$ würde gegen x konvergieren. Nach Annahme wäre dann aber $x \in A$, Widerspruch. $\square$

Wir nennen einen metrischen Raum X *beschränkt*, wenn es ein $x \in X$ und ein $R > 0$ gibt mit $X = B_R(x)$. (Insbesondere ist jeder total beschränkte metrische Raum beschränkt, aber nicht umgekehrt.) Ein wichtiger Spezialfall ist der einer Teilmenge $U \subset X$ eines gegebenen metrischen Raums: U selbst ist nämlich ein metrischer Raum, sofern die Metrik d (definiert auf $X \times X$) auf $U \times U$ eingeschränkt wird.

Existiert nun ein $X \supset B_R(x) \supset U$, so ist U beschränkt, denn: Erstens dürfen wir o.B.d.A. annehmen $x \in U$ (falls nämlich $x \notin U$, so wähle $x' \in U \subset B_R(x)$, und dann ist nach Dreiecksungleichung $U \subset B_{R+d(x,x')}(x')$). Zweitens ist die Kugel um x mit Radius R im metrischen Raum $(U, d|_{U \times U})$ genau $B_R(x) \cap U$. Wir verwenden außerdem die Konvention, daß die leere Menge beschränkt ist.

PROPOSITION 6.26. *Sei X ein metrischer Raum und $U \subset X$. Ist U kompakt, so ist es auch beschränkt und abgeschlossen.*

BEWEIS. Sei U kompakt und $x \in U$ beliebig (wir nehmen an $U \neq \emptyset$ – falls doch, ist die Behauptung trivial). Es ist $U \subset \bigcup_{R>0} B_R(x)$ eine offene Überdeckung von X (da die Distanz eines Punktes in U von x stets endlich ist), die eine endliche Teilüberdeckung $U \subset \bigcup_{i=1}^n B_{R_i}(x)$ zuläßt. Mit $R := \max_{i=1,\dots,n} R_i$ ist dann aber $U \subset B_R(x)$, also ist U beschränkt.

Sei $(x_n)_{n \in \mathbb{N}} \subset U$ eine Folge mit $\lim_{n \rightarrow \infty} x_n = x$. Da U nach Satz 6.24 folgenkompakt ist, enthält $(x_n)_{n \in \mathbb{N}}$ eine in U konvergente Teilfolge. Deren Limes kann aber nur x sein, also ist $x \in U$, und die Abgeschlossenheit von U folgt aus Proposition 6.25. $\square$

KOROLLAR 6.27 (vgl. Korollar 3.17). Sei X ein kompakter topologischer Raum und $f: X \rightarrow \mathbb{R}$ stetig. Dann nimmt f sein Maximum und Minimum an.

BEWEIS. Nach Satz 6.15 ist $f(X) \subset \mathbb{R}$ kompakt, also insbesondere beschränkt und abgeschlossen. Da eine nichtleere beschränkte Menge ein endliches Supremum hat (Satz 2.28), gibt es eine Folge $(x_n)_{n \in \mathbb{N}} \subset f(X)$, die gegen $\sup f(X)$ konvergiert. Da aber $f(X)$ abgeschlossen ist, ist $\sup f(X) \in f(X)$ (siehe Proposition 6.25), also ist das Supremum sogar das Maximum. Analog argumentiert man für das Minimum. $\square$

Schließlich noch einige Aussagen über stetige Funktionen:

SATZ 6.28. Seien X, Y metrische Räume. Eine Abbildung $f: X \rightarrow Y$ ist genau dann stetig in $x \in X$, wenn sie folgenstetig in x ist.

BEWEIS. Sei f stetig in x und $\lim_{n \rightarrow \infty} x_n = x$. Sei $U \ni f(x)$ eine offene Umgebung, dann ist wegen der Stetigkeit $f^{-1}(U)$ eine offene Umgebung von x . Also liegen alle bis auf endlich viele x_n in $f^{-1}(U)$. Daher liegen auch alle bis auf endlich viele $f(x_n)$ in U , und somit konvergiert $(f(x_n))_{n \in \mathbb{N}}$ gegen $f(x)$.

Sei nun umgekehrt f in x folgenstetig und $U \subset Y$ eine offene Umgebung von $f(x)$. Sei $(x_n)_{n \in \mathbb{N}} \subset f^{-1}(U)^c$ konvergent gegen $\tilde{x}$. Da f folgenstetig ist, gilt $f(\tilde{x}) = \lim_{n \rightarrow \infty} f(x_n)$. Da aber $f(x_n) \in U^c$, kann $f(x_n)$ nicht gegen ein Element von U konvergieren, also ist $\tilde{x} \in f^{-1}(U)^c$, und nach Proposition 6.25 ist $f^{-1}(U)^c$ abgeschlossen, also $f^{-1}(U)$ offen. $\square$

KOROLLAR 6.29. Seien X, Y, Z metrische Räume und $f: X \rightarrow Y$, $g: Y \rightarrow Z$ stetig. Dann ist auch $g \circ f: X \rightarrow Z$ stetig.

BEWEIS. Ist $(x_n)_{n \in \mathbb{N}}$ eine Folge in X mit $\lim_{n \rightarrow \infty} x_n = x$, so folgt aus der (Folgen-)Stetigkeit von f auch $\lim_{n \rightarrow \infty} f(x_n) = f(x)$, und wegen der Stetigkeit von g schließlich $\lim_{n \rightarrow \infty} g(f(x_n)) = g(f(x))$. $\square$

Es ist leicht zu sehen (Übung), daß eine Funktion von einem metrischen Raum (X, d_X) in einen weiteren metrischen Raum (Y, d_Y) genau dann stetig ist, wenn gilt:

$$\forall \epsilon > 0 \quad \forall x \in X \quad \exists \delta > 0 \quad \forall y \in X: \quad d_X(x, y) < \delta \Rightarrow d_Y(f(x), f(y)) < \epsilon.$$

Wie im Reellen heißt eine Funktion von einem metrischen Raum (X, d_X) in einen metrischen Raum (Y, d_Y) sogar *gleichmäßig stetig*, wenn gilt:

$$\forall \epsilon > 0 \quad \exists \delta > 0 \quad \forall x, y \in X: \quad d_X(x, y) < \delta \Rightarrow d_Y(f(x), f(y)) < \epsilon.$$

SATZ 6.30 (Vgl. Satz 3.20). Sei (X, d_X) ein kompakter metrischer Raum und (Y, d_Y) ein metrischer Raum. Dann ist jede stetige Abbildung $X \rightarrow Y$ sogar gleichmäßig stetig.

BEWEIS. Sei $f : X \rightarrow Y$ stetig und $\epsilon > 0$. Da f stetig ist, existiert zu jedem $x \in X$ ein $\delta(x) > 0$, sodaß aus $d_X(x, y) < \delta(x)$ folgt $d_Y(f(x), f(y)) < \frac{\epsilon}{2}$. Nun ist

$$X = \bigcup_{x \in X} B_{\delta(x)/2}(x)$$

eine offene Überdeckung von X , und da X kompakt ist, existiert eine endliche Teilüberdeckung $X = \bigcup_{i=1}^N B_{\delta(x_i)/2}(x_i)$. Setze $\delta := \frac{1}{2} \min_{i=1, \dots, N} \delta(x_i) > 0$, dann ist jedes $x \in X$ in einem $B_{\delta(x_j)/2}(x_j)$ enthalten, und falls $d_X(x, y) < \delta$, so ist auch $y \in B_{\delta(x_j)/2}(x_j)$. Daher gilt

$$d_Y(f(x), f(y)) \leq d_Y(f(x), f(x_j)) + d_Y(f(x_j), f(y)) < \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon.$$

□

Eine Folge $(f_n)_{n \in \mathbb{N}}$ von Funktionen von einem metrischen Raum X in einen metrischen Raum Y heißt *gleichmäßig konvergent* gegen $f : X \rightarrow Y$, falls

$$\lim_{n \rightarrow \infty} \sup_{x \in X} d(f(x), f_n(x)) = 0.$$

SATZ 6.31 (Gleichmäßige Limites stetiger Funktionen sind stetig, vgl. Satz 5.31). Sei $(f_n)_{n \in \mathbb{N}}$ eine Folge stetiger Funktionen $X \rightarrow Y$, die gleichmäßig gegen $f : X \rightarrow Y$ konvergiert. Dann ist f stetig.

BEWEIS. Wörtlich wie für Satz 5.31, sofern Ausdrücke der Form $|f - g|$ durch $d(f, g)$ ersetzt werden. □

6.2.4. Vollständigkeit. Im Gegensatz zu $\mathbb{Q}$ hat $\mathbb{R}$ die Eigenschaft, daß jede Cauchyfolge konvergiert. Wir haben diese Eigenschaft als *Vollständigkeit* bezeichnet. Ebenso wie in $\mathbb{R}$ heißt in einem metrischen Raum X eine Folge $(x_n)_{n \in \mathbb{N}}$ *Cauchy*, wenn gilt:

$$\forall \epsilon > 0 \quad \exists N \in \mathbb{N} \quad \forall n, m \geq N : \quad d(x_n, x_m) < \epsilon.$$

Wie im Reellen ist klar, daß jede konvergente Folge Cauchy ist.

DEFINITION 6.32. Ein metrischer Raum heißt *vollständig*, wenn in ihm jede Cauchyfolge konvergiert.

BEISPIEL 6.33 (Vollständigkeit von $BC(X; \mathbb{R}^n)$). Sei X ein metrischer Raum und $(BC(X; \mathbb{R}^n); d_\infty)$ der Raum der beschränkten stetigen Funktionen $X \rightarrow \mathbb{R}^n$, versehen mit der Metrik $d_\infty(f, g) := \sup_{x \in X} |f(x) - g(x)|$. Man beachte, daß die Konvergenz bezüglich dieser Metrik genau die gleichmäßige Konvergenz ist. Wir zeigen, daß $(BC(X; \mathbb{R}^n); d_\infty)$ vollständig ist.

BEWEIS. Sei $(f_k)_{k \in \mathbb{N}} \subset BC(X; \mathbb{R}^n)$ Cauchy. Nach Voraussetzung gibt es also zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$, sodaß für alle $k, l \geq N$ gilt $d_\infty(f_k, f_l) < \epsilon$. Insbesondere gilt dies punktweise, d.h. für jedes $x \in X$ gilt $|f_k(x) - f_l(x)| < \epsilon$. Da $\mathbb{R}^n$ vollständig ist (wie wir unten zeigen werden), folgt die punktweise Konvergenz $f_k(x) \rightarrow f(x)$ für eine Abbildung $f : X \rightarrow \mathbb{R}^n$. Wir zeigen, daß diese Konvergenz sogar gleichmäßig ist: In der Tat, für N wie oben und für alle $k, l \geq N$ gilt für alle $x \in X$

$$|f_k(x) - f_l(x)| < \epsilon.$$

Für festes x nehmen wir den Grenzwert $k \rightarrow \infty$, nutzen die punktweise Konvergenz $\lim_{k \rightarrow \infty} f_k(x) = f(x)$ aus und erhalten so

$$\lim_{k \rightarrow \infty} |f_k(x) - f_l(x)| = |f(x) - f_l(x)| \leq \epsilon.$$

Da N nicht von x abhng, erhalten wir wie behauptet die gleichmäßige Konvergenz $f_l \rightarrow f$.

Als gleichmäßiger Limes beschränkter Funktionen ist f beschränkt (wähle etwa N so groß, daß $\sup_{x \in X} |f(x) - f_N(x)| < 1$, dann ist nach Dreiecksungleichung $\sup_{x \in X} |f(x)| \leq \sup_{x \in X} |f(x) - f_N(x)| + \sup_{x \in X} |f_N(x)| < \infty$), und aus Satz 6.31 folgt die Stetigkeit von f . □

Der vielleicht wichtigste Satz über vollständige metrische Räume stammt von BANACH:

SATZ 6.34 (BANACHScher Fixpunktsatz). Sei (X, d) ein vollständiger metrischer Raum und $T : X \rightarrow X$ eine *Kontraktion*, d.h. es existiert ein $0 \leq \theta < 1$, sodaß

$$d(T(x), T(y)) \leq \theta d(x, y) \quad \forall x, y \in X.$$

Dann hat T genau einen *Fixpunkt*, d.h. es existiert genau ein $\bar{x} \in X$ mit $T(\bar{x}) = \bar{x}$.

BEWEIS. *Schritt 1.* Sei $x^0 \in X$ beliebig. Wir definieren rekursiv eine Folge $(x_k)_{k \in \mathbb{N}}$ durch

$$x_0 = x^0, \quad x_{k+1} = T(x_k) \quad (k \geq 0).$$

Wir zeigen, daß diese Folge Cauchy ist. Seien dazu $k, l \in \mathbb{N}$ mit $k \geq l$, so ist

$$\begin{aligned} d(x_k, x_l) &\leq \sum_{j=l}^{k-1} d(x_{j+1}, x_j) \\ &= \sum_{j=l}^{k-1} d(T^j(x_1), T^j(x^0)) \\ &\leq \sum_{j=l}^{k-1} \theta^j d(x_1, x^0) \\ &= \theta^l \sum_{j=0}^{k-l-1} \theta^j d(x_1, x^0) \leq d(x_1, x^0) \frac{\theta^l}{1-\theta}, \end{aligned}$$

wobei wir zuletzt die Summenformel für die geometrische Reihe verwendet haben. Da $\lim_{l \rightarrow \infty} \theta^l = 0$, folgt die Cauchy-Eigenschaft der Folge.

Schritt 2. Da X vollständig ist, konvergiert $(x_k)_{k \in \mathbb{N}}$ gegen einen Grenzwert $\bar{x} \in X$. Wir zeigen, daß $\bar{x}$ Fixpunkt von T ist. Zunächst bemerken wir, daß T als Kontraktion stetig ist, denn für $\epsilon > 0$ folgt aus $d(x, y) < \epsilon$ auch $d(T(x), T(y)) \leq \theta d(x, y) < \epsilon$. Daher ist

$$T(\bar{x}) = T\left(\lim_{k \rightarrow \infty} x_k\right) = \lim_{k \rightarrow \infty} T(x_k) = \lim_{k \rightarrow \infty} x_{k+1} = \bar{x}.$$

Schritt 3. Wir zeigen noch die Eindeutigkeit. Seien dazu $\bar{x}, \bar{y}$ Fixpunkte von T , so gilt

$$d := d(\bar{x}, \bar{y}) = d(T(\bar{x}), T(\bar{y})) \leq \theta d(\bar{x}, \bar{y}) = \theta d,$$

und wegen $\theta < 1$ folgt $d = 0$, also $\bar{x} = \bar{y}$. □

BEISPIEL 6.35. Sei $f : [0, 1] \rightarrow [0, 1]$ eine stetig differenzierbare Funktion mit $|f'(x)| < \frac{2}{3}$ für alle $x \in [0, 1]$. Dann ist f eine Kontraktion, denn nach Mittelwertsatz gibt es für $x, y \in [0, 1]$ ein $\xi \in (0, 1)$ mit

$$|f(x) - f(y)| = |f'(\xi)| |x - y| \leq \frac{2}{3} |x - y|.$$

Nach dem Satz von Banach existiert also genau ein $x \in [0, 1]$ mit $f(x) = x$.

Wählt man hier nur das *offene* Intervall $(0, 1)$, so ist der Definitionsbereich von f nicht vollständig (z.B. ist die Folge $(1/n)_{n \in \mathbb{N}}$ Cauchy, aber nicht konvergent in $(0, 1)$), und das Beispiel $f(x) = \frac{1}{2}x$ zeigt, daß die Aussage des Banachschen Fixpunktsatzes nun nicht mehr wahr ist.

Das Beispiel $f(x) = x$ zeigt, daß für $\theta = 1$ die Aussage des Satzes ebenfalls nicht mehr gültig ist (hier die Eindeutigkeit, im allgemeinen aber auch die Existenz).

6.3. Die Topologie des $\mathbb{R}^n$

Im weiteren Verlauf der Vorlesung werden wir uns hauptsächlich mit dem $\mathbb{R}^n$ beschäftigen. Da dieser ein (wie wir gleich zeigen werden vollständiger) metrischer Raum ist, treffen alle bisher gemachten Aussagen auch auf $\mathbb{R}^n$ zu. Darüber hinaus gibt es einige Eigenschaften, die für den $\mathbb{R}^n$ spezifisch sind.

Zunächst trägt der $\mathbb{R}^n$ nicht nur eine metrische, sondern auch eine algebraische Struktur: Es handelt sich nämlich um einen $\mathbb{R}$ -Vektorraum. Metrische Räume, deren Topologie mit der Vektorraumstruktur verträglich sind, heißen *normierte Räume*; genauer:

DEFINITION 6.36 (Normierter Raum). Sei $\mathbb{K}$ ein Körper. Ein *normierter Raum* ist ein $\mathbb{K}$ -Vektorraum X zusammen mit einer Abbildung (der *Norm*) $\|\cdot\| : X \rightarrow \mathbb{R}$, sodaß gilt:

- i) *Positivität*: $\|x\| \geq 0$ für alle $x \in X$, mit Gleichheit genau dann, wenn $x = 0$;
- ii) *Homogenität*: $\|\lambda x\| = |\lambda| \|x\|$ für alle $\lambda \in \mathbb{K}$ und $x \in X$;
- iii) *Dreiecksungleichung*: $\|x + y\| \leq \|x\| + \|y\|$ für alle $x, y \in X$.

Jeder normierte Raum ist insbesondere ein metrischer Raum, denn durch $d(x, y) := \|x - y\|$ ist eine Metrik definiert, wie man unschwer nachprüft. Eine Folge $(x_n)_{n \in \mathbb{N}}$ in einem normierten Raum konvergiert also gegen x , wenn für jedes $\epsilon > 0$ ein $N \in \mathbb{N}$ existiert, sodaß für alle $n \geq N$ gilt $\|x_n - x\| < \epsilon$.

Wir verzichten im folgenden auf allgemeine Aussagen über normierte Räume und halten lediglich fest, daß $\mathbb{R}^n$ mit der euklidischen Norm $|x| = (\sum_{i=1}^n x_i^2)^{1/2}$ die Axiome eines normierten Raums erfüllt, wie aus der linearen Algebra bekannt sein dürfte.

Zur Notation: Für einen Vektor $x \in \mathbb{R}^n$ schreiben wir von nun an stets x_i für die i -te Komponente; betrachten wir Folgen von Vektoren, so notieren wir den Folgenindex oben. Ist also $(x^k)_{k \in \mathbb{N}}$ eine Folge von Vektoren in $\mathbb{R}^n$, so bezeichnet x_i^k die i -te Komponente des k -ten Folgenglieds.

PROPOSITION 6.37. *Eine Folge $(x^k)_{k \in \mathbb{N}} \subset \mathbb{R}^n$ konvergiert gegen $x \in \mathbb{R}^n$ genau dann, wenn $\lim_{k \rightarrow \infty} x_i^k = x_i$ für jedes $i = 1, \dots, n$.*

BEWEIS. Gelte zunächst $\lim_{k \rightarrow \infty} x^k = x$, das heißt $\lim_{k \rightarrow \infty} (\sum_{i=1}^n |x_i^k - x_i|^2)^{1/2} = 0$. Da aber für jedes $i = 1, \dots, n$ gilt $|x_i^k - x_i| \leq (\sum_{i=1}^n |x_i^k - x_i|^2)^{1/2} = 0$, folgt auch $\lim_{k \rightarrow \infty} x_i^k = x_i$.

Gelte umgekehrt $\lim_{k \rightarrow \infty} x_i^k = x_i$ für jedes i , und sei $\epsilon > 0$. Dann gibt es zu jedem $i = 1, \dots, n$ ein N_i , sodaß für alle $k \geq N_i$ gilt $|x_i^k - x_i| < \epsilon$. Sei $N := \max_{i=1, \dots, n} N_i$. Dann gilt für jedes $k \geq N$

$$|x^k - x|^2 = \sum_{i=1}^n |x_i^k - x_i|^2 < n\epsilon^2,$$

und da $\epsilon > 0$ beliebig war, folgt die behauptete Konvergenz. $\square$

KOROLLAR 6.38. Sei X ein metrischer Raum. Eine Abbildung $f : X \rightarrow \mathbb{R}^n$ ist genau dann stetig, wenn jede Komponente $f_i : X \rightarrow \mathbb{R}$ stetig ist.

BEWEIS. Sei $(x^k)_{k \in \mathbb{N}}$ eine konvergente Folge in X . Dann folgt die Aussage durch Anwendung von Proposition 6.37 auf die Folge $(f(x^k))_{k \in \mathbb{N}} \subset \mathbb{R}^n$. $\square$

KOROLLAR 6.39. $\mathbb{R}^n$ ist vollständig.

BEWEIS. Sei $(x^k)_{k \in \mathbb{N}}$ Cauchy. Wegen $|x_i^k - x_i^l| \leq |x^k - x^l|$ für alle $i = 1, \dots, n$ ist dann auch die reelle Folge $(x_i^k)_{k \in \mathbb{N}}$ Cauchy für jedes i , und aufgrund der Vollständigkeit von $\mathbb{R}$ konvergiert diese Folge. Nach Proposition 6.37 konvergiert daher auch $(x^k)_{k \in \mathbb{N}}$ in $\mathbb{R}^n$. $\square$

KOROLLAR 6.40 (Schachtelungsprinzip). Sei $(A^k)_{k \in \mathbb{N}}$ eine Folge nichtleerer abgeschlossener Teilmengen von $\mathbb{R}^n$, sodaß $A^k \supset A^{k+1}$ für alle $k \in \mathbb{N}$, und

$$\lim_{k \in \mathbb{N}} \sup_{x, y \in A^k} |x - y| = 0$$

(d.h. der Durchmesser der A^k konvergiert gegen null). Dann existiert genau ein $x^\infty \in \mathbb{R}^n$, das in jedem A^k enthalten ist.

BEWEIS. Wähle zu jedem $k \in \mathbb{N}$ ein $x^k \in A^k$. Sind $k, l \geq N$, so gilt wegen $A^k, A^l \subset A^N$:

$$|x^k - x^l| \leq \sup_{x, y \in A^k} |x - y|,$$

was nach Voraussetzung beliebig klein wird, wenn N hinreichend groß ist. Also ist die Folge $(x^k)_{k \in \mathbb{N}}$ Cauchy und nach der vorigen Proposition konvergent gegen x^∞ . Da für beliebiges N gilt $x^k \in A^N$ für alle $k \geq N$, und da A^N abgeschlossen ist, liegt x^∞ in A^N (Proposition 6.25).

Zur Eindeutigkeit: Liegen x, x' beide in A^k für alle k , so gilt

$$|x - x'| \leq \sup_{y, z \in A^k} |y - z|,$$

und da letzteres mit $k \rightarrow \infty$ gegen null konvergiert, folgt $x = x'$. $\square$

SATZ 6.41 (HEINE-BOREL). Eine Teilmenge des $\mathbb{R}^n$ ist genau dann kompakt, wenn sie beschränkt und abgeschlossen ist.

BEWEIS. Die erste Implikation ist Proposition 6.26. Sei also umgekehrt $K \subset \mathbb{R}^n$ beschränkt und abgeschlossen. Da K beschränkt ist, existiert ein Würfel $[-R, R]^n$, der K enthält. Wir zeigen zunächst die Kompaktheit dieses Würfels: Sei $(x^k)_{k \in \mathbb{N}} \subset [-R, R]^n$, so folgt für alle $i = 1, \dots, n$, daß $(x_i^k)_{k \in \mathbb{N}} \subset [-R, R]$. Da das Intervall $[-R, R]$ kompakt ist (Beispiel 6.14), gibt es eine konvergente Teilfolge $(x_1^{k_l})_{l \in \mathbb{N}}$. Aus der Teilfolge $(x^{k_l})_{l \in \mathbb{N}}$ können wir aber eine weitere Teilfolge auswählen, deren zweite Komponente ebenfalls konvergiert. Die n -malige Wiederholung dieser Teilfolgenauswahl (für jede der n Komponenten) liefert schließlich eine Teilfolge, deren Komponenten alle konvergieren. Nach Proposition 6.37 konvergiert also diese Teilfolge in $[-R, R]^n$.

Sei schließlich $(x^k)_{k \in \mathbb{N}}$ eine Folge in K . Da $K \subset [-R, R]^n$, konvergiert die Folge. Da aber K abgeschlossen ist, liegt der Grenzwert sogar in K . Damit ist die Kompaktheit von K bewiesen. $\square$

Die folgende Aussage über Stetigkeit sei zur Übung empfohlen:

PROPOSITION 6.42. Sei X ein metrischer Raum und $f, g : X \rightarrow \mathbb{R}$ stetig. Dann sind auch $f \pm g$ und fg stetig. Ist $X' := X \setminus \{x \in X : g(x) = 0\}$, so ist außerdem $\frac{f}{g} : X' \rightarrow \mathbb{R}$ stetig.

Differentialrechnung in mehreren Variablen

Im Bezug auf Differentiation und Integration haben wir uns bisher auf Funktionen einer Variablen fokussiert. In diesem Kapitel soll es um Funktionen $f : \Omega \subset \mathbb{R}^n \rightarrow \mathbb{R}^m$ gehen, die also von n unabhängigen Variablen abhängen. Die Relevanz solcher Abbildungen ist offenkundig: Physikalische Felder (z.B. elektromagnetische Potentiale, Temperaturfelder etc.) hängen von drei Raum- und einer Zeitvariablen ab, der Preis eines Produkts variiert unter anderem mit den Rohstoffkosten und den Lohnkosten, und die Beschreibung höherdimensionaler geometrischer Objekte kann (zumindest lokal) mittels Funktionen mehrerer Veränderlicher erfolgen.

7.1. Partielle Differentiation

7.1.1. Einführung. Sei stets $\Omega \subset \mathbb{R}^n$ eine offene Teilmenge. Wir betrachten Abbildungen $f = (f_1, \dots, f_m) : \Omega \rightarrow \mathbb{R}^m$. Die reellwertige Abbildung $f_k : \Omega \rightarrow \mathbb{R}$ heißt *k-te Komponente* der (vektorwertigen) Abbildung f . Im Falle $m = 1$ nennt man f bisweilen ein *Skalarfeld*, für $n = m > 1$ bezeichnet man es als *Vektorfeld*.

Der *Graph* von f ist die Teilmenge $G_f := \{(x, y) \in \Omega \times \mathbb{R}^m : f(x) = y\}$. Im skalaren Falle ($m = 1$) bezeichnet man die Mengen

$$N_c(f) := \{x \in \Omega : f(x) = c\}, \quad c \in \mathbb{R}$$

als *Niveaumengen*, für $n = 2$ auch als *Niveau- bzw. Höhenlinien*. Ob diese Mengen tatsächlich ‚Linien‘ sind, wird uns noch beschäftigen.

- BEISPIEL 7.1. (1) Sei $\mathbb{R}^2 \supset \Omega = B_1(0) := \{(x, y) \in \mathbb{R}^2 : x^2 + y^2 < 1\}$ und $f : \Omega \rightarrow \mathbb{R}$ gegeben durch $f(x, y) = \sqrt{1 - x^2 - y^2}$. Dann ist der Graph von f die obere zweidimensionale Halbsphäre, und die Niveaulinien sind Kreise in der (x, y) -Ebene.
- (2) Sei $\Omega = (0, \infty) \times \mathbb{R}$ und $f(t, x) = \sin(x - ct)$ für einen Parameter $c \in \mathbb{R}$. Diese Funktion kann interpretiert werden als Sinuswelle, die sich mit der Geschwindigkeit c fortbewegt. Die Niveaulinien sind Geraden in der (t, x) -Ebene mit Steigung c .
- (3) Die Abbildung $f : [0, 2\pi) \rightarrow \mathbb{R}^3$, $f(t) = (\cos t, \sin t, t)$ beschreibt eine *Schraubenlinie*.
- (4) Man skizziere die beiden Vektorfelder $f, g : \mathbb{R}^2 \rightarrow \mathbb{R}^2$, gegeben durch

$$f(x, y) = (x, y)^\perp := (-y, x),$$

$$g(x, y) = (x, y).$$

7.1.2. Partielle Ableitungen. Betrachte für $f : \Omega \rightarrow \mathbb{R}$, $l = 1, \dots, n$ und $x \in \Omega$ die Funktion

$$\xi \mapsto f(x_1, \dots, x_{l-1}, \xi, x_{l+1}, \dots, x_n),$$

die in einer Umgebung von x_l definiert ist. Ist diese Funktion (von einer Variablen ξ) in x_l differenzierbar, so heißt f in x *partiell differenzierbar* nach der Variablen x_l , und

$$\lim_{\xi \rightarrow x_l} \frac{f(x_1, \dots, x_{l-1}, \xi, x_{l+1}, \dots, x_n) - f(x)}{\xi - x_l}$$

heißt *partielle Ableitung* von f nach x_l an der Stelle $x \in \Omega$. Man schreibt dafür auch

$$\frac{\partial f}{\partial x_l}(x), \quad \partial_{x_l} f(x), \quad \partial_l f(x), \quad D_{x_l} f(x), \quad f_{x_l}(x), \quad \text{etc.}$$

Ist f am Punkt $x \in \Omega$ partiell differenzierbar in jeder Richtung $l = 1, \dots, n$, so heißt f partiell differenzierbar in x .

Bei der partiellen Differentiation hält man also alle außer der l -ten Variablen fest und differenziert wie im Eindimensionalen.

DEFINITION 7.2. Sei $f: \Omega \rightarrow \mathbb{R}^m$ in $x \in \Omega$ partiell differenzierbar (d.h. jede Komponente von f ist partiell differenzierbar). Dann heißt die Matrix

$$Df(x) = (\partial_l f_k(x))_{\substack{k=1, \dots, m \\ l=1, \dots, n}} = \begin{pmatrix} \partial_1 f_1(x) & \partial_2 f_1(x) & \dots & \partial_n f_1(x) \\ \partial_1 f_2(x) & \partial_2 f_2(x) & \dots & \partial_n f_2(x) \\ \vdots & \vdots & \ddots & \vdots \\ \partial_1 f_m(x) & \partial_2 f_m(x) & \dots & \partial_n f_m(x) \end{pmatrix} \in \mathbb{R}^{m \times n}$$

Jacobi-Matrix von f an der Stelle x .

Ist $m = 1$, so ist die Jacobi-Matrix einfach ein Vektor in $\mathbb{R}^n$, der als *Gradient* von f in x bezeichnet wird. Man verwendet dafür die Schreibweisen $\nabla f(x)$, $\text{grad } f(x)$ oder $Df(x)$.

BEISPIEL 7.3. Betrachte die Funktion $r: \mathbb{R}^n \rightarrow \mathbb{R}$,

$$r(x) = |x| = \sqrt{x_1^2 + \dots + x_n^2}.$$

Dann ist r in jedem Punkt $x \neq 0$ partiell differenzierbar, und es gilt

$$\partial_l r(x) = \frac{2x_l}{2\sqrt{x_1^2 + \dots + x_n^2}} = \frac{x_l}{r},$$

und daher $\nabla r = \frac{x}{r}$.

Betrachte weiterhin die Funktion $f: \mathbb{R}^2 \rightarrow \mathbb{R}$,

$$f(x) = \begin{cases} \frac{x_1 x_2}{r^4} & \text{falls } x \neq 0, \\ 0 & \text{falls } x = 0. \end{cases}$$

Für $x \neq 0$ ergibt die Ableitung nach x_1

$$\partial_1 f(x) = \frac{r^4 x_2 - 4x_1 x_2 r^3 \frac{x_1}{r}}{r^8} = \frac{x_2(r^2 - 4x_1^2)}{r^6}$$

und, da die Rollen von x_1 und x_2 vertauscht werden können,

$$\partial_2 f(x) = \frac{x_1(r^2 - 4x_2^2)}{r^6}.$$

In $x = 0$ ist f aber auch partiell differenzierbar, denn

$$\frac{f(h, 0) - f(0, 0)}{h} = 0$$

für jedes $h \neq 0$, und somit ist $\partial_1 f(0) = 0$ und analog $\partial_2 f(0) = 0$. Die Funktion f ist also auf ganz $\mathbb{R}^2$ partiell differenzierbar. Andererseits ist sie in $x = 0$ nicht stetig: Betrachte die Folge $(x^k)_{k \in \mathbb{N}}$ mit

$$x^k = \left(\frac{1}{k}, \frac{1}{k} \right),$$

sodaß $x^k \rightarrow 0$, aber $f(x^k) = \frac{k^2}{4}$, was nicht konvergiert.

Dieses Beispiel zeigt, daß (anders als im eindimensionalen Falle) partielle Differenzierbarkeit *nicht* die Stetigkeit impliziert.

7.1.3. Höhere partielle Ableitungen. Ist $f : \Omega \rightarrow \mathbb{R}^m$ an jedem Punkt $x \in \Omega$ nach x_l partiell differenzierbar, so kann man wie im eindimensionalen Falle die partielle Ableitung $\partial_l f$ ihrerseits als Funktion $\Omega \rightarrow \mathbb{R}^m$ auffassen. Ist diese in $x \in \Omega$ stetig, so heißt f in x stetig partiell differenzierbar nach x_l ; ist f in $x \in \Omega$ nach jeder Richtung x_l , $l = 1, \dots, n$ stetig partiell differenzierbar, so heißt sie stetig partiell differenzierbar in x ; ist sie schließlich in jedem $x \in \Omega$ stetig partiell differenzierbar, so heißt sie stetig partiell differenzierbar in Ω .

Ist f stetig partiell differenzierbar, so kann die Funktion $\partial_l f : \Omega \rightarrow \mathbb{R}^m$ selbst wieder partiell differenzierbar sein; ist dies nach allen Richtungen der Fall, so heißt f zweimal partiell differenzierbar, und sind alle zweiten Ableitungen $\partial_{kl} f := \partial_k(\partial_l f)$ ($k, l = 1, \dots, n$) stetig in Ω , so heißt die Funktion zweimal stetig partiell differenzierbar. Analog definiert man höhere partielle Ableitungen.

Wieder gibt es zahlreiche Schreibweisen für höhere partielle Ableitungen:

$$\frac{\partial^2 f}{\partial x_k \partial x_l}, \quad \partial_k \partial_l f, \quad \partial_{kl}^2 f, \quad f_{x_k x_l}, \quad \text{etc.}$$

Partielle Ableitungen kommutieren miteinander:

SATZ 7.4 (Satz von SCHWARZ). Ist $f : \Omega \rightarrow \mathbb{R}^m$ zweimal stetig partiell differenzierbar, so gilt für jedes $x \in \Omega$

$$\partial_{kl} f(x) = \partial_{lk} f(x).$$

BEWEIS. Ohne Beschränkung der Allgemeinheit können wir annehmen $n = 2$, denn sonst betrachte für $x \in \Omega$ die Funktion

$$\tilde{f} : \tilde{\Omega} \subset \mathbb{R}^2 \rightarrow \mathbb{R}^m, \quad \tilde{f}(\xi, \eta) := f(x_1, \dots, x_{k-1}, \xi, \dots, x_{l-1}, \eta, \dots, x_n),$$

die in einer Umgebung $\tilde{\Omega} \ni (x_k, x_l)$ definiert ist.

Schritt 1. Sei also $x = (x_1, x_2) \in \Omega$. Da Ω offen ist, existiert ein $\delta > 0$ mit

$$[x_1 - \delta, x_1 + \delta] \times [x_2 - \delta, x_2 + \delta] \subset \Omega.$$

Für festes $\eta \in (x_2 - \delta, x_2 + \delta)$ definiere die Funktion $F_\eta : (x_1 - \delta, x_1 + \delta) \rightarrow \mathbb{R}^m$ durch

$$F_\eta(\xi) = f(\xi, \eta) - f(\xi, x_2).$$

Dann ist F_η stetig differenzierbar auf $(x_1 - \delta, x_1 + \delta)$, und nach dem Mittelwertsatz existiert für jedes ξ ein $\xi' \in (x_1, \xi)^1$ mit

$$F'_\eta(\xi')(\xi - x_1) = F_\eta(\xi) - F_\eta(x_1) = f(\xi, \eta) - f(\xi, x_2) - f(x_1, \eta) + f(x_1, x_2). \quad (7.1)$$

Nun ist aber $F'_\eta(\xi') = \partial_1 f(\xi', \eta) - \partial_1 f(\xi', x_2)$, und erneute Anwendung des Mittelwertsatzes auf die Funktion $\eta \mapsto \partial_1 f(\xi', \eta)$ liefert ein $\eta' \in (x_2, \eta)$, sodaß

$$\partial_{21} f(\xi', \eta')(\eta - x_2) = \partial_1 f(\xi', \eta) - \partial_1 f(\xi', x_2),$$

und mit (7.1) folgt

$$\partial_{21} f(\xi', \eta')(\eta - x_2)(\xi - x_1) = f(\xi, \eta) - f(\xi, x_2) - f(x_1, \eta) + f(x_1, x_2).$$

Schritt 2. Betrachte nun für $\xi \in (x_1 - \delta, x_1 + \delta)$ die Funktion

$$G_\xi(\eta) = f(\xi, \eta) - f(x_1, \eta).$$

Nach dem Mittelwertsatz existiert für jedes η ein $\eta'' \in (x_2, \eta)$ mit

$$G'_\xi(\eta'')(\eta - x_2) = f(\xi, \eta) - f(x_1, \eta) - f(\xi, x_2) + f(x_1, x_2).$$

Nochmalige Anwendung des Mittelwertsatzes auf $\xi \mapsto \partial_2 f(\xi, \eta'')$ ergibt unter Berücksichtigung von $G'_\xi(\eta'') = \partial_2 f(\xi, \eta'') - \partial_2 f(x_1, \eta'')$ die Existenz eines $\xi'' \in (x_1, \xi)$, sodaß

$$\partial_{12} f(\xi'', \eta'')(\xi - x_1)(\eta - x_2) = f(\xi, \eta) - f(x_1, \eta) - f(\xi, x_2) + f(x_1, x_2),$$

¹Mit (x_1, ξ) ist hier das Intervall (x_1, ξ) gemeint, falls $x_1 \leq \xi$, und das Intervall (ξ, x_1) , falls $\xi < x_1$.

also (vgl. Schritt 1) für $\xi \neq x_1$, $\eta \neq x_2$

$$\partial_{12}f(\xi'', \eta'') = \partial_{21}f(\xi', \eta').$$

Da $\partial_{12}f$ und $\partial_{21}f$ nach Voraussetzung stetig sind und mit $(\xi, \eta) \rightarrow (x_1, x_2)$ auch (ξ', η') , $(\xi'', \eta'') \rightarrow (x_1, x_2)$ konvergieren, folgt in der Tat

$$\partial_{12}f(x_1, x_2) = \partial_{21}f(x_1, x_2).$$

□

DEFINITION 7.5. Sei $f : \Omega \rightarrow \mathbb{R}$ in $x \in \Omega$ zweimal partiell differenzierbar. Dann heißt

$$D^2f(x) = (\partial_{kl}f(x))_{\substack{k=1,\dots,n \\ l=1,\dots,n}} = \begin{pmatrix} \partial_{11}f(x) & \partial_{12}f(x) & \dots & \partial_{1n}f(x) \\ \partial_{21}f(x) & \ddots & \ddots & \partial_{2n}f(x) \\ \vdots & \ddots & \ddots & \vdots \\ \partial_{n1}f(x) & \partial_{n2}f(x) & \dots & \partial_{nn}f(x) \end{pmatrix} \in \mathbb{R}^{n \times n}$$

HESSE-Matrix von f an der Stelle x .

Nach dem Satz von SCHWARZ ist die HESSE-Matrix also an jedem Punkt symmetrisch, sofern f zweimal stetig partiell differenzierbar ist.

7.1.4. Spezielle Differentialoperatoren. Für ein partiell differenzierbares Vektorfeld $f : \mathbb{R}^n \supset \Omega \rightarrow \mathbb{R}^n$ definiert man die *Divergenz* $\operatorname{div} f : \Omega \rightarrow \mathbb{R}$,

$$\operatorname{div} f(x) = \sum_{l=1}^n \partial_l f_l(x),$$

und speziell im Falle $n = 3$ definiert man weiterhin die *Rotation*² $\operatorname{rot} f : \Omega \rightarrow \mathbb{R}^3$,

$$\operatorname{rot} f(x) = (\partial_2 f_3(x) - \partial_3 f_2(x), \quad \partial_3 f_1(x) - \partial_1 f_3(x), \quad \partial_1 f_2(x) - \partial_2 f_1(x)).$$

Interpretiert man das Vektorfeld etwa als Strömungsfeld fließenden Wassers, so kann man $\operatorname{div} f$ als infinitesimalen Zu- bzw. Abfluß (Quelle bzw. Senke) auffassen und $\operatorname{rot} f$ bzw. $\nabla^\perp f$ als infinitesimale Rotation des Wassers. Ähnliche Interpretationen treffen z.B. für elektrische und magnetische Felder zu. Man berechne die jeweiligen Operatoren für die in Abschnitt 7.1.1 genannten Vektorfelder, um diese Intuition zu entwickeln.

Von großer Bedeutung ist darüber hinaus der LAPLACE-Operator, der für eine zweimal partiell differenzierbare Funktion $f : \Omega \rightarrow \mathbb{R}$ definiert ist als

$$\Delta f(x) := \operatorname{div} \nabla f(x) = \operatorname{spur} D^2f(x) = \sum_{l=1}^n \partial_{ll}f(x)$$

(die zwei letzten Identitäten sollten als Übung nachgerechnet werden).

7.2. Totale Differentiation

7.2.1. Totale Differenzierbarkeit. In Analysis I haben Sie gesehen, daß eine Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ genau dann im Punkt x_0 differenzierbar ist, wenn sie an diesem Punkt *linear approximiert* werden kann, wenn es also ein $a \in \mathbb{R}$ gibt, sodaß

$$f(x) = f(x_0) + a(x - x_0) + \phi(x), \quad \text{wobei} \quad \lim_{x \rightarrow x_0} \frac{\phi(x)}{|x - x_0|} = 0.$$

In diesem Falle ist $a = f'(x_0)$. Dies motiviert folgende Definition:

²Im Englischen schreibt man meist $\operatorname{curl} f$ für die Rotation von f .

DEFINITION 7.6. Sei $\Omega \subset \mathbb{R}^n$ offen und $f : \Omega \rightarrow \mathbb{R}^m$. Dann heißt f (*total*) *differenzierbar* in $x_0 \in \Omega$, wenn es eine lineare Abbildung $A : \mathbb{R}^n \rightarrow \mathbb{R}^m$ gibt, sodaß für alle $x \in \Omega$

$$f(x) = f(x_0) + A(x - x_0) + \phi(x), \quad (7.2)$$

wobei $\phi : \Omega \rightarrow \mathbb{R}^m$ eine Abbildung ist mit

$$\lim_{x \rightarrow x_0} \frac{\phi(x)}{|x - x_0|} = 0.$$

In diesem Falle heißt A die *totale Ableitung* von f an der Stelle x_0 .

Die Gleichheit (7.2) ist natürlich äquivalent zu

$$f(x_0 + \xi) = f(x_0) + A\xi + \psi(\xi), \quad \lim_{\xi \rightarrow 0} \frac{\psi(\xi)}{|\xi|} = 0.$$

Die totale Ableitung ist, falls sie existiert, eindeutig bestimmt: Sind nämlich A und A' totale Ableitungen von f in x_0 , und sind ϕ, ϕ' die zugehörigen Funktionen höherer Ordnung aus der Definition, so gilt für alle x

$$A(x - x_0) + \phi(x) = A'(x - x_0) + \phi'(x).$$

Division dieser Gleichung durch $|x - x_0|$ ergibt

$$\lim_{x \rightarrow x_0} (A - A') \frac{x - x_0}{|x - x_0|} = 0,$$

und mit der Wahl $x^k = x_0 + \frac{1}{k}e_j$ für einen beliebigen Standardbasisvektor $e_j \in \mathbb{R}^n$ folgt $(A - A')e_j = 0$, also $A = A'$.

BEISPIEL 7.7. Sei $B \in \mathbb{R}^{n \times n}$ eine Matrix und $f : \mathbb{R}^n \rightarrow \mathbb{R}$ gegeben durch $f(x) = (x, Bx)$, wobei $(\cdot, \cdot)$ das Standardskalarprodukt in $\mathbb{R}^n$ bezeichnet. Seien $x, x_0 \in \mathbb{R}^n$, so gilt

$$\begin{aligned} f(x) - f(x_0) &= (x, Bx) - (x_0, Bx_0) \\ &= (x, Bx_0) - (x_0, Bx_0) + (x, B(x - x_0)) \\ &= (x - x_0, Bx_0) + (x_0, B(x - x_0)) + (x - x_0, B(x - x_0)). \end{aligned}$$

Definieren wir $A : \mathbb{R}^n \rightarrow \mathbb{R}$ als $A\xi := (\xi, Bx_0) + (x_0, B\xi)$, so ist A linear. Setzen wir

$$\phi(x) := (x - x_0, B(x - x_0)),$$

so ist

$$\frac{|\phi(x)|}{|x - x_0|} \leq \frac{|(x - x_0, B(x - x_0))|}{|x - x_0|} \leq |B(x - x_0)| \rightarrow 0$$

für $x \rightarrow x_0$.

Also ist f differenzierbar, und die totale Ableitung ist die lineare Abbildung A , die geschrieben werden kann als

$$A\xi = ((B + B^t)x_0, \xi).$$

Sie kann also mit dem Vektor $(B + B^t)x_0 \in \mathbb{R}^n$ identifiziert werden.

Berechnen wir die partiellen Ableitungen von f , so erhalten wir

$$\partial_l f = \partial_l \sum_{j,k=1}^n B_{jk} x_j x_k = \sum_{j=1}^n B_{lj} x_j + \sum_{k=1}^n B_{kl} x_k = [(B + B^t)x]_l,$$

d.h. der Gradient von f an der Stelle x_0 ist genau $(B + B^t)x_0$.

Aus diesem Beispiel kann man die Vermutung herleiten, daß die Jacobi-Matrix (für $m = 1$ also der Gradient) stets mit der totalen Ableitung identifiziert werden kann. Dies ist in der Tat der Fall:

SATZ 7.8. Sei $f : \Omega \rightarrow \mathbb{R}^m$ im Punkt $x_0 \in \Omega$ total differenzierbar mit totaler Ableitung $A : \mathbb{R}^n \rightarrow \mathbb{R}^m$. Dann ist f in x_0 stetig und partiell differenzierbar, und es gilt für alle $\xi \in \mathbb{R}^n$

$$A\xi = Df(x_0)\xi,$$

d.h. die Jacobi-Matrix $Df(x_0) \in \mathbb{R}^{m \times n}$ ist die darstellende Matrix der totalen Ableitung an der Stelle x_0 bezüglich der Standardbasis von $\mathbb{R}^n$ bzw. $\mathbb{R}^m$.

Wir unterscheiden von nun an nicht mehr zwischen einer linearen Abbildung und ihrer darstellenden Matrix bzgl. der Standardbasis. In der Differentialgeometrie ist diese Unterscheidung allerdings wichtig.

BEWEIS. Die Stetigkeit folgt sofort aus der Definition, denn für $x \rightarrow x_0$ ist

$$\lim_{x \rightarrow x_0} f(x) = \lim_{x \rightarrow x_0} (f(x_0) + A(x - x_0) + \phi(x)) = f(x_0).$$

Sei nun $A \in \mathbb{R}^{m \times n}$ die totale Ableitung (hier haben wir bereits die lineare Abbildung mit ihrer darstellenden Matrix identifiziert). Ist $(e_j)_{j=1, \dots, n}$ die Standardbasis in $\mathbb{R}^n$, so gilt für jedes $l = 1, \dots, n$ und $h \in \mathbb{R}$ so klein, daß $x_0 + he_l \in \Omega$ nach Definition der totalen Ableitung

$$f(x_0 + he_l) = f(x_0) + Ahe_l + \phi(x_0 + he_l),$$

und damit

$$\partial_l f(x_0) = \lim_{h \rightarrow 0} \frac{f(x_0 + he_l) - f(x_0)}{h} = \lim_{h \rightarrow 0} \frac{Ahe_l + \phi(x_0 + he_l)}{h} = Ae_l,$$

also in der Tat $Df(x_0) = A$. □

Nach dem gerade bewiesenen Satz ist also jede total differenzierbare Funktion partiell differenzierbar, aber nicht umgekehrt: Wir haben in Abschnitt 7.1.2 eine Funktion studiert, die partiell differenzierbar, aber nicht stetig und damit auch nicht total differenzierbar ist. Der nächste Satz zeigt aber, daß die Stetigkeit der partiellen Ableitungen hinreichend für die totale Differenzierbarkeit ist:

SATZ 7.9. Sei $f : \Omega \rightarrow \mathbb{R}$ in Ω partiell differenzierbar, und seien alle partiellen Ableitungen in x_0 stetig. Dann ist f in x_0 total differenzierbar.

BEWEIS. Wir betrachten nur den Fall $m = 1$, da man für allgemeines m einfach komponentenweise argumentieren kann.

Sei $r > 0$ so klein, daß die Kugel mit Radius r und Mittelpunkt x_0 in Ω liegt. Sei $\xi \in \mathbb{R}^n$ ein Vektor mit $0 < |\xi| < r$. Für $j = 1, \dots, n$ setze

$$x^j := x_0 + \sum_{k=1}^j \xi_k e_k,$$

sodaß $x^0 = x_0$ und $x^n = x_0 + \xi$. Nach dem Mittelwertsatz (angewendet auf f in Abhängigkeit von der j -ten Variablen, wobei die anderen Variablen fest bleiben) gibt es für jedes $j = 1, \dots, n$ ein $\theta^j \in (0, 1)$, sodaß

$$f(x^j) - f(x^{j-1}) = \partial_j f(x^{j-1} + \theta^j \xi_j e_j) \xi_j.$$

Wir schreiben $y^j := x^{j-1} + \theta^j \xi_j e_j$ und summieren über alle $j = 1, \dots, n$:

$$f(x_0 + \xi) - f(x_0) = \sum_{j=1}^n \partial_j f(y^j) \xi_j = Df(x_0)\xi + \phi(\xi),$$

wo wir

$$\phi(\xi) := \sum_{j=1}^n \partial_j f(y^j) \xi_j - Df(x_0)\xi$$

gesetzt haben. Für dieses ϕ gilt aber

$$\begin{aligned} \lim_{\xi \rightarrow 0} \frac{|\phi(\xi)|}{|\xi|} &= \lim_{\xi \rightarrow 0} \frac{\sum_{j=1}^n (\partial_j f(y_j) - \partial_j f(x_0)) \xi_j}{|\xi|} \\ &\leq \lim_{\xi \rightarrow 0} \left(\sum_{j=1}^n |\partial_j f(y_j) - \partial_j f(x_0)|^2 \right)^{1/2} = 0. \end{aligned}$$

Hier haben wir im vorletzten Schritt die CAUCHY-SCHWARZ-Ungleichung und im Grenzübergang die Stetigkeit von $\partial_j f$ und die Tatsache $\lim_{\xi \rightarrow 0} y^j = x_0$ verwendet (die y^j hängen ja von ξ ab!). Damit ist die totale Differenzierbarkeit von f in x_0 gezeigt. $\square$

Wir lassen von nun an das Attribut ‚total‘ weg und sprechen nur noch von differenzierbaren Funktionen. Eine differenzierbare Funktion, deren JACOBI-Matrix stetig ist, nennen wir stetig differenzierbar.

7.2.2. Kettenregel.

SATZ 7.10 (Kettenregel). Seien $\Omega_1 \subset \mathbb{R}^n$, $\Omega_2 \subset \mathbb{R}^m$ offen und $f : \Omega_1 \rightarrow \mathbb{R}^m$, $g : \Omega_2 \rightarrow \mathbb{R}^k$ Abbildungen mit $f(\Omega_1) \subset \Omega_2$. Ist f im Punkt $x \in \Omega_1$ und g im Punkt $y := f(x) \in \Omega_2$ differenzierbar, so ist die Komposition $g \circ f : \Omega_1 \rightarrow \mathbb{R}^k$ im Punkt x differenzierbar, und es gilt

$$D(g \circ f)(x) = Dg(f(x))Df(x). \quad (7.3)$$

Beachte, daß das Matrixprodukt in (7.3) wohldefiniert ist: Es ist nämlich $Dg(y) \in \mathbb{R}^{k \times m}$ und $Df(x) \in \mathbb{R}^{m \times n}$, mithin $D(g \circ f)(x) \in \mathbb{R}^{k \times n}$.

BEWEIS. Nach Voraussetzung haben wir

$$f(x + \xi) = f(x) + Df(x)\xi + \phi(\xi), \quad g(y + \eta) = g(y) + Dg(y)\eta + \psi(\eta)$$

für alle hinreichend kleinen $\xi \in \mathbb{R}^n$ und $\eta \in \mathbb{R}^m$, wobei

$$\lim_{\xi \rightarrow 0} \frac{\phi(\xi)}{|\xi|} = 0, \quad \lim_{\eta \rightarrow 0} \frac{\psi(\eta)}{|\eta|} = 0.$$

Daher ist

$$\begin{aligned} (g \circ f)(x + \xi) &= g(f(x) + Df(x)\xi + \phi(\xi)) \\ &= g(f(x)) + Dg(f(x))(Df(x)\xi + \phi(\xi)) + \psi(Df(x)\xi + \phi(\xi)) \\ &= g(f(x)) + Dg(f(x))Df(x)\xi + Dg(f(x))\phi(\xi) + \psi(Df(x)\xi + \phi(\xi)), \end{aligned}$$

wobei wir im zweiten Schritt $\eta := Df(x)\xi + \phi(\xi)$ gesetzt haben. Man beachte dabei, daß mit $\xi \rightarrow 0$ auch $\eta \rightarrow 0$. Mit der Matrixnorm $|A| := \sup\{|A\xi| : \xi \in \mathbb{R}^n, |\xi| = 1\} < \infty$ gilt genauer

$$|\eta| \leq |Df(x)||\xi| + |\phi(\xi)|$$

und somit

$$\limsup_{\xi \rightarrow 0} \frac{|\psi(\eta)|}{|\xi|} = \limsup_{\xi \rightarrow 0} \frac{|\psi(\eta)|}{|\eta|} \frac{|\eta|}{|\xi|} \leq \limsup_{\xi \rightarrow 0} \frac{|\psi(\eta)|}{|\eta|} \left(|Df(x)| + \frac{|\phi(\xi)|}{|\xi|} \right) = 0.$$

Außerdem haben wir

$$\limsup_{\xi \rightarrow 0} \frac{|Dg(f(x))\phi(\xi)|}{|\xi|} \leq \limsup_{\xi \rightarrow 0} \frac{|Dg(f(x))||\phi(\xi)|}{|\xi|} = 0.$$

Insgesamt erhalten wir

$$(g \circ f)(x + \xi) = (g \circ f)(x) + Dg(f(x))Df(x)\xi + \chi(\xi)$$

mit $\lim_{\xi \rightarrow 0} \frac{\chi(\xi)}{|\xi|} = 0$, und die Behauptung folgt. $\square$

BEISPIEL 7.11. Betrachte eine Abbildung $\gamma: \mathbb{R} \rightarrow \mathbb{R}$ der Form

$$\gamma: t \mapsto F(t, x_1(t), x_2(t), \dots, x_\mu(t))$$

für differenzierbare Funktionen $F: \mathbb{R}^{\mu+1} \rightarrow \mathbb{R}$, $x = (x_1, \dots, x_\mu): \mathbb{R} \rightarrow \mathbb{R}^\mu$. Wir wollen $\gamma'(t)$ berechnen und verwenden dazu die Kettenregel mit $n = 1$, $m = \mu + 1$, $k = 1$,

$$f: \mathbb{R} \rightarrow \mathbb{R}^m, \quad f(t) = (t, x_1(t), x_2(t), \dots, x_\mu(t)),$$

und $g = F$. Dann ergibt die Kettenregel

$$\gamma'(t) = Dg(f(t)) \cdot f'(t) = \partial_t F(t, x(t)) + \sum_{l=1}^{\mu} \partial_{x_l} F(t, x(t)) x'_l(t).$$

Späterhin werden wir eine Anwendung dieser Formel in der klassischen Mechanik kennenlernen.

Als Anwendung erhalten wir die folgende höherdimensionale Version des Mittelwertsatzes:

SATZ 7.12. Sei $f: \Omega \rightarrow \mathbb{R}^m$ stetig differenzierbar, $x \in \Omega$, und $\xi \in \mathbb{R}^n$ dergestalt, daß die Strecke $\{x + t\xi: t \in [0, 1]\}$ in Ω liegt. Dann gilt

$$f(x + \xi) - f(x) = \left(\int_0^1 Df(x + t\xi) dt \right) \xi.$$

Hierbei meint das Integral diejenige Matrix, deren (i, j) -ter Eintrag gleich

$$\int_0^1 \partial_j f_i(x + t\xi) dt$$

ist.

BEWEIS. Für $i = 1, \dots, m$ setze $g_i(t) := f_i(x + t\xi)$. Dann ist nach Kettenregel g_i auf $[0, 1]$ differenzierbar mit Ableitung

$$g'_i(t) = Df_i(x + t\xi) \cdot \xi = \sum_{j=1}^n \partial_j f_i(x + t\xi) \xi_j,$$

was nach Voraussetzung der stetigen Differenzierbarkeit von f stetig ist (in t). Daher können wir den Hauptsatz der Differential- und Integralrechnung anwenden und erhalten wie behauptet

$$\begin{aligned} f_i(x + \xi) - f_i(x) &= g_i(1) - g_i(0) \\ &= \int_0^1 g'_i(t) dt \\ &= \int_0^1 \partial_j f_i(x + t\xi) \xi_j dt \\ &= \left(\int_0^1 Df(x + t\xi) dt \right) \xi. \end{aligned}$$

□

Als weitere Anwendung der Kettenregel werden wir Richtungsableitungen charakterisieren.

DEFINITION 7.13 (Richtungsableitung). Sei $f: \Omega \rightarrow \mathbb{R}^m$ in $x \in \Omega$ stetig differenzierbar und $e \in \mathbb{R}^n$ ein Vektor mit $|e| = 1$. Dann heißt

$$\lim_{h \rightarrow 0} \frac{f(x + he) - f(x)}{h}$$

die *Richtungsableitung* von f in Richtung e im Punkt x .

Die Definition suggeriert die Existenz der Richtungsableitung. Dies ist tatsächlich der Fall:

SATZ 7.14. Sei $f : \Omega \rightarrow \mathbb{R}^m$ stetig differenzierbar und $e \in \mathbb{R}^n$ ein Vektor mit $|e| = 1$. Dann existiert der Limes aus Definition 7.13 und ist gleich

$$Df(x)e.$$

BEWEIS. Definiere die Abbildung

$$g : \mathbb{R} \rightarrow \mathbb{R}^n, \quad h \mapsto x + he,$$

dann ist (da Ω offen!) für hinreichend kleines $\delta > 0$ das Bild von $(-\delta, \delta)$ unter g in Ω enthalten, und somit ist die Abbildung $F : (-\delta, \delta) \rightarrow \mathbb{R}^m$, $F = f \circ g$, wohldefiniert. Nun ist einerseits

$$F'(0) = \lim_{h \rightarrow 0} \frac{f(g(h)) - f(g(0))}{h} = \lim_{h \rightarrow 0} \frac{f(x + he) - f(x)}{h}$$

und andererseits nach Kettenregel

$$F'(0) = Df(g(0))g'(0) = Df(x)e,$$

und die Behauptung folgt. $\square$

Die geometrische Interpretation für $m = 1$ lautet wie folgt: Nach der Cauchy-Schwarz-Ungleichung gilt

$$|Df(x) \cdot e| \leq |Df(x)| |e|,$$

und die Wahl $e = \frac{Df(x)}{|Df(x)|}$ saturiert diese Ungleichung (d.h. sie ist dann mit Gleichheit erfüllt). Die Richtung des Gradienten ist also (falls $Df(x) \neq 0$) die Richtung des höchsten Anstiegs von f im Punkt x .

Um unsere geometrische Intuition für die Bedeutung des Gradienten weiter zu festigen, zeigen wir noch, daß dieser senkrecht auf den Niveaumengen steht. Sei nämlich $f : \Omega \rightarrow \mathbb{R}$ stetig differenzierbar, $c \in \mathbb{R}$, und $x \in \Omega$ ein Punkt mit $f(x) = c$. Sei weiter $\phi : (-\delta, \delta) \rightarrow \Omega$ eine stetig differenzierbare Abbildung mit $\phi(0) = x$ und $\phi((-\delta, \delta)) \subset N_f(c)$, wobei wie gehabt

$$N_f(c) = \{x \in \Omega : f(x) = c\}.$$

Dann ist der Vektor $\phi'(0)$ tangential zur Kurve ϕ am Punkt $\phi(0) = x$, und es gilt

$$Df(x) \cdot \phi'(0) = 0. \tag{7.4}$$

Dies folgt aus der Kettenregel: Die Komposition $f \circ \phi$ ist auf $(-\delta, \delta)$ wohldefiniert und konstant gleich c , da $\phi((-\delta, \delta)) \subset N_f(c)$. Deshalb ist ihre Ableitung in null gleich null:

$$0 = (f \circ \phi)'(0) = Df(\phi(0)) \cdot \phi'(0) = Df(x) \cdot \phi'(0),$$

wie behauptet.

7.3. Der Satz von TAYLOR und lokale Extrema

7.3.1. Multiindexnotation. In der höherdimensionalen Differentialrechnung ist es zweckmäßig, die folgenden Schreibweisen einzuführen. Sei $f : \Omega \rightarrow \mathbb{R}^m$, wobei wie immer $\Omega \subset \mathbb{R}^n$. Ein *Multiindex* ist ein Element $\alpha = (\alpha_1, \dots, \alpha_n) \in \mathbb{N}_0^n$. Wir setzen

$$\partial_\alpha f := \partial_1^{\alpha_1} \partial_2^{\alpha_2} \dots \partial_n^{\alpha_n} f,$$

d.h. die Zahl α_l gibt an, wie oft f nach der l -ten Variablen partiell differenziert wird. Wir schreiben weiter

$$|\alpha| := \sum_{l=1}^n \alpha_l$$

und bemerken, daß $\partial_\alpha f$ als stetige Funktion wohldefiniert ist, sofern f $|\alpha|$ -mal stetig differenzierbar ist. Schließlich setzen wir

$$\alpha! := \alpha_1! \alpha_2! \cdots \alpha_n!$$

sowie, für $x = (x_1, \dots, x_n) \in \mathbb{R}^n$,

$$x^\alpha := x_1^{\alpha_1} x_2^{\alpha_2} \cdots x_n^{\alpha_n}.$$

Ist beispielsweise $n = 3$ und $\alpha = (1, 0, 2)$, so ist $|\alpha| = 3$, und für eine dreimal stetig differenzierbare Funktion f gilt $\partial_\alpha f = \partial_1 \partial_{33} f$. Statt ∂_{33} können wir auch ∂_3^2 schreiben. Für dieses α ist außerdem $\alpha! = 1! \cdot 0! \cdot 2! = 2$.

7.3.2. Der Satz von TAYLOR. Wir benötigen ein vorbereitendes Resultat:

LEMMA 7.15. *Sei $f : \Omega \rightarrow \mathbb{R}$ k -mal stetig differenzierbar. Seien weiter $x \in \Omega$ und $\xi \in \mathbb{R}^n$ so gewählt, daß*

$$\{x + h\xi : h \in [0, 1]\} \subset \Omega.$$

Ist $g : [0, 1] \rightarrow \mathbb{R}$, $g(h) := f(x + h\xi)$, so ist g k -mal stetig differenzierbar, und es gilt

$$g^{(k)}(h) = \sum_{|\alpha|=k} \frac{k!}{\alpha!} \partial_\alpha f(x + h\xi) \xi^\alpha.$$

BEWEIS. Wir beweisen zunächst durch Induktion nach k die Formel

$$g^{(k)}(h) = \sum_{l_1, \dots, l_k=1}^n \partial_{l_1} \cdots \partial_{l_k} f(x + h\xi) \xi_{l_1} \cdots \xi_{l_k}. \quad (7.5)$$

Für $k = 1$ haben wir nach Kettenregel

$$g'(h) = Df(x + h\xi) \cdot \xi = \sum_{l=1}^n \partial_l f(x + h\xi) \xi_l.$$

Gilt die Behauptung für ein k , so ist wiederum nach Kettenregel

$$\begin{aligned} g^{(k+1)}(h) &= \frac{d}{dh} \left(\sum_{l_1, \dots, l_k=1}^n \partial_{l_1} \cdots \partial_{l_k} f(x + h\xi) \xi_{l_1} \cdots \xi_{l_k} \right) \\ &= \sum_{l_1, \dots, l_{k+1}=1}^n \partial_{l_1} \cdots \partial_{l_{k+1}} f(x + h\xi) \xi_{l_1} \cdots \xi_{l_{k+1}}, \end{aligned}$$

wobei wir den Satz von SCHWARZ über die Vertauschbarkeit partieller Ableitungen verwendet haben. Damit ist (7.5) gezeigt.

Nun läßt sich jeder Summand $\partial_{l_1} \cdots \partial_{l_k} f(x + h\xi) \xi_{l_1} \cdots \xi_{l_k}$ schreiben als $\partial_\alpha f(x + h\xi) \xi^\alpha$ für einen Multiindex α mit $|\alpha| = k$; allerdings gehören dann mehrere solcher Summanden zum selben Multiindex, nämlich diejenigen, für die die Anzahl der Indizes l_m , die den Wert $j \in \{1, \dots, n\}$ annehmen, für jedes j gleich ist.³ Genauer gibt es für jedes α mit $|\alpha| = k$

$$\frac{k!}{\alpha_1! \cdots \alpha_n!} = \frac{k!}{\alpha!}$$

Möglichkeiten, $\partial_\alpha f(x + h\xi) \xi^\alpha$ zu schreiben als $\partial_{l_1} \cdots \partial_{l_k} f(x + h\xi) \xi_{l_1} \cdots \xi_{l_k}$. Es folgt

$$\sum_{l_1, \dots, l_k=1}^n \partial_{l_1} \cdots \partial_{l_k} f(x + h\xi) \xi_{l_1} \cdots \xi_{l_k} = \sum_{|\alpha|=k} \frac{k!}{\alpha!} \partial_\alpha f(x + h\xi) \xi^\alpha$$

und daraus mit (7.5) die Behauptung. $\square$

³So sind etwa für $k = 3$, $n = 2$ die Ausdrücke $\partial_2 \partial_2 \partial_1 f \xi_2 \xi_2 \xi_1$ und $\partial_1 \partial_2 \partial_2 f \xi_1 \xi_2 \xi_2$ gleich und lassen sich als $\partial_\alpha f \xi^\alpha$ schreiben für $\alpha = (1, 2)$.

SATZ 7.16 (Satz von TAYLOR). Sei $f : \Omega \rightarrow \mathbb{R}$ $(k+1)$ -mal stetig differenzierbar sowie $x \in \Omega$ und $\xi \in \mathbb{R}^n$ so gewählt, daß

$$\{x + h\xi : h \in [0, 1]\} \subset \Omega.$$

Dann existiert ein $h \in [0, 1]$, sodaß

$$f(x + \xi) = \sum_{|\alpha|=0}^k \frac{\partial_\alpha f(x)}{\alpha!} \xi^\alpha + \sum_{|\alpha|=k+1} \frac{\partial_\alpha f(x + h\xi)}{\alpha!} \xi^\alpha.$$

BEWEIS. Setze $g : [0, 1] \rightarrow \mathbb{R}$, $g(h) := f(x + h\xi)$, so ist g nach Lemma 7.15 $(k+1)$ -mal stetig differenzierbar. Nach dem eindimensionalen Satz von TAYLOR (Korollar 4.49) existiert daher ein $h \in [0, 1]$, sodaß

$$g(1) = \sum_{m=0}^k \frac{g^{(m)}(0)}{m!} + \frac{g^{(k+1)}(h)}{(k+1)!}. \quad (7.6)$$

Wiederum nach Lemma 7.15 ist aber

$$g^{(m)}(0) = \sum_{|\alpha|=m} \frac{m!}{\alpha!} \partial_\alpha f(x) \xi^\alpha$$

sowie

$$g^{(k+1)}(h) = \sum_{|\alpha|=k+1} \frac{(k+1)!}{\alpha!} \partial_\alpha f(x + h\xi) \xi^\alpha.$$

Einsetzen in (7.6) ergibt unter Berücksichtigung von $g(1) = f(x + \xi)$ die Behauptung. $\square$

KOROLLAR 7.17. Sei $f : \Omega \rightarrow \mathbb{R}$ k -mal stetig differenzierbar sowie $x \in \Omega$ und $\delta > 0$ so gewählt, daß $B_\delta(x) \subset \Omega$. Dann gilt für alle $\xi \in \mathbb{R}^n$ mit $|\xi| < \delta$

$$f(x + \xi) = \sum_{|\alpha|=0}^k \frac{\partial_\alpha f(x)}{\alpha!} \xi^\alpha + \phi(\xi)$$

für ein $\phi : \mathbb{B}(0, \delta) \rightarrow \mathbb{R}$ mit der Eigenschaft

$$\lim_{\xi \rightarrow 0} \frac{\phi(\xi)}{|\xi|^k} = 0.$$

BEWEIS. Nach Satz 7.16 gibt es zu jedem ξ mit $|\xi| < \delta$ ein $h \in [0, 1]$, sodaß

$$\begin{aligned} f(x + \xi) &= \sum_{|\alpha|=0}^{k-1} \frac{\partial_\alpha f(x)}{\alpha!} \xi^\alpha + \sum_{|\alpha|=k} \frac{\partial_\alpha f(x + h\xi)}{\alpha!} \xi^\alpha \\ &= \sum_{|\alpha|=0}^k \frac{\partial_\alpha f(x)}{\alpha!} \xi^\alpha + \sum_{|\alpha|=k} \frac{\partial_\alpha f(x + h\xi) - \partial_\alpha f(x)}{\alpha!} \xi^\alpha \\ &=: \sum_{|\alpha|=0}^k \frac{\partial_\alpha f(x)}{\alpha!} \xi^\alpha + \phi(\xi). \end{aligned}$$

Wir haben aber

$$\frac{|\phi(\xi)|}{|\xi|^k} = \sum_{|\alpha|=k} \frac{|\partial_\alpha f(x + h\xi) - \partial_\alpha f(x)|}{\alpha!} \frac{|\xi^\alpha|}{|\xi|^k} \leq \sum_{|\alpha|=k} \frac{|\partial_\alpha f(x + h\xi) - \partial_\alpha f(x)|}{\alpha!} \rightarrow 0$$

mit $\xi \rightarrow 0$ aufgrund der Stetigkeit von $\partial_\alpha f$. $\square$

Wir betrachten noch die Spezialfälle $k = 1$ und $k = 2$: Ist $k = 1$, so ist $\alpha! = 1$ für jedes α mit $|\alpha| = 1$, und

$$\sum_{|\alpha|=0}^1 \frac{\partial_\alpha f(x)}{\alpha!} \xi^\alpha = f(x) + Df(x) \cdot \xi.$$

In diesem Falle gibt Korollar 7.17 also einfach die Definition der Differenzierbarkeit wieder.

Interessanter ist der Fall $k = 2$: Für die Terme zweiter Ordnung erhalten wir $\alpha! = 1$ für gemischte zweite Ableitungen und $\alpha! = 2$ für zweite Ableitungen der Form $\partial_{ll}f$, und somit

$$\sum_{|\alpha|=2} \frac{\partial_\alpha f(x)}{\alpha!} \xi^\alpha = \frac{1}{2} \sum_{m,l=1}^n \partial_{ml}f(x) \xi_m \xi_l = \frac{1}{2} (\xi, D^2 f(x) \xi)$$

(man beachte, daß für $m \neq l$ die Terme $\partial_{ml}f(x) \xi_m \xi_l$ und $\partial_{lm}f(x) \xi_l \xi_m$ einander gleichen und zum selben Multiindex α gehören, daher der Faktor $1/2$).

Die quadratische Approximation einer zweimal stetig differenzierbaren Funktion lautet also

$$f(x + \xi) \sim f(x) + Df(x) \cdot \xi + \frac{1}{2} (\xi, D^2 f(x) \xi)$$

für kleine ξ .

BEMERKUNG 7.18. Wir haben in Abschnitt 5.4 (im eindimensionalen Falle) gesehen, daß es Funktionen gibt, die sich innerhalb einer festen Umgebung eines bestimmten Punktes als Taylorreihe darstellen lassen. An einem festen Punkt innerhalb des Konvergenzbereiches konnte man somit den Funktionswert approximieren, indem man die Ordnung des Taylorpolynoms ausreichend groß wählte. In diesem Abschnitt ist die Idee eine andere: Man fixiert die Ordnung des Taylorpolynoms (entweder weil die Funktion nur endlich oft differenzierbar ist, oder weil der Rechenaufwand bei Approximation höherer Ordnung zu groß würde) und erhält eine Approximation, die desto besser ist, je näher der zu evaluierende Punkt am Entwicklungspunkt liegt.

7.3.3. Lokale Extrema.

DEFINITION 7.19. Eine Funktion $f : \Omega \rightarrow \mathbb{R}$ hat im Punkt $x_0 \in \Omega$ ein *lokales Maximum*, wenn es eine Umgebung U von x_0 gibt, sodaß

$$f(x) \leq f(x_0) \quad \text{für alle } x \in U.$$

Eine Funktion f hat in x_0 ein *lokales Minimum*, wenn $-f$ in x_0 ein lokales Maximum hat.

Der Oberbegriff zu Maximum und Minimum heißt *Extremum* (Plural *Extrema*).

Ziel dieses Unterabschnitts ist es, notwendige und auch hinreichende Bedingungen für das Vorliegen eines lokalen Extremums herzuleiten.

SATZ 7.20. Sei $f : \Omega \rightarrow \mathbb{R}$ partiell differenzierbar. Besitzt f in $x_0 \in \Omega$ ein lokales Extremum, so gilt

$$Df(x_0) = 0.$$

BEWEIS. Wir nehmen an, x_0 sei eine lokale Maximalstelle (der Fall des Minimums folgt dann durch Betrachten der Funktion $-f$), d.h. es existiert eine offene Menge $x_0 \in U \subset \Omega$, sodaß $f(x) \leq f(x_0)$ für alle $x \in U$. Für $l = 1, \dots, n$ ist dann insbesondere $f(x_0 + h e_l) \leq f(x_0)$ für hinreichend kleines $|h|$, d.h. die (nach Voraussetzung differenzierbare) Funktion

$$h \mapsto f(x_0 + h e_l)$$

hat ein lokales Maximum in $h = 0$. Aus der Analysis I wissen wir, daß die Ableitung dieser Funktion an der Stelle $h = 0$ verschwindet, also gilt $\partial_l f(x_0) = 0$. Da dies für jedes $l = 1, \dots, n$ der Fall ist, folgt die Behauptung. $\square$

Doch nicht an jedem Punkt, an dem $Df(x_0) = 0$ ist, muß ein Extremum vorliegen. Dies ist bereits im Eindimensionalen der Fall: Die Funktion $\mathbb{R} \rightarrow \mathbb{R}$, $x \mapsto x^3$ hat verschwindende Ableitung in $x = 0$, dort ist die Funktion allerdings weder minimal noch maximal. Um hinreichende Bedingungen für das Vorliegen einer Extremalstelle zu finden, müssen wir die zweiten Ableitungen, also die HESSE-Matrix, konsultieren. Dazu wiederholen wir ein paar Begriffe aus der Linearen Algebra:

DEFINITION 7.21. Sei $A \in \mathbb{R}^{n \times n}$ symmetrisch.

- (1) A heißt *positiv definit*, falls $(\xi, A\xi) > 0$ für jedes $\xi \in \mathbb{R}^n \setminus \{0\}$;
- (2) A heißt *positiv semidefinit*, falls $(\xi, A\xi) \geq 0$ für alle $\xi \in \mathbb{R}^n$;
- (3) A heißt *negativ (semi-)definit*, wenn $-A$ positiv (semi-)definit ist;
- (4) A heißt *indefinit*, wenn es $\xi, \eta \in \mathbb{R}^n$ gibt, sodaß $(\xi, A\xi) > 0$, aber $(\eta, A\eta) < 0$.

Wir schreiben auch $A > 0$ bzw. $A \geq 0$, falls A positiv (semi-)definit ist, und analog für negativ (semi-)definite Matrizen⁴.

Eine Charakterisierung der positiven bzw. negativen (Semi-)Definitheit einer Matrix liefern ihre Eigenwerte: Eine symmetrische reelle $(n \times n)$ -Matrix ist stets diagonalisierbar und besitzt n (nicht notwendig paarweise verschiedene) reelle Eigenwerte. Dann ist $A > 0$ genau dann, wenn alle Eigenwerte positiv sind, $A \geq 0$ genau dann, wenn alle Eigenwerte nichtnegativ sind, und analog $A < 0$ bzw. $A \leq 0$. Eine Matrix ist indefinit genau dann, wenn sie sowohl positive als auch negative Eigenwerte besitzt.

SATZ 7.22. Sei $f : \Omega \rightarrow \mathbb{R}$ zweimal stetig differenzierbar und sei $x_0 \in \Omega$ ein Punkt mit $Df(x_0) = 0$.

- (1) Ist $D^2f(x_0) > 0$, so hat f bei x_0 ein lokales Minimum;
- (2) Ist $D^2f(x_0) < 0$, so hat f bei x_0 ein lokales Maximum;
- (3) Ist $D^2f(x_0)$ indefinit, so ist x_0 weder lokale Maximal- noch Minimalstelle von f .

BEWEIS. Nach Korollar 7.17 und den darauf folgenden Ausführungen existiert ein $\phi : \Omega \rightarrow \mathbb{R}$ mit

$$\lim_{\xi \rightarrow 0} \frac{\phi(\xi)}{|\xi|^2} = 0$$

und

$$f(x_0 + \xi) = f(x_0) + \frac{1}{2}(\xi, D^2f(x_0)\xi) + \phi(\xi)$$

für hinreichend kleine $|\xi|$. Nehmen wir zunächst an $D^2f(x_0) > 0$. Wir zeigen, daß dann ein $\lambda > 0$ existiert mit

$$(\xi, D^2f(x_0)\xi) \geq \lambda|\xi|^2 \quad \text{für alle } \xi \in \mathbb{R}^n. \quad (7.7)$$

Betrachte nämlich die Menge $\mathbb{S}^{n-1} := \{\xi \in \mathbb{R}^n : |\xi| = 1\}$. Sie ist kompakt, und daher nimmt die stetige Funktion $\xi \mapsto (\xi, D^2f(x_0)\xi)$ auf $\mathbb{S}^{n-1}$ ihr Minimum

$$\lambda := \min_{\mathbb{S}^{n-1}} (\xi, D^2f(x_0)\xi)$$

an, welches wegen der positiven Definitheit von $D^2f(x_0)$ positiv ist. Für beliebiges $\xi \in \mathbb{R}^n$ ist aber $\xi = |\xi|\tilde{\xi}$ mit $\tilde{\xi} \in \mathbb{S}^{n-1}$, und somit gilt

$$(\xi, D^2f(x_0)\xi) = |\xi|^2(\tilde{\xi}, D^2f(x_0)\tilde{\xi}) \geq \lambda|\xi|^2,$$

womit (7.7) gezeigt ist.

Sei schließlich $\delta > 0$ so klein, daß einerseits $B_\delta(x_0) \subset \Omega$ und andererseits

$$\phi(\xi) \geq -\frac{1}{4}\lambda|\xi|^2 \quad \text{für alle } |\xi| < \delta.$$

⁴Setzt man noch $A \leq B$, falls $A - B \leq 0$, so definiert $\leq$ eine Halbordnung auf der Menge der symmetrischen reellen $(n \times n)$ -Matrizen.

Dann gilt für $|\xi| < \delta$

$$f(x_0 + \xi) = f(x_0) + \frac{1}{2}(\xi, D^2 f(x_0)\xi) + \phi(\xi) \geq f(x_0) + \frac{1}{2}\lambda|\xi|^2 - \frac{1}{4}\lambda|\xi|^2 \geq f(x_0),$$

und damit ist (1) gezeigt. Es folgt außerdem (2) durch Übergang zu $-f$.

Sei nun $D^2 f(x_0)$ indefinit und $\xi, \eta \in \mathbb{R}^n$ zwei Vektoren mit $\alpha := (\xi, A\xi) > 0$ und $\beta := (\eta, A\eta) < 0$. Ohne Einschränkung dürfen wir annehmen $|\xi| = |\eta| = 1$. Sei weiterhin $\delta > 0$, dann gilt für

$$\xi_\delta := \delta\xi, \quad \eta_\delta := \delta\eta$$

$(\xi_\delta, D^2 f(x_0)\xi_\delta) = \delta^2\alpha > 0$ und $(\eta_\delta, D^2 f(x_0)\eta_\delta) = \delta^2\beta < 0$. Ist schließlich δ so klein gewählt, daß

$$-\frac{1}{4}\alpha|\xi|^2 \leq \phi(\xi) \leq -\frac{1}{4}\beta|\xi|^2$$

für alle $|\xi| \leq \delta$, so gilt einerseits

$$f(x_0 + \xi_\delta) = f(x_0) + \frac{1}{2}(\xi_\delta, D^2 f(x_0)\xi_\delta) + \phi(\xi_\delta) \geq f(x_0) + \frac{1}{4}\alpha\delta^2 > f(x_0)$$

und andererseits

$$f(x_0 + \eta_\delta) = f(x_0) + \frac{1}{2}(\eta_\delta, D^2 f(x_0)\eta_\delta) + \phi(\eta_\delta) \leq f(x_0) + \frac{1}{4}\beta\delta^2 < f(x_0).$$

Wir haben also gezeigt, daß jede Umgebung von x_0 Punkte enthält, an denen f einen kleineren bzw. größeren Wert als $f(x_0)$ annimmt, und es folgt (3). $\square$

BEISPIEL 7.23. Sei eine symmetrische Matrix $A \in \mathbb{R}^{2 \times 2}$ gegeben, dann ist die HESSE-Matrix der Funktion $f: \mathbb{R}^2 \rightarrow \mathbb{R}$, $x \mapsto \frac{1}{2}(x, Ax)$, gegeben durch die konstante Matrix A . In der Tat, gemäß Beispiel nach Satz 7.8 lautet der Gradient $Df(x) = Ax$ (beachte nach Voraussetzung $A = A^t$), und daher ist

$$(D^2 f(x))_{kl} = \partial_k(Ax)_l = \partial_k \sum_{j=0}^n A_{lj}x_j = A_{lk}.$$

Beachte, daß $Df(0) = 0$.

Wir wählen nun spezielle Matrizen A und untersuchen, ob bei $x = 0$ jeweils ein lokales Extremum vorliegt. Ist z.B.

$$A = \begin{pmatrix} 2 & 0 \\ 0 & 1 \end{pmatrix},$$

so ist A positiv definit (die Eigenwerte 2 und 1 sind beide positiv), also hat f bei null ein lokales Minimum. Explizit lautet die Funktionsgleichung $f(x, y) = 2x^2 + y^2$, und der Graph von f ist ein nach oben geöffnetes Paraboloid. Die Höhenlinien $\{2x^2 + y^2 = c\}$ sind Ellipsen.

Betrachte nun die Matrix

$$A = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix},$$

die offensichtlich indefinit ist. Dann ist die Funktionsgleichung gegeben durch $f(x, y) = y^2 - x^2$ und f hat bei null kein Extremum, sondern einen sogenannten *Sattelpunkt*. Die Höhenlinien sind nunmehr Hyperbeln.

Ist die HESSE-Matrix an einem Punkt mit $Df(x_0) = 0$ (positiv oder negativ) semi-definit, so läßt sich keine Aussage über das Vorliegen eines lokalen Extremums treffen. Betrachte dazu die Funktionen

$$f(x, y) = x^4 + y^4, \quad g(x, y) = -x^4 - y^4, \quad h(x, y) = x^3,$$

deren Gradienten bei $(x, y) = 0$ sämtlich verschwinden. Ebenso sind die HESSE-Matrizen jeweils null (und insbesondere positiv und negativ semidefinit). Es hat f bei $(x, y) = 0$ ein lokales (sogar globales) Minimum, g ein Maximum, und h weder das eine noch das andere.

Es gilt die folgende Umkehrung von Satz 7.22:

PROPOSITION 7.24. *Sei $f : \Omega \rightarrow \mathbb{R}$ zweimal stetig differenzierbar und sei $x_0 \in \Omega$ ein Punkt mit $Df(x_0) = 0$.*

- (1) *Hat f bei x_0 ein lokales Maximum, so ist $D^2f(x_0) \leq 0$;*
- (2) *Hat f bei x_0 ein lokales Minimum, so ist $D^2f(x_0) \geq 0$.*

BEWEIS. Wir zeigen nur die erste Aussage, da die zweite wieder durch Übergang zu $-f$ folgt. Habe also f bei x_0 ein lokales Maximum. Wäre $D^2f(x_0)$ nicht negativ semidefinit, so gäbe es ein $\eta \in \mathbb{R}^n$ mit $|\eta| = 1$ und

$$(\eta, D^2f(x_0)\eta) =: \alpha > 0.$$

Nach Korollar 7.17 existiert $\phi : \Omega \rightarrow \mathbb{R}$ mit

$$\lim_{\xi \rightarrow 0} \frac{\phi(\xi)}{|\xi|^2} = 0$$

und

$$f(x_0 + \xi) = f(x_0) + \frac{1}{2}(\xi, D^2f(x_0)\xi) + \phi(\xi)$$

für kleine $|\xi|$. Ist $\delta > 0$ so klein gewählt, daß

$$\phi(\xi) \geq -\frac{1}{4}\alpha|\xi|^2 \quad \text{für alle } |\xi| \leq \delta,$$

so folgt

$$f(x_0 + \delta\eta) = f(x_0) + \frac{1}{2}\delta^2(\eta, D^2f(x_0)\eta) + \phi(\delta\eta) \geq f(x_0) + \frac{1}{2}\alpha\delta^2|\eta|^2 - \frac{1}{4}\alpha\delta^2|\eta|^2 > f(x_0),$$

also existiert in jeder Umgebung von x_0 ein Punkt, an dem f einen Wert größer als $f(x_0)$ annimmt. Dies ist der gewünschte Widerspruch dazu, daß x_0 Maximalstelle ist. $\square$

Diese Proposition nutzen wir nun, um ein wichtiges Resultat über *harmonische Funktionen* zu zeigen. Eine zweimal stetig differenzierbare Funktion $f : \Omega \rightarrow \mathbb{R}$ heißt *harmonisch*, wenn

$$\Delta f(x) = 0 \quad \text{für jedes } x \in \Omega. \quad (7.8)$$

Wir erinnern uns an den LAPLACE-Operator $\Delta = \sum_{l=1}^n \partial_{l,l}$. Die Gleichung $\Delta f = 0$ heißt dementsprechend LAPLACE-Gleichung und ist eine der wichtigsten partiellen Differentialgleichungen. Man kann sie z.B. zur Modellierung von Seifenhäuten oder statischen Wärmeverteilungen heranziehen. Das folgende Resultat ist der Prototyp eines *Maximumsprinzips*, wie es in der Theorie partieller Differentialgleichungen größte Bedeutung besitzt:

SATZ 7.25 (Schwach Maximumsprinzip für harmonische Funktionen). Sei $\Omega \subset \mathbb{R}^n$ offen und beschränkt. Sei außerdem $f : \bar{\Omega} \rightarrow \mathbb{R}$ stetig in $\bar{\Omega}$ und zweimal stetig differenzierbar in Ω^5 . Gilt

$$\Delta f(x) = 0 \quad \text{für jedes } x \in \Omega, \quad (7.9)$$

so nimmt f sein Maximum und sein Minimum jeweils auf dem Rand $\partial\Omega$ an, d.h.

$$\max_{\bar{\Omega}} f = \max_{\partial\Omega} f, \quad \min_{\bar{\Omega}} f = \min_{\partial\Omega} f.$$

⁵Das bedeutet, daß f stetig *bis zum Rand* ist. Die Funktion $f(x, y) = \frac{1}{1-x^2-y^2}$ ist zum Beispiel zweimal stetig differenzierbar in $B_1(0)$, aber nicht stetig fortsetzbar nach $\overline{B_1(0)}$.

BEWEIS. Wir zeigen nur die Aussage über das Maximum, da die über das Minimum durch Übergang zu $-f$ folgt. Zuerst bemerken wir, daß f als stetige Funktion auf der kompakten (da beschränkten und abgeschlossenen) Menge $\bar{\Omega}$ tatsächlich sein Maximum annimmt.

Schritt 1. Wir nehmen zunächst an, daß

$$\Delta f(x) > 0 \quad \text{für jedes } x \in \Omega.$$

Nähme f sein Maximum im Punkte $x_0 \in \Omega$ an, so wäre x_0 insbesondere eine lokale Maximalstelle, und es gälte dort nach Proposition 7.24 $D^2 f(x_0) \leq 0$. Aus der Linearen Algebra ist bekannt, daß die Spur einer symmetrischen Matrix gleich der Summe ihrer Eigenwerte ist. Aus $D^2 f(x_0) \leq 0$ folgt die Nichtpositivität aller Eigenwerte, und es folgt

$$\Delta f(x_0) = \text{spur } D^2 f(x_0) \leq 0,$$

Widerspruch. Es folgt $\max_{\bar{\Omega}} f = \max_{\partial\Omega} f$.

Schritt 2. Sei nun $\Delta f = 0$ in Ω . Wir führen diesen Fall auf Schritt 1 zurück, indem wir für $\epsilon > 0$

$$f^\epsilon(x) := f(x) + \epsilon \frac{1}{2} |x|^2$$

setzen. Man prüft leicht nach $\Delta f^\epsilon(x) = n\epsilon > 0$ für alle $x \in \Omega$, und daher gilt nach Schritt 1 für alle $x \in \Omega$

$$f^\epsilon(x) \leq \max_{\partial\Omega} f^\epsilon \leq \max_{\partial\Omega} f + C\epsilon,$$

wobei $C := \max\{\frac{1}{2}|x|^2 : x \in \bar{\Omega}\} < \infty$ (da Ω beschränkt ist), und insbesondere

$$f(x) \leq \max_{\partial\Omega} f^\epsilon \leq \max_{\partial\Omega} f + C\epsilon,$$

da $f \leq f^\epsilon$. Da dies für jedes $\epsilon > 0$ gilt, haben wir sogar

$$f(x) \leq \max_{\partial\Omega} f$$

für alle $x \in \Omega$, und die Behauptung folgt. $\square$

KOROLLAR 7.26 (Eindeutigkeit des DIRICHLET-Problems). Sei $\Omega \subset \mathbb{R}^n$ offen und beschränkt, und seien $f_1, f_2 \in C(\bar{\Omega})$ beide zweimal stetig differenzierbar in Ω . Seien weiterhin $h \in C(\Omega)$ und $g \in C(\partial\Omega)$. Gilt für $j = 1, 2$

$$\begin{aligned} \Delta f_j(x) &= h(x) \quad \text{für alle } x \in \Omega, \\ f_j(x) &= g(x) \quad \text{für alle } x \in \partial\Omega, \end{aligned} \tag{7.10}$$

so gilt $f_1 = f_2$ in $\bar{\Omega}$.

BEWEIS. Sei $\bar{f} := f_1 - f_2$, dann erfüllt $\bar{f}$

$$\begin{aligned} \Delta \bar{f}(x) &= 0 \quad \text{für alle } x \in \Omega, \\ \bar{f}(x) &= 0 \quad \text{für alle } x \in \partial\Omega. \end{aligned}$$

Nach Satz 7.25 gilt für alle $x \in \Omega$ $\bar{f}(x) \leq \max_{\partial\Omega} \bar{f} = 0$, aber auch $\bar{f}(x) \geq \min_{\partial\Omega} \bar{f} = 0$, also ist $\bar{f}$ identisch null in $\bar{\Omega}$. Es folgt die Behauptung. $\square$

Das Randwertproblem (7.10) (bekannt als DIRICHLET-*Problem* für die POISSON-Gleichung) besitzt also *höchstens* eine Lösung. Für die Frage, ob überhaupt eine Lösung existiert, verweisen wir auf Vorlesungen über partielle Differentialgleichungen⁶.

⁶Die Antwort auf diese Frage lautet für „anständige“ Gebiete Ω und Inhomogenitäten h „ja“.

7.4. Implizite Funktionen und lokale Invertierbarkeit

7.4.1. Der Satz über implizite Funktionen. Im Gegensatz zu einer explizit in der Form $y = f(x)$ gegebenen Funktion betrachten wir nun *implizit* definierte Funktionen der Form $F(x, y) = 0$. Die Frage ist dann, ob man diese Gleichung nach y auflösen kann, ob es also ausgehend von $F(x, y) = 0$ eine Darstellung $y = f(x)$ gibt.

Ein uns bereits bekanntes Beispiel ist $F(x, y) = x^2 + y^2 - 1$: Die Menge $\{(x, y) \in \mathbb{R}^2 : F(x, y) = 0\}$ ist genau der Einheitskreis. Dieser ist *nicht* der Graph einer Funktion, da jedem x mit $|x| < 1$ zwei y -Werte zugeordnet werden müßten, nämlich $\pm\sqrt{1-x^2}$. Andererseits läßt sich der Einheitskreis zumindest *lokal* in der Umgebung jedes Punktes (x_0, y_0) mit $x_0^2 + y_0^2 = 1$ und $|x_0| < 1$ als Graph darstellen: für $y_0 > 0$ liegt (x_0, y_0) auf dem Graphen von $f_+(x) = \sqrt{1-x^2}$, für $y_0 < 0$ auf dem von $f_-(x) = -\sqrt{1-x^2}$. Lediglich an den Stellen $(\pm 1, 0)$ kann man nicht nach y auflösen; beachte, daß genau an diesen beiden Stellen gilt $\partial_y F(x, y) = 0$.

Im allgemeinen wird man eine implizite Gleichung nicht in geschlossener Darstellung nach der gewünschten Variablen auflösen können (denken Sie z.B. an $y(1+y) - x \exp(xy) = 0$). Trotzdem haben wir

SATZ 7.27 (Satz über implizite Funktionen). Seien $\Omega_1 \subset \mathbb{R}^n$ und $\Omega_2 \subset \mathbb{R}^m$ offen und $F : \Omega_1 \times \Omega_2 \rightarrow \mathbb{R}^m$ stetig differenzierbar. Sei $(x_0, y_0) \in \Omega_1 \times \Omega_2$ ein Punkt mit $F(x_0, y_0) = 0$ und mit

$$\det D_y F(x_0, y_0) \neq 0.$$

Dann existieren Umgebungen $\Omega_1 \supset U_1 \ni x_0$ und $\Omega_2 \supset U_2 \ni y_0$ und eine stetig differenzierbare Funktion $f : U_1 \rightarrow U_2$, sodaß für alle $(x, y) \in U_1 \times U_2$ gilt

$$F(x, y) = 0 \quad \text{genau dann, wenn} \quad y = f(x).$$

Für die JACOBI-Matrix von f gilt

$$Df(x) = -(D_y F(x, f(x)))^{-1} D_x F(x, f(x)). \quad (7.11)$$

BEMERKUNG 7.28. Hier haben wir mit $D_y F(x, y)$ die JACOBI-Matrix $(m \times m)$ der Funktion $y \mapsto F(x, y)$ bei festem x bezeichnet, d.h. die Matrix bestehend aus den partiellen Ableitungen der Komponenten von F in Richtung $y_1, \dots, y_m$.

Außerdem werden wir im Beweis die bereits (im Beweis der Kettenregel) eingeführte Matrixnorm $|A| := \sup\{|A\xi| : \xi \in \mathbb{R}^m, |\xi| = 1\} < \infty$ verwenden, für die offensichtlich $|A\xi| \leq |A||\xi|$ für alle $\xi \in \mathbb{R}^m$ gilt.

BEWEIS. *Schritt 1.* Wir nehmen ohne Einschränkung an $(x_0, y_0) = 0$, ansonsten betrachte die Funktion $F^0(x, y) := F(x + x_0, y + y_0)$. Schreibe außerdem

$$B := D_y F(0, 0),$$

dann ist B nach Voraussetzung invertierbar, und wir können definieren

$$G : \Omega_1 \times \Omega_2 \rightarrow \mathbb{R}^m, \quad G(x, y) := y - B^{-1}F(x, y).$$

Dann ist $F(x, y) = 0$ äquivalent zu $G(x, y) = y$. Außerdem ist G stetig differenzierbar mit

$$D_y G(0, 0) = I - B^{-1}D_y F(0, 0) = 0,$$

und da mit F auch G stetig differenzierbar ist, existieren Umgebungen $V_1 \ni 0$, $V_2 \ni 0$, sodaß

$$|D_y G(x, y)| < \frac{1}{2} \quad \forall (x, y) \in V_1 \times V_2. \quad (7.12)$$

Ohne Einschränkung dürfen wir annehmen, daß V_1 beschränkt ist und

$$V_2 = B_r(0) \quad \text{für ein } r > 0.$$

Schritt 2. Wir zeigen: Es gibt Umgebungen $V_1 \supset U_1 \ni 0$ und $V_2 \supset U_2 \ni 0$, sodaß einerseits

$$|G(x, y_2) - G(x, y_1)| \leq \frac{1}{2}|y_2 - y_1| \quad \forall x \in U_1, \quad y_1, y_2 \in U_2, \quad (7.13)$$

und andererseits

$$|G(x, y)| < r \quad \forall x \in U_1, \quad y \in U_2. \quad (7.14)$$

Wir setzen dazu einfach $U_2 := V_2 = B_r(0)$ und erhalten mit dem Mittelwertsatz für Funktionen mehrerer Veränderlicher (Satz 7.12) für beliebige $x \in V_1$, $y_1, y_2 \in U_2$

$$\begin{aligned} |G(x, y_2) - G(x, y_1)| &\leq \int_0^1 |D_y G(x, y_1 + t(y_2 - y_1))| dt |y_2 - y_1| \\ &\leq \frac{1}{2}|y_2 - y_1|, \end{aligned}$$

wo wir im letzten Schritt (7.12) verwendet haben. Damit ist (7.13) gezeigt. Für (7.14) bemerken wir, daß als Spezialfall von (7.13) (mit $x = 0$, $y_1 = 0$) folgt

$$|G(0, y)| \leq \frac{1}{2}|y| < \frac{1}{2}r \quad \forall y \in U_2.$$

Da aber G auf der kompakten Menge $\overline{V_1} \times \overline{U_2}$ gleichmäßig stetig ist, gibt es ein $\delta > 0$, sodaß $\mathbb{B}(0, \delta) \subset V_1$ und

$$|G(x, y)| \leq |G(x, y) - G(0, y)| + |G(0, y)| < \frac{r}{2} + \frac{r}{2} \quad \forall x \in \mathbb{B}(0, \delta) \times U_2.$$

Damit ist (7.14) mit $U_1 := B_\delta(0)$ gezeigt.

Schritt 3. Betrachte den Raum X der beschränkten stetigen Funktionen $f : U_1 \rightarrow \overline{U_2}$ mit $f(0) = 0$. Nach Beispiel 6.33 ist X bezüglich der Metrik d_∞ vollständig⁷. Wir definieren

$$T : X \rightarrow X, \quad T(f)(x) = G(x, f(x)).$$

Zuerst müssen wir prüfen, daß T tatsächlich nach X abbildet. Klar ist, daß $T(f)$ stetig ist und $T(f)(0) = G(0, f(0)) = G(0, 0) = 0$. Außerdem haben wir für alle $x \in U_1$

$$|T(f)(x)| = |G(x, f(x))| \leq r$$

wegen (7.14) und $f(x) \in \overline{U_2}$. Also ist in der Tat $T(f)(x) \in \overline{U_2}$ für alle $x \in U_1$.

Außerdem ist T eine Kontraktion, denn nach (7.13) gilt

$$d_\infty(T(f) - T(g)) = \sup_{x \in U_1} |G(x, f(x)) - G(x, g(x))| \leq \frac{1}{2}d_\infty(f, g).$$

Die Anwendung des BANACHschen Fixpunktsatzes liefert also eine eindeutige stetige Funktion $f : U_1 \rightarrow \overline{U_2}$, für die $G(x, f(x)) = f(x)$ für alle $x \in U_1$, für die also $F(x, f(x)) = 0$ für alle $x \in U_1$. Nach eventueller Vergrößerung von U_2 dürfen wir sogar annehmen, daß $f : U_1 \rightarrow U_2$.

Schritt 4. Wir zeigen, daß $x, y \in U_1 \times U_2$ und $F(x, y) = 0$ sogar $y = f(x)$ impliziert (bisher haben wir nur gezeigt, daß $y = f(x)$ die implizite Gleichung $F(x, y) = 0$ impliziert). Gäbe es nämlich für ein $x \in U_1$ neben $y_1 := f(x)$ noch ein weiteres $y_2 \in U_2$ mit $F(x, y_2) = 0$, so wäre also $G(x, y_1) = y_1$, $G(x, y_2) = y_2$, und damit

$$|G(x, y_1) - G(x, y_2)| = |y_1 - y_2|,$$

und aus (7.13) folgt dann $|y_1 - y_2| \leq \frac{1}{2}|y_1 - y_2|$, d.h. $y_1 = y_2$.

Schritt 5. Wir zeigen, daß die soeben konstruierte Funktion f sogar stetig differenzierbar in U_1 ist, und daß (7.11) gilt. Weiterhin sei $(x_0, y_0) = (0, 0)$. Schreibe dazu wie gehabt

⁷Erinnerung: $d_\infty(f, g) := \sup_{x \in U_1} |f(x) - g(x)|$. Man beachte hierbei, daß die Bedingung $f(0) = 0$ unter gleichmäßiger Konvergenz erhalten bleibt.

$B = \partial_y F(0, 0)$ sowie $A := \partial_x F(0, 0)$, dann gilt nach Definition der (totalen) Differenzierbarkeit unter Beachtung von $F(x, f(x)) = 0$

$$0 = Ax + Bf(x) + \phi(x, f(x)), \quad (7.15)$$

wobei $\phi : U_1 \times U_2 \rightarrow \mathbb{R}^m$ eine Funktion ist mit

$$\lim_{x, y \rightarrow 0} \frac{\phi(x, y)}{\sqrt{|x|^2 + |y|^2}} = 0. \quad (7.16)$$

Umstellen von (7.15) ergibt

$$f(x) = -B^{-1}Ax - B^{-1}\phi(x, f(x)), \quad (7.17)$$

und die Differenzierbarkeit von f an der Stelle $(0, 0)$ ist bewiesen, sobald wir zeigen können

$$\lim_{x \rightarrow 0} \frac{B^{-1}\phi(x, f(x))}{|x|} = 0. \quad (7.18)$$

Zunächst folgt aus (7.17) die Abschätzung

$$|f(x)| \leq |B^{-1}A||x| + |B^{-1}\phi(x, f(x))|, \quad (7.19)$$

und nach (7.16) dürfen wir (nach eventueller Verkleinerung von U_1, U_2) annehmen

$$|\phi(x, y)| \leq \frac{1}{2|B^{-1}|} \sqrt{|x|^2 + |y|^2} \leq \frac{1}{2|B^{-1}|} (|x| + |y|),$$

wobei wir in der letzten Ungleichung verwendet haben $a^2 + b^2 \leq (a+b)^2$ für $a, b \geq 0$. Einsetzen in (7.19) ergibt

$$|f(x)| \leq C|x| + \frac{1}{2}|f(x)|$$

für eine Konstante C , mithin

$$|f(x)| \leq 2C|x|. \quad (7.20)$$

Sei nun $\epsilon > 0$ und $\delta > 0$ so klein, daß

$$|\phi(x, y)| \leq \epsilon \sqrt{|x|^2 + |y|^2} \leq \epsilon(|x| + |y|) \quad (7.21)$$

für alle $x, y \in U_1 \times U_2$ mit $|x| + |y| < \delta$.

Ist dann $|x| < \frac{\delta}{2C+1}$, so ist $|x| + |f(x)| < \delta$, und es gilt mit (7.21) und (7.20)

$$\frac{|B^{-1}\phi(x, f(x))|}{|x|} \leq |B^{-1}| \epsilon \frac{(|x| + |f(x)|)}{|x|} \leq |B^{-1}| \epsilon (1 + 2C),$$

und damit ist (7.18) gezeigt.

Schließlich dürfen wir annehmen, daß $D_y F(x, f(x))$ nach eventueller Verkleinerung von U_1 invertierbar ist für jedes $x \in U_1$ (dies folgt aus der Stetigkeit von $D_y F$ und der Stetigkeit der Determinante). Dann können wir Schritt 5 am Punkt $(x_0, y_0) = (x, f(x))$ anwenden und erhalten so die Differenzierbarkeit von f in x und Formel (7.11). Da die partiellen Ableitungen von F stetig sind, folgt aus (7.11) zudem die Stetigkeit von Df . $\square$

BEMERKUNG 7.29. Die Formel (7.11) kann man sich mithilfe der Kettenregel herleiten: Differentiation der Identität $F(x, f(x)) = 0$ ergibt

$$\partial_x F(x, f(x)) + \partial_y F(x, f(x)) Df(x) = 0,$$

und nach Umstellen

$$Df(x_0) = -(D_y F(x, f(x)))^{-1} D_x F(x, f(x)).$$

BEISPIEL 7.30. Für $F(x, y) = x^2 + y^2 - 1$ haben wir $\partial_y F(x, y) = 2y$, und dies ist ungleich null genau dann, wenn $y \neq 0$. Es ist z.B. $F(0, 1) = 0$, und der Satz über implizite Funktionen liefert eine stetig differenzierbare Funktion $y = f(x)$ in einer Umgebung von $x = 0$, sodaß $x^2 + f(x)^2 = 1$ für alle x in dieser Umgebung. Die Ableitung errechnet sich mit (7.11) zu

$$f'(x) = -\frac{\partial_x F(x, f(x))}{\partial_y F(x, y)} = -\frac{x}{f(x)};$$

insbesondere erhalten wir $f'(0) = 0$. Diese Beobachtungen sind konsistent mit unseren Erörterungen zu Beginn dieses Abschnitts, wo wir $f(x) = \sqrt{1 - x^2}$ ermittelt haben.

BEISPIEL 7.31. Sei $F(x, y) = y(1+y) - x \exp(xy)$ (definiert auf ganz $\mathbb{R}^2$). Es ist $F(0, 0) = 0$ und $\partial_y F(0, 0) = 1$. Also existiert in einer Umgebung von null eine explizite Darstellung $y = f(x)$ der impliziten Relation $F(x, y) = 0$. Für die Ableitung erhalten wir

$$f'(0) = -\frac{\partial_x F(0, 0)}{\partial_y F(0, 0)} = -\frac{1}{1} = -1.$$

BEISPIEL 7.32 (Höhenlinien). Die Niveaumenge einer stetig differenzierbaren Funktion $F : \Omega \subset \mathbb{R}^2 \rightarrow \mathbb{R}$ zum Wert $c \in \mathbb{R}$ ist gegeben durch

$$N_c(f) = \{x \in \Omega : f(x, y) = c\}.$$

Ist $Df(x_0, y_0) \neq 0$, so ist $\partial_x f(x_0, y_0) \neq 0$ oder $\partial_y f(x_0, y_0) \neq 0$, also ist $N_c(f)$ nach dem Satz über implizite Funktionen (evtl. mit Vertauschung der Rollen von x und y) in einer Umgebung von (x_0, y_0) der Graph einer stetig differenzierbaren Funktion $x = h(y)$ oder $y = g(x)$. Natürlich kann beides gleichzeitig der Fall sein (nämlich wenn keine der Komponenten von $Df(x_0, y_0)$ null ist), dann ist $h = g^{-1}$. Diese Beobachtungen rechtfertigen die Bezeichnung *Niveaulinie* bzw. *Höhenlinie*.

7.4.2. Lokale Invertierbarkeit. Eng verwandt mit der Frage nach der Auflösbarkeit impliziter Gleichungen ist die nach der Invertierbarkeit einer Abbildung $f : \Omega_1 \subset \mathbb{R}^n \rightarrow \mathbb{R}^n$. Gibt es also eine Umkehrabbildung $f^{-1} : \Omega_2 \rightarrow \mathbb{R}^n$, sodaß $f \circ f^{-1}$ die Identität auf Ω_2 und $f^{-1} \circ f$ die Identität auf Ω_1 ist? Für *lineare* Funktionen $f(x) = Ax$, $A \in \mathbb{R}^{n \times n}$, ist die Antwort aus der linearen Algebra bekannt: Eine Umkehrabbildung existiert genau dann, wenn $\det A \neq 0$, und ist in diesem Fall durch $f^{-1}(y) = A^{-1}y$ gegeben.

Falls f nichtlinear, aber stetig differenzierbar ist, kann man f in einer Umgebung eines Punktes linear approximieren und würde dann erwarten, daß es sich nahe dieses Punktes wie die lineare Approximation verhält. Insbesondere sollte f in einer Umgebung von x_0 invertierbar sein, sofern $\det Df(x_0) \neq 0$. Dies ist genau die Aussage des nächsten Satzes:

SATZ 7.33 (Lokale Invertierbarkeit). Sei $\Omega \subset \mathbb{R}^n$ offen und $f : \Omega \rightarrow \mathbb{R}^n$ stetig differenzierbar. Ist $x_0 \in \Omega$ ein Punkt mit $\det Df(x_0) \neq 0$, so existieren Umgebungen $U_1 \ni x_0$ und $U_2 \ni f(x_0) := y_0$, sodaß $f : U_1 \rightarrow U_2$ bijektiv ist. Die Umkehrabbildung $f^{-1} : U_2 \rightarrow U_1$ ist stetig differenzierbar mit

$$Df^{-1}(y) = Df(f^{-1}(y))^{-1} \quad \forall y \in U_2. \quad (7.22)$$

BEWEIS. Betrachte die stetig differenzierbare Funktion $F : \mathbb{R}^n \times \Omega \rightarrow \mathbb{R}^n$, $F(y, x) = y - f(x)$. Wegen $y_0 = f(x_0)$ ist $F(y_0, x_0) = 0$, und $D_x F(y_0, x_0) = -Df(x_0)$ ist nach Voraussetzung invertierbar. Nach dem Satz über implizite Funktionen existieren daher offene Umgebungen $V_1 \ni x_0$, $U_2 \ni y_0$ und eine stetig differenzierbare Funktion $g : U_2 \rightarrow V_1$ mit $x = g(y)$ genau dann, wenn $F(y, x) = 0$, d.h. wenn $y = f(x)$. Setze nun $U_1 := V_1 \cap f^{-1}(U_2)$, dann ist U_1 offen, und $f : U_1 \rightarrow U_2$ ist bijektiv: f ist surjektiv, denn für $y \in U_2$ ist $f(g(y)) = f(x) = y$, und $x = g(y) \in V_1$ sowie $x \in f^{-1}(U_2)$; und f ist injektiv, denn aus $x_1, x_2 \in U_1$ folgt $f(x_1), f(x_2) \in U_2$, und aus $y := f(x_1) = f(x_2) \in U_2$ folgt $g(y) = x_1 = x_2$. Es ist klar, daß $g = f^{-1}$.

Die Formel für die Ableitung der Umkehrfunktion ergibt sich schließlich aus der Kettenregel, angewandt auf $y \mapsto f(f^{-1}(y)) = y$:

$$I = Df(f^{-1}(y))Df^{-1}(y),$$

was nach Multiplikation beider Seiten von links mit $Df(f^{-1}(y))^{-1}$ (7.22) ergibt. $\square$

BEMERKUNG 7.34. Für $n = 1$ ist (7.22) die bekannte Formel für die Ableitung der Umkehrfunktion in einer Variablen:

$$(f^{-1})'(y) = \frac{1}{f'(f^{-1}(y))}.$$

Man kann den Satz über die lokale Invertierbarkeit anwenden, um nichtlineare Gleichungssysteme zu lösen⁸:

BEISPIEL 7.35. Betrachte das nichtlineare Gleichungssystem

$$\begin{aligned} x(y+1)z - y \exp(z) &= a \\ 2 \sin(x) - xy^2 + yz &= b \\ x^2 - y^2 + 5z &= c. \end{aligned} \tag{7.23}$$

Setzen wir $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$,

$$f(x, y, z) = \begin{pmatrix} x(y+1)z - y \exp(z) \\ 2 \sin(x) - xy^2 + yz \\ x^2 - y^2 + 5z \end{pmatrix},$$

so ist offenbar $f(0, 0, 0) = 0$, und

$$Df(0, 0, 0) = \begin{pmatrix} 0 & -1 & 0 \\ 2 & 0 & 0 \\ 0 & 0 & 5 \end{pmatrix},$$

was invertierbar ist. Nach dem Satz über die lokale Inverse existieren also Umgebungen U_1 und U_2 von $(0, 0, 0)$, sodaß $f: U_1 \rightarrow U_2$ invertierbar ist. Das bedeutet: Für betragsmäßig hinreichend kleine $a, b, c \in \mathbb{R}$ hat das Gleichungssystem (7.23) eine eindeutige Lösung.

Die Formel für die Ableitung der Umkehrfunktion erlaubt uns zudem, die *Sensitivität* der Lösung von den Daten (a, b, c) zu messen: Nach (7.22) ist nämlich

$$Df^{-1}(a, b, c) = Df(x, y, z)^{-1},$$

wobei (x, y, z) die eindeutig bestimmte Lösung von (7.23) mit rechter Seite (a, b, c) ist. Bei null gilt also

$$Df^{-1}(0, 0, 0) = \begin{pmatrix} 0 & \frac{1}{2} & 0 \\ -1 & 0 & 0 \\ 0 & 0 & \frac{1}{5} \end{pmatrix};$$

die Interpretation lautet: In linearer Näherung hängt die Lösung x nicht von a und c ab, wohl aber von b (ein Anstieg von b um db bewirkt einen Anstieg von x um $db/2$). Die Lösungskomponente y fällt mit a und hängt (in linearer Näherung) nicht von b und c ab; und z hängt nur signifikant von c ab und verkleinert Änderungen in c um den Faktor $1/5$.

⁸„Lösen“ heißt hier: die Existenz einer (eindeutigen) Lösung zeigen, oder allenfalls ein Näherungsverfahren zur Lösung angeben. Eine Darstellung der Lösung in geschlossener Form ist typischerweise weder möglich noch wichtig.

KAPITEL 8

Integralrechnung in $\mathbb{R}^n$

So wie in einer Dimension das Integral die Fläche unter einem Graphen angibt, kann man das Volumen unter dem Graphen einer Funktion $\mathbb{R}^2 \supset \Omega \rightarrow \mathbb{R}$ mithilfe eines *Doppelintegrals* bestimmen. Die Integration bezüglich mehrerer Variabler gestattet auch die Berechnung von Volumina bekannter geometrischer Objekte, etwa von Kugeln, Ellipsoiden, Zylindern und Kegeln.

Als ‚Nebenprodukt‘ unserer Entwicklung des Integralbegriffs erhalten wir außerdem die Vertauschbarkeit der Integrationsreihenfolge (ein Vorläufer des Satzes von FUBINI) sowie die Vertauschbarkeit von Differentiation und Integration bezüglich verschiedener Variabler. Als Exkurs erhalten wir so die EULER-LAGRANGE-Gleichungen der Variationsrechnung.

Die hier vorgestellte Integrationstheorie ist im Wesentlichen nur auf stetige Integranden anwendbar. Eine viel allgemeinere Integrationstheorie lernen Sie im nächsten Semester in der Maßtheorie kennen.

8.1. Mehrfachintegrale

8.1.1. Differentiation bezüglich eines Parameters. Sei $f : [a, b] \times \Omega \rightarrow \mathbb{R}$ eine stetige Funktion, wobei $[a, b] \subset \mathbb{R}$ ein kompaktes Intervall und $\Omega \subset \mathbb{R}^{n-1}$ eine Teilmenge ist. Für jedes $x' \in \Omega$ ist dann die Funktion

$$[a, b] \rightarrow \mathbb{R}, \quad x \mapsto f(x, x')$$

stetig und somit integrierbar. Die Funktion

$$F : \Omega \rightarrow \mathbb{R}, \quad x' \mapsto \int_a^b f(x, x') dx$$

ist also wohldefiniert. Wir interessieren uns für die Eigenschaften dieser Funktion, insbesondere für ihre Stetigkeit und Differenzierbarkeit (sofern f selbst differenzierbar ist).

SATZ 8.1 (Stetige Abhängigkeit des Integrals von Parametern). Sei $f : [a, b] \times \Omega \rightarrow \mathbb{R}$ stetig, dann ist

$$F : \Omega \rightarrow \mathbb{R}, \quad x' \mapsto \int_a^b f(x, x') dx$$

ebenfalls stetig.

BEWEIS. Sei $(x'_k)_{k \in \mathbb{N}}$ eine Folge in Ω , die gegen $x' \in \Omega$ konvergiert. Wir werden zeigen, daß die Folge $(f_k)_{k \in \mathbb{N}}$ von Funktionen $f_k : [a, b] \rightarrow \mathbb{R}$, $f_k(x) := f(x, x'_k)$ gleichmäßig auf $[a, b]$ gegen $\bar{f} : x \mapsto f(x, x')$ konvergiert. Nach Satz 5.32 gilt dann nämlich

$$\lim_{k \rightarrow \infty} F(x'_k) = \lim_{k \rightarrow \infty} \int_a^b f_k(x) dx = \int_a^b \bar{f}(x) dx = F(x'),$$

was zu beweisen ist.

Um die gleichmäßige Konvergenz zu zeigen, sei $\epsilon > 0$. Die Menge $\{x'\} \cup \{x'_k : k \in \mathbb{N}\}$ ist beschränkt und abgeschlossen, also (nach dem Satz von Heine-Borel) kompakt, und somit ist f auf der Menge $[a, b] \times (\{x'\} \cup \{x'_k : k \in \mathbb{N}\})$ sogar gleichmäßig stetig. Es existiert also ein $\delta > 0$, sodaß

$$|f(x, x') - f(y, y')| < \epsilon \quad \text{falls } |x - y| + |x' - y'| < \delta.$$

Insbesondere existiert aufgrund der Konvergenz $x'_k \rightarrow x'$ ein $K \in \mathbb{N}$, sodaß

$$|\bar{f}(x) - f_k(x)| = |f(x, x') - f(x, x'_k)| < \epsilon \quad \forall k \geq K, \forall x \in [a, b];$$

es genügt dafür, K so groß zu wählen, daß $|x' - x'_k| < \delta$ für $k \geq K$. Somit ist die gleichmäßige Konvergenz gezeigt. $\square$

SATZ 8.2 (Differenzierbare Abhängigkeit des Integrals von Parametern). Seien $[a, b], [c, d] \subset \mathbb{R}$ kompakte Intervalle und $f : [a, b] \times [c, d] \rightarrow \mathbb{R}$ stetig und bezüglich der zweiten Variablen stetig partiell differenzierbar. Dann ist

$$F : [c, d] \rightarrow \mathbb{R}, \quad x' \mapsto \int_a^b f(x, x') dx$$

stetig differenzierbar, und es gilt

$$F'(x') = \int_a^b \partial_{x'} f(x, x') dx.$$

BEWEIS. Sei $x' \in [c, d]$ und $(x'_k)_{k \in \mathbb{N}} \subset [c, d]$ eine Folge, die gegen x' konvergiert. Wir setzen

$$f_k(x) := \frac{f(x, x') - f(x, x'_k)}{x' - x'_k}, \quad \bar{f}(x) := \partial_{x'} f(x, x').$$

Wenn wir zeigen können, daß f_k gleichmäßig gegen $\bar{f}$ konvergiert, so ist die Differenzierbarkeit von F gezeigt, denn dann haben wir nach Satz 5.32

$$\begin{aligned} \lim_{k \rightarrow \infty} \frac{F(x') - F(x'_k)}{x' - x'_k} &= \lim_{k \rightarrow \infty} \int_a^b \frac{f(x, x') - f(x, x'_k)}{x' - x'_k} dx \\ &= \lim_{k \rightarrow \infty} \int_a^b f_k(x) dx \\ &= \int_a^b \bar{f}(x) dx = \int_a^b \partial_{x'} f(x, x') dx. \end{aligned}$$

Wir zeigen nun die behauptete gleichmäßige Konvergenz: Nach dem Mittelwertsatz der Differentialrechnung existiert zu jedem $k \in \mathbb{N}$ ein $\xi'_k \in (x'_k, x')$, sodaß

$$f_k(x) = \partial_{x'} f(x, \xi'_k).$$

Nun ist aber die Funktion $\partial_{x'} f$ auf dem Kompaktum $[a, b] \times [c, d]$ gleichmäßig stetig. Zu $\epsilon > 0$ existiert also ein $\delta > 0$, sodaß

$$|\partial_{x'} f(x, x') - \partial_{x'} f(y, y')| < \epsilon \quad \text{falls } |x - y| + |x' - y'| < \delta.$$

Ist aber $K \in \mathbb{N}$ so groß gewählt, daß $|x'_k - x'| < \delta$ für $k \geq K$, so ist auch $|\xi'_k - x'| < \delta$ für $k \geq K$, und somit ist für alle $x \in [a, b]$

$$|f_k(x) - \bar{f}(x)| = |\partial_{x'} f(x, \xi'_k) - \partial_{x'} f(x, x')| < \epsilon \quad \text{falls } k \geq K,$$

und es folgt die gleichmäßige Konvergenz.

Die Stetigkeit von F' folgt schließlich aus der Stetigkeit von $\partial_{x'} f$ zusammen mit Satz 8.1. $\square$

Als Anwendung zeigen wir einen einfachen Spezialfall des Lemmas von POINCARÉ: Sei nämlich $\phi : B_1(0) \subset \mathbb{R}^3 \rightarrow \mathbb{R}^3$ ein zweimal stetig partiell differenzierbares Vektorfeld, dann gilt $\text{rot } D\phi = 0$ (Übung). Die Umkehrung ist ebenfalls gültig:

SATZ 8.3 (Lemma von POINCARÉ). Sei $f : B_1(0) \subset \mathbb{R}^3 \rightarrow \mathbb{R}^3$ ein stetig differenzierbares Vektorfeld mit $\text{rot } f(x) = 0$ für jedes $x \in \mathbb{R}^3$. Dann existiert eine zweimal stetig partiell differenzierbare Funktion $\phi : \mathbb{R}^3 \rightarrow \mathbb{R}$, sodaß $f = D\phi$.

BEWEIS. Wir setzen

$$\phi(x) := \int_0^1 f(tx) \cdot x dt$$

und behaupten $D\phi = f$. Zunächst ist ϕ wohldefiniert, da f stetig ist und mit x auch tx im Definitionsbereich liegt für alle $t \in [0, 1]$. Fixieren wir $x \in B_1(0)$ und eine Richtung x_l , so können wir den Integranden als Funktion von (t, x_l) betrachten (alle anderen x -Komponenten werden fixiert), mit kompaktem Definitionsbereich $(t, x_l) \in [0, 1] \times [-1 + \delta, 1 - \delta]$ für geeignetes $\delta > 0$. Dadurch dürfen wir Satz 8.2 anwenden und erhalten mithilfe der Produktregel

$$\begin{aligned} \partial_l \phi(x) &= \int_0^1 \partial_{x_l} (f(tx) \cdot x) dt \\ &= \int_0^1 t \partial_{x_l} f(tx) \cdot x dt + \int_0^1 f_l(tx) dt. \end{aligned} \quad (8.1)$$

Nun ist andererseits nach Produkt- und Kettenregel

$$\frac{d}{dt}(t f_l(tx)) = f_l(tx) + t D f_l(tx) \cdot x = f_l(tx) + t \partial_{x_l} f(tx) \cdot x,$$

wobei wir in der letzten Umformung rot $f = 0$ verwendet haben (daraus folgt nämlich $\partial_k f_l = \partial_l f_k$ und somit $\partial_l f \cdot x = D f_l \cdot x$). Der Vergleich mit (8.1) liefert schließlich mit dem Hauptsatz der Differential- und Integralrechnung

$$\partial_l \phi(x) = \int_0^1 \frac{d}{dt}(t f_l(tx)) dt = t f_l(tx)|_{t=0}^1 = f_l(x).$$

□

Man nennt in dieser Situation die Funktion ϕ auch (skalares) *Potential* von f , und bezeichnet f als *Gradienten(vektor)feld* oder *konserve(vektor-)Feld*. Das Lemma von POINCARÉ behält seine Gültigkeit auch auf anderen Definitionsbereichen als einer Kugel, genauer gesagt auf sogenannten *zusammenziehbaren* Gebieten. Dazu gehören etwa konvexe Gebiete. Auf Gebieten ‚mit Löchern‘, etwa auf $\mathbb{R}^3 \setminus \{x_3 = 0\}$, kann die Aussage des POINCARÉ-Lemmas allerdings falsch sein (Übung).

8.1.2. Mehrfachintegrale. Seien $[a, b], [c, d] \subset \mathbb{R}$ zwei kompakte Intervalle und $f : [a, b] \times [c, d]$ stetig. Nach Satz 8.1 ist die Abbildung

$$y \mapsto \int_a^b f(x, y) dx$$

auf $[c, d]$ stetig und kann daher ihrerseits integriert werden. Der resultierende Ausdruck

$$\int_c^d \left(\int_a^b f(x, y) dx \right) dy$$

heißt *Doppelintegral* und kann als Volumen unter dem Graphen von f interpretiert werden. Dies ist konsistent mit der Beobachtung, daß der Wert des Doppelintegrals für $f \equiv 1$ gleich $(b-a)(d-c)$ ist, also gleich dem anschaulichen Volumen des Quaders $[a, b] \times [c, d] \times [0, 1]$. Man beachte allerdings, daß das Konzept des Volumens einer Teilmenge von $\mathbb{R}^3$ noch nicht definiert wurde (dies ist der Ausgangspunkt der Maßtheorie, die Sie im nächsten Semester kennenlernen werden) – ebensowenig übrigens wie der Flächeninhalt einer Teilmenge von $\mathbb{R}^2$. Man kann allerdings den Flächeninhalt einer Menge, die vom Graphen einer stetigen Funktion begrenzt wird, einfach als das Integral dieser Funktion definieren, und ebenso kann man gewisse Volumina als Doppel- oder Mehrfachintegrale definieren. Wir werden auf diese Weise das Volumen der (dreidimensionalen) Einheitskugel bestimmen. Zunächst aber zeigen wir die Beliebigkeit der Integrationsreihenfolge:

SATZ 8.4. Seien $[a, b], [c, d] \subset \mathbb{R}$ zwei kompakte Intervalle und $f : [a, b] \times [c, d]$ stetig. Dann gilt

$$\int_c^d \left(\int_a^b f(x, y) dx \right) dy = \int_a^b \left(\int_c^d f(x, y) dy \right) dx.$$

BEWEIS. Setze

$$F(y) := \int_a^b \left(\int_c^y f(x, t) dt \right) dx,$$

dann ist $F(c) = 0$. Nach Hauptsatz ist $y \mapsto \int_c^y f(x, t) dt$ stetig differenzierbar mit Ableitung $f(x, y)$, und nach Satz 8.2 ist somit auch F stetig differenzierbar mit

$$F'(y) = \int_a^b f(x, y) dx.$$

Daher gilt, wieder nach Hauptsatz,

$$\int_c^d \left(\int_a^b f(x, y) dx \right) dy = \int_c^d F'(y) dy = F(d) - F(c) = \int_a^b \left(\int_c^d f(x, y) dy \right) dx.$$

□

Natürlich kann man analog für eine stetige Funktion $f : \prod_{l=1}^n [a_l, b_l] \rightarrow \mathbb{R}$ von n Variablen das Mehrfachintegral definieren, und die Integrationsreihenfolge beliebig wählen.

BEISPIEL 8.5 (Volumen der Einheitskugel). Der Graph der Funktion $f : B_1(0) \subset \mathbb{R}^2 \rightarrow \mathbb{R}$, $f(x, y) = \sqrt{1 - x^2 - y^2}$, beschreibt die obere Halbsphäre $\mathbb{S}^2 \cap \{z > 0\} \subset \mathbb{R}^3$. Man kann f stetig nach $[-1, 1] \times [-1, 1]$ fortsetzen, indem man definiert

$$f(x, y) = \begin{cases} \sqrt{1 - x^2 - y^2} & \text{falls } x^2 + y^2 < 1, \\ 0 & \text{sonst.} \end{cases}$$

Das Doppelintegral von f ist dann das Volumen der Halbkugel. Wir integrieren in der Reihenfolge $dx dy$ und bemerken, daß für festes $y \in [-1, 1]$ der Funktionswert von f nur für $x \in (-\sqrt{1 - y^2}, \sqrt{1 - y^2})$ nicht null ist. Bezeichnen wir $a_y := \sqrt{1 - y^2}$, dann ergibt sich

$$\int_{-1}^1 \int_{-1}^1 f(x, y) dx dy = \int_{-1}^1 \int_{-a_y}^{a_y} \sqrt{a_y^2 - x^2} dx dy.$$

Im zweiten Beispiel nach Satz 4.39 haben wir errechnet $\int_{-1}^1 \sqrt{1 - z^2} dz = \frac{\pi}{2}$. Dann liefert die Substitution $x = a_y z$

$$\int_{-a_y}^{a_y} \sqrt{a_y^2 - x^2} dx = a_y \int_{-1}^1 \sqrt{a_y^2 - a_y^2 z^2} dz = a_y^2 \int_{-1}^1 \sqrt{1 - z^2} dz = a_y^2 \frac{\pi}{2}.$$

Wir erhalten demnach

$$\int_{-1}^1 \int_{-1}^1 f(x, y) dx dy = \frac{\pi}{2} \int_{-1}^1 (1 - y^2) dy = \frac{\pi}{2} \left(y - \frac{y^3}{3} \right) \Big|_{-1}^1 = \frac{2}{3} \pi.$$

Die dreidimensionale Einheitskugel hat also Volumen $\frac{4}{3} \pi$.

Interessiert man sich allgemeiner für Kugeln mit Radius r , so betrachte man stattdessen die Funktion $f(x, y) = \sqrt{r^2 - x^2 - y^2}$, die auf $B_r(0)$ definiert ist und wie oben stetig auf $[-r, r]^2$ fortgesetzt werden kann. Mit den Substitutionen $y = ry'$, $x = rx'$ erhält man so

$$\int_{-r}^r \int_{-r}^r \sqrt{r^2 - x^2 - y^2} dx dy = r^2 \int_{-1}^1 \int_{-1}^1 \sqrt{r^2 - r^2 x'^2 - r^2 y'^2} dx' dy' = r^3 \frac{2}{3} \pi,$$

also ist das Volumen einer Kugel mit Radius r durch $\frac{4}{3} \pi r^3$ gegeben.

Man kann diese Rechnung iterieren und somit das Volumen ω_n der n -dimensionalen Einheitskugel auf ω_{n-1} zurückführen. Man erhält damit recht sperrige Formeln für ω_n . Wichtiger als die Kenntnis dieser Formeln ist meist der Umstand, daß das Volumen mit r^n skaliert, d.h. das Volumen einer n -dimensionalen Kugel mit Radius r ist proportional zu r^n .

8.1.3. Exkurs: Variationsprobleme. Betrachte eine zweimal stetig differenzierbare Funktion

$$L : [0, 1] \times \mathbb{R}^m \times \mathbb{R}^m \rightarrow \mathbb{R}, \quad (t, x, p) \mapsto L(t, x, p).$$

Dann definieren wir ein *Funktional*, d.h. eine Abbildung vom Raum der zweimal stetig differenzierbaren Funktionen $[0, 1] \rightarrow \mathbb{R}^m$ in die reellen Zahlen, durch

$$\mathcal{L}[x] := \int_0^1 L(t, x(t), x'(t)) dt. \quad (8.2)$$

Ein *Variationsproblem* besteht darin, unter gegebenen Randbedingungen $x(0) = x^0$ und $x(1) = x^1$ das Funktional $\mathcal{L}$ zu minimieren, das heißt: Finde unter allen zweimal stetig differenzierbaren Abbildungen $x : [0, 1] \rightarrow \mathbb{R}^m$, die $x(0) = x^0$ und $x(1) = x^1$ erfüllen, diejenige, für die (8.2) minimal wird.

SATZ 8.6 (EULER-LAGRANGE-Gleichungen). Sei $L : [0, 1] \times \mathbb{R}^m \times \mathbb{R}^m \rightarrow \mathbb{R}$ zweimal stetig differenzierbar. Ist

$$\mathcal{L}[x] = \inf \{ \mathcal{L}[y] \},$$

wobei das Infimum über alle zweimal stetig differenzierbaren Abbildungen $y : [0, 1] \rightarrow \mathbb{R}^m$ genommen wird, die $y(0) = x^0$ und $y(1) = x^1$ erfüllen, so sind die *Euler-Lagrange-Gleichungen*

$$D_x L(t, x(t), x'(t)) - \frac{d}{dt} D_p L(t, x(t), x'(t)) = 0. \quad (8.3)$$

erfüllt.

Man beachte, daß es sich um ein System von m Gleichungen handelt: Mit $D_x L$ wird nämlich der Vektor der m partiellen Ableitungen $\partial_{x_1} L, \dots, \partial_{x_m} L$ bezeichnet und analog $D_p L$.

BEWEIS. Sei $\eta : [0, 1] \rightarrow \mathbb{R}^m$ eine zweimal stetig differenzierbare Funktion mit $\eta(0) = \eta(1) = 0$, dann ist für jedes $\epsilon \in \mathbb{R}$ die Funktion $x + \epsilon\eta$ zweimal stetig differenzierbar in $(0, 1)$, stetig in $[0, 1]$, und erfüllt die Randbedingungen. Wegen Minimalität von x hat insbesondere die Funktion $\mathbb{R} \rightarrow \mathbb{R}, \epsilon \mapsto \mathcal{L}[x + \epsilon\eta]$ ein Minimum bei $\epsilon = 0$. Diese Funktion ist wegen Kettenregel stetig differenzierbar, und daher gilt

$$\begin{aligned} 0 &= \left. \frac{d}{d\epsilon} \right|_{\epsilon=0} \mathcal{L}[x + \epsilon\eta] \\ &= \left. \frac{d}{d\epsilon} \right|_{\epsilon=0} \int_0^1 L(t, x(t) + \epsilon\eta(t), x'(t) + \epsilon\eta'(t)) dt \\ &\stackrel{\text{Satz 8.2}}{=} \int_0^1 \left. \frac{d}{d\epsilon} \right|_{\epsilon=0} L(t, x(t) + \epsilon\eta(t), x' + \epsilon\eta'(t)) dt \\ &\stackrel{\text{Kettenregel}}{=} \int_0^1 D_x L(t, x(t), x'(t)) \cdot \eta(t) + D_p L(t, x(t), x'(t)) \cdot \eta'(t) dt \\ &= \int_0^1 \left[D_x L(t, x(t), x'(t)) - \frac{d}{dt} D_p L(t, x(t), x'(t)) \right] \cdot \eta(t) dt, \end{aligned}$$

wobei im letzten Schritt unter Beachtung von $\eta(0) = \eta(1) = 0$ partiell integriert wurde.

Eine stetige Funktion ist identisch null, wenn das Integral ihres Produkts mit jeder zweimal stetig differenzierbaren Funktion mit Nullrandwerten verschwindet (Übung). Es folgt also

$$D_x L(t, x(t), x'(t)) - \frac{d}{dt} D_p L(t, x(t), x'(t)) = 0.$$

□

Ähnlich der Übersetzung des Optimierungsproblems für eine Funktion $f: \mathbb{R}^n \rightarrow \mathbb{R}$ in das algebraische Gleichungssystem $Df(x) = 0$ kann man also ein Optimierungsproblem für ein Funktional in ein System gewöhnlicher Differentialgleichungen, nämlich die EULER-LAGRANGE-Gleichungen, transformieren.

BEISPIEL 8.7 (Kürzeste Verbindung zwischen zwei Punkten). Eine stetige Abbildung $x: [0, 1] \rightarrow \mathbb{R}^m$ heißt *Kurve*. Falls x in $[0, 1]$ sogar stetig differenzierbar ist, so ist das Integral

$$\int_0^1 |x'(t)| dt$$

wohldefiniert und heißt *Länge* der Kurve. Wir suchen die zweimal stetig differenzierbare Kurve(n) zwischen zwei Punkten $x^0, x^1 \in \mathbb{R}^m$, deren Länge minimal ist.

Dazu nehmen wir zusätzlich an $x'(t) \neq 0$ für alle $t \in [0, 1]$ ¹. Setzen wir $L(t, x, p) = |p|$, so ist L im Wertebereich von x' zweimal stetig differenzierbar (den Wert $x' = 0$ haben wir ja ausgeschlossen), es gilt $D_x L \equiv 0$, $D_p L(p) = \frac{p}{|p|}$, und deshalb gemäß (8.3)

$$\frac{d}{dt} \frac{x'(t)}{|x'(t)|} = 0.$$

Da der Vektor $x'(t)$ tangential zur Kurve am Punkt $x(t)$ liegt, besagt die Euler-Lagrange-Gleichung also, daß die Richtung der Tangente entlang der Kurve konstant ist. Die Kurve ist somit eine gerade Strecke und wir haben gezeigt, daß – wie zu erwarten war – die gerade Strecke die kürzeste Verbindung zwischen zwei Punkten im Raum ist.

BEISPIEL 8.8 (Klassische Mechanik). Ein punktförmiger Körper der Masse m befinde sich zur Zeit $t \in \mathbb{R}$ am Punkt $x(t) \in \mathbb{R}^3$. Er sei einer Kraft $F(x(t))$ unterworfen, wobei $F: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ ein konservatives Vektorfeld ist, d.h. $F = DV$ für ein Potential $V: \mathbb{R}^3 \rightarrow \mathbb{R}$ (vgl. Satz 8.3). Die Größen x' und x'' werden als Geschwindigkeit bzw. Beschleunigung interpretiert.² Man bezeichnet mit $\frac{1}{2}m|x'|^2$ die *kinetische Energie* des Körpers.

Dann postuliert man das *Prinzip der kleinsten Wirkung*³:

Ein Körper, der sich zur Zeit t^0 am Ort x^0 und zur Zeit t^1 am Ort x^1 befindet, durchläuft dazwischen diejenige Kurve $x(t)$, die die *Wirkung* minimiert:

$$\mathcal{L}[x] = \int_{t^0}^{t^1} L(x(t), x'(t)) dt,$$

wobei $L(x, x') := \frac{1}{2}m|x'|^2 - V(x)$ die LAGRANGE-Funktion ist. Die Euler-Lagrange-Gleichungen lauten dann

$$0 = DV(x(t)) - \frac{d}{dt}(mx'(t)) = F(x(t)) - mx''(t),$$

was genau das zweite NEWTONsche Gesetz ist.

¹Dies ist keine echte Einschränkung, da eine Kurve diese Bedingung ggf. nach Reparametrisierung stets erfüllt.

²In der Physik verwendet man meist die Notation $\dot{x}$ bzw. $\ddot{x}$.

³Dieses Prinzip wird von Wissenschaftshistorikerinnen entweder LEIBNIZ, MAUPERTUIS oder EULER zugerechnet. In jedem Falle stammt es aus der ersten Hälfte des 18. Jahrhunderts und spielt seitdem in der theoretischen Physik – nicht nur in der klassischen Mechanik – eine überragende Rolle.

8.2. Die Transformationsformel

Ziel dieses Abschnitts ist der Beweis der Transformationsformel, die als höherdimensionale Verallgemeinerung der Substitutionsregel angesehen werden kann. Während aber letztere sehr einfach aus der Kettenregel und dem Hauptsatz folgt, ist der Beweis der Transformationsformel erheblich aufwendiger. Wir folgen der Darstellung in FORSTER [10].

8.2.1. Stetige Funktionen mit kompaktem Träger. Der *Träger* einer Funktion $f : \mathbb{R}^n \rightarrow \mathbb{R}$ ist definiert als⁴

$$\text{supp}(f) := \overline{\{x \in \mathbb{R}^n : f(x) \neq 0\}}.$$

Hierbei bedeutet der Querbalken wie zuvor den topologischen Abschluß. Außerhalb ihres Trägers ist eine Funktion also identisch null.

Wir bezeichnen mit $C_c(\mathbb{R}^n)$ die Menge der stetigen Funktionen, deren Träger kompakt ist. Ist $f \in C_c(\mathbb{R}^n)$, so ist $\text{supp}(f)$ insbesondere beschränkt, und es gibt somit einen Quader $Q = \prod_{j=1}^n [a_j, b_j]$, der $\text{supp}(f)$ enthält. Wir definieren daher für $f \in C_c(\mathbb{R}^n)$

$$\int_{\mathbb{R}^n} f(x) dx := \int_{a_n}^{b_n} \int_{a_{n-1}}^{b_{n-1}} \cdots \int_{a_1}^{b_1} f(x_1, x_2, \dots, x_n) dx_1 dx_2 \cdots dx_n.$$

Diese Definition ist offenbar unabhängig von der Wahl des Quaders (solange er $\text{supp}(f)$ enthält).

Aus den Eigenschaften des eindimensionalen Integrals erhalten wir unmittelbar

PROPOSITION 8.9 (Linearität und Monotonie). *Seien $f, g \in C_c(\mathbb{R}^n)$ und $\lambda \in \mathbb{R}$, so gilt*

- (1) $\int_{\mathbb{R}^n} (f + g)(x) dx = \int_{\mathbb{R}^n} f(x) dx + \int_{\mathbb{R}^n} g(x) dx;$
- (2) $\int_{\mathbb{R}^n} (\lambda f)(x) dx = \lambda \int_{\mathbb{R}^n} f(x) dx.$

Ist außerdem $f \leq g$ (das bedeutet $f(x) \leq g(x)$ für alle $x \in \mathbb{R}^n$), so ist

- (3) $\int_{\mathbb{R}^n} f(x) dx \leq \int_{\mathbb{R}^n} g(x) dx.$

PROPOSITION 8.10 (Translationsinvarianz). *Für jedes $a \in \mathbb{R}^n$ und $f \in C_c(\mathbb{R}^n)$ gilt*

$$\int_{\mathbb{R}^n} f(x - a) dx = \int_{\mathbb{R}^n} f(x) dx.$$

BEWEIS. Dies folgt durch n -malige Verwendung der entsprechenden Eigenschaft des eindimensionalen Integrals. $\square$

Ein *Funktional* auf $C_c(\mathbb{R}^n)$ ist eine Abbildung $C_c(\mathbb{R}^n) \rightarrow \mathbb{R}$. Man kann den Inhalt der beiden Propositionen also auch so formulieren: Das Integral ist ein lineares, monotones, translationsinvariantes Funktional auf $C_c(\mathbb{R}^n)$. Im nächsten Abschnitt zeigen wir, daß das Integral bis auf eine multiplikative Konstante sogar das *einzige* Funktional mit diesen Eigenschaften ist.

8.2.2. Das Integral als lineares, monotones, translationsinvariantes Funktional. Ziel dieses Abschnitts ist der Beweis der folgenden Aussage:

SATZ 8.11. Sei $I : C_c(\mathbb{R}^n) \rightarrow \mathbb{R}$ ein lineares, monotones, translationsinvariantes Funktional, das heißt: Für alle $\lambda \in \mathbb{R}$, $f, g \in C_c(\mathbb{R}^n)$ gilt

- (1) $I(f + g) = I(f) + I(g);$
- (2) $I(\lambda f) = \lambda I(f);$
- (3) aus $f \leq g$ folgt $I(f) \leq I(g);$
- (4) Für $T_a f \in C_c(\mathbb{R}^n)$, definiert durch $T_a f(x) := f(x - a)$ für ein $a \in \mathbb{R}^n$, gilt $I \circ T_a = I.$

⁴supp wegen engl. *support*.

Dann gibt es eine reelle Zahl $c \geq 0$, sodaß für alle $f \in C_c(\mathbb{R}^n)$ gilt

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

Der Beweis dieses Satzes bedarf dreier Lemmata, die wir nun nacheinander zeigen werden.

LEMMA 8.12. *Sei $I : C_c(\mathbb{R}^n) \rightarrow \mathbb{R}$ ein lineares monotones Funktional und sei $(f_k)_{k \in \mathbb{N}}$ eine Folge in $C_c(\mathbb{R}^n)$, die gleichmäßig gegen $f \in C_c(\mathbb{R}^n)$ konvergiert, und die gleichmäßig kompakten Träger hat (d.h. es gibt eine kompakte Menge $K \subset \mathbb{R}^n$, sodaß $\text{supp}(f_k) \subset K$ für alle $k \in \mathbb{N}$). Dann gilt*

$$\lim_{k \rightarrow \infty} I(f_k) = I(f).$$

BEWEIS. Sei $K \subset \mathbb{R}^n$ eine kompakte Menge, die die Träger aller f_k enthält. Wir wählen eine nichtnegative Funktion $\phi \in C_c(\mathbb{R}^n)$ dergestalt, daß $\phi|_K \equiv 1$ (warum existiert eine solche?). Setze $\alpha_k := \sup_{x \in K} |f_k(x) - f(x)|$, sodaß wegen der gleichmäßigen Konvergenz $\lim_{k \rightarrow \infty} \alpha_k = 0$. Wegen $\text{supp}(f_k - f) \subset K$ ist dann

$$-\alpha_k \phi \leq f_k - f \leq \alpha_k \phi,$$

und aufgrund der Linearität und Monotonie von I auch

$$-\alpha_k I(\phi) \leq I(f_k) - I(f) \leq \alpha_k I(\phi),$$

mithin $|I(f_k) - I(f)| \leq I(\phi) \alpha_k$. Da $\alpha_k \rightarrow 0$, folgt die Behauptung. $\square$

Das nächste Lemma befaßt sich mit dem Skalierungsverhalten eines linearen translationsinvarianten Funktionals. Wir betrachten dazu die Funktion $\psi \in C_c(\mathbb{R})$, gegeben durch

$$\psi(t) := \begin{cases} 1 - |t| & \text{falls } |t| \leq 1, \\ 0 & \text{sonst.} \end{cases} \quad (8.4)$$

Offenbar ist $\int_{\mathbb{R}} \psi(t) dt = 1$. Wir definieren außerdem $\Psi \in C_c(\mathbb{R}^n)$ durch

$$\Psi(x_1, \dots, x_n) = \psi(x_1) \cdot \dots \cdot \psi(x_n),$$

sodaß $\text{supp}(\Psi) = [-1, 1]^n$. Eine *skalierte* Version dieser Funktion ist dann

$$\Psi_\delta(x) := \Psi\left(\frac{x}{\delta}\right).$$

Es handelt sich um die um den Faktor $\delta > 0$ ‚gestauchte‘ Funktion Ψ , für die also gilt $\text{supp}(\Psi_\delta) = [-\delta, \delta]^n$, und man berechnet leicht mithilfe der Substitution $\tilde{x} = \frac{x}{\delta}$, daß

$$\int_{\mathbb{R}^n} \Psi_\delta(x) dx = \delta^n.$$

Wir zeigen nun, daß *jedes* lineare translationsinvariante Funktional in der gleichen Weise skaliert:

LEMMA 8.13. *Sei $I : C_c(\mathbb{R}^n) \rightarrow \mathbb{R}$ ein lineares translationsinvariantes Funktional und Ψ, Ψ_δ wie oben. Dann gilt für jedes $\delta > 0$*

$$I(\Psi_{\delta/2}) = 2^{-n} I(\Psi_\delta).$$

BEWEIS. *Schritt 1.* Wir erinnern uns an die Notation $T_a f(x) := f(x - a)$, sodaß also T_a die Translation einer Funktion um a bezeichnet. Man mache sich klar, daß

$$\psi_\delta = \frac{1}{2} T_{-\delta/2} \psi_{\delta/2} + \psi_{\delta/2} + \frac{1}{2} T_{\delta/2} \psi_{\delta/2}. \quad (8.5)$$

Schritt 2. Sei die Dimension n nun beliebig, dann ist nach Definition $\Psi_\delta(x_1, \dots, x_n) = \prod_{k=1}^n \psi_\delta(x_k)$. Die Formel (8.5) läßt sich schreiben als

$$\psi_\delta(x_k) = \sum_{l=-1}^1 2^{-|l|} (T_{l\delta/2} \psi_{\delta/2})(x_k). \quad (8.6)$$

Um daraus Ψ_δ zu erhalten, müssen wir die Identitäten (8.6) über alle $k = 1, \dots, n$ aufmultiplizieren und erhalten so

$$\Psi_\delta(x_1, \dots, x_n) = \sum_{\alpha \in A} \gamma_\alpha (T_{\alpha\delta/2} \Psi_{\delta/2})(x_1, \dots, x_n),$$

wobei $A := \{(\alpha_1, \dots, \alpha_n) : \alpha_k \in \{-1, 0, 1\}\}$ und $\gamma_\alpha := \prod_{k=1}^n 2^{-|\alpha_k|}$. Unter Beachtung der Translationsinvarianz und Linearität von I erhalten wir daraus

$$I(\Psi_\delta) = \sum_{\alpha \in A} \gamma_\alpha I(T_{\alpha\delta/2} \Psi_{\delta/2}) = I(\Psi_{\delta/2}) \sum_{\alpha \in A} \gamma_\alpha.$$

Schritt 3. Wir zeigen durch Induktion nach n , daß

$$\sum_{\alpha \in A} \gamma_\alpha = 2^n.$$

Ist dies gezeigt, so ist das Lemma bewiesen.

Für $n = 1$ erhalten wir dies durch direkte Rechnung. Angenommen, es ist $\sum_{\alpha \in A} \gamma_\alpha = 2^n$ für ein $n \in \mathbb{N}$. Wir bezeichnen mit A' die entsprechende Menge für $n + 1$, also $A := \{(\alpha_1, \dots, \alpha_{n+1}) : \alpha_k \in \{-1, 0, 1\}\}$. Dann können wir jedes $\alpha' \in A'$ schreiben als (α, α_{n+1}) mit $\alpha \in A$, und wir haben

$$\gamma_{\alpha'} = \gamma_{(\alpha, \alpha_{n+1})} = 2^{-|\alpha_{n+1}|} \prod_{k=1}^n 2^{-|\alpha_k|} = 2^{-|\alpha_{n+1}|} \gamma_\alpha.$$

Somit erhalten wir

$$\sum_{\alpha' \in A'} \gamma_{\alpha'} = \sum_{\alpha_{n+1}=-1}^1 \sum_{\alpha \in A} 2^{-|\alpha_{n+1}|} \gamma_\alpha = \left(\frac{1}{2} + 1 + \frac{1}{2}\right) \sum_{\alpha \in A} \gamma_\alpha = 2 \cdot 2^n = 2^{n+1}.$$

□

Zur Vorbereitung des nächsten Lemmas bemerken wir, daß $\sum_{l \in \mathbb{Z}} T_l \psi(x) = 1$ für alle $x \in \mathbb{R}$, wobei die Summe an jeder Stelle $x \in \mathbb{R}$ höchstens zwei Summanden ungleich null enthält. Man bezeichnet die Familie $(T_l \psi)_{l \in \mathbb{Z}}$ daher auch als eine *Teilung der Eins* von $\mathbb{R}$.

Daher gilt auch

$$\sum_{\alpha \in \mathbb{Z}^n} (T_\alpha \Psi)(x_1, \dots, x_n) = \prod_{k=1}^n \sum_{\alpha_k \in \mathbb{Z}} (T_{\alpha_k} \psi)(x_k) = 1$$

und somit für jedes $\delta > 0$

$$\sum_{\alpha \in \mathbb{Z}^n} (T_{\delta\alpha} \Psi_\delta)(x) = \sum_{\alpha \in \mathbb{Z}^n} \Psi_\delta(x - \delta\alpha) = \sum_{\alpha \in \mathbb{Z}^n} \Psi\left(\frac{x}{\delta} - \alpha\right) = 1.$$

Die Familie $(T_{\delta\alpha} \Psi_\delta)_{\alpha \in \mathbb{Z}^n}$ ist also eine Teilung der Eins von $\mathbb{R}^n$.

LEMMA 8.14. *Sei $f \in C_c(\mathbb{R}^n)$. Dann existiert zu jedem $\epsilon > 0$ ein $\delta_0 > 0$, sodaß für alle $0 < \delta \leq \delta_0$ gilt:*

$$\sup_{x \in \mathbb{R}^n} \left| f(x) - \sum_{\alpha \in \mathbb{Z}^n} f(\delta\alpha) (T_{\delta\alpha} \Psi_\delta)(x) \right| \leq \epsilon.$$

BEWEIS. Da f kompakten Träger hat, ist es sogar gleichmäßig stetig. Sei also zu gegebenem $\epsilon > 0$ ein $\tilde{\delta} > 0$ so gewählt, daß

$$|f(x) - f(y)| < \epsilon \quad \text{falls} \quad |x - y| < \tilde{\delta}.$$

Wir setzen nun $\delta_0 := \frac{\tilde{\delta}}{\sqrt{n}}$. Sei nun $0 < \delta \leq \tilde{\delta}$ und $x \in \mathbb{R}^n$ beliebig. Da $(T_{\delta\alpha}\Psi_\delta)_{\alpha \in \mathbb{Z}^n}$ eine Teilung der Eins ist, gilt $f(x) = \sum_{\alpha \in \mathbb{Z}^n} f(x)(T_{\delta\alpha}\Psi_\delta)(x)$ und deshalb auch

$$\begin{aligned} \left| f(x) - \sum_{\alpha \in \mathbb{Z}^n} f(\delta\alpha)(T_{\delta\alpha}\Psi_\delta)(x) \right| &= \left| \sum_{\alpha \in \mathbb{Z}^n} (f(x) - f(\delta\alpha))(T_{\delta\alpha}\Psi_\delta)(x) \right| \\ &= \left| \sum_{\alpha \in A} (f(x) - f(\delta\alpha))(T_{\delta\alpha}\Psi_\delta)(x) \right|, \end{aligned}$$

wobei A die Menge aller Elemente von $\mathbb{Z}^n$ bezeichnet, für die $x \in \text{supp}(T_{\delta\alpha}\Psi_\delta)$. Für jedes $\alpha \in A$ ist aber $|x - \delta\alpha| \leq \sqrt{n}\delta < \tilde{\delta}$ und somit wegen der gleichmäßigen Stetigkeit auch $|f(x) - f(\delta\alpha)| < \epsilon$. Es folgt

$$\begin{aligned} \left| f(x) - \sum_{\alpha \in \mathbb{Z}^n} f(\delta\alpha)(T_{\delta\alpha}\Psi_\delta)(x) \right| &= \left| \sum_{\alpha \in A} (f(x) - f(\delta\alpha))(T_{\delta\alpha}\Psi_\delta)(x) \right| \\ &< \epsilon \sum_{\alpha \in A} (T_{\delta\alpha}\Psi_\delta)(x) = \epsilon. \end{aligned}$$

□

BEWEIS VON SATZ 8.11. Sei $I: \mathbb{C}_c(\mathbb{R}^n) \rightarrow \mathbb{R}$ linear, monoton, und translationsinvariant und setze $c := I(\Psi)$ für Ψ wie oben. Setze $J(f) := c \int_{\mathbb{R}^n} f(x) dx$. Wir müssen zeigen $I = J$.

Beachte dafür zunächst, daß für jedes $m \in \mathbb{N}$ gilt

$$I(\Psi_{2^{-m}}) = 2^{-mn} I(\Psi) = 2^{-mn} J(\Psi) = J(\Psi_{2^{-m}}).$$

Dies folgt nämlich durch m -malige Anwendung von Lemma 8.13 (mit $\delta := 2^{-m}$) und aus der Eigenschaft $\int_{\mathbb{R}^n} \Psi(x) dx = 1$.

Sei nun $f \in C_c(\mathbb{R}^n)$ beliebig und setze

$$f_m := \sum_{\alpha \in \mathbb{R}^n} f(2^{-m}\alpha) T_{2^{-m}\alpha} \Psi_{2^{-m}}.$$

Da $I(\Psi_{2^{-m}}) = J(\Psi_{2^{-m}})$, folgt aus der Linearität und Translationsinvarianz der beiden Funktionalen $I(f_m) = J(f_m)$. Nach Lemma 8.14 konvergiert $(f_m)_{m \in \mathbb{N}}$ für $m \rightarrow \infty$ gleichmäßig gegen f , und die Träger der f_m sind allesamt in der kompakten Menge

$$\{x \in \mathbb{R}^n : |x - y| \leq \sqrt{n} \text{ für ein } y \in \text{supp}(f)\}.$$

enthalten. Daher folgt aus Lemma 8.12 $I(f) = J(f)$.

□

8.2.3. Die Transformationsformel für lineare Transformationen. Wir wollen in diesem Abschnitt das Integral von $f \circ A$ auf das von f zurückführen, wobei A eine invertierbare lineare Abbildung $\mathbb{R}^n \rightarrow \mathbb{R}^n$ sei (die wir mit einer Matrix $A \in \mathbb{R}^{n \times n}$ identifizieren). Wir behandeln erst orthogonale Matrizen und führen dann den allgemeinen Fall auf diese zurück.

PROPOSITION 8.15. *Sei $I: C_c(\mathbb{R}^n) \rightarrow \mathbb{R}$ ein lineares, monotones, translationsinvariantes Funktional und $A \in \mathbb{R}^{n \times n}$ invertierbar. Dann ist auch $J(f) := I(f \circ A)$ linear, monoton und translationsinvariant.*

BEWEIS. Beachte, daß $f \circ A \in C_c(\mathbb{R}^n)$ genau dann, wenn $f \in C_c(\mathbb{R}^n)$ (warum?). Zur Linearität: Sind $f, g \in C_c(\mathbb{R}^n)$, so gilt

$$J(f + g) = I((f + g) \circ A) = I(f \circ A + g \circ A) = I(f \circ A) + I(g \circ A) = J(f) + J(g).$$

Monotonie folgt daraus, daß $f \leq g$ auch $f \circ A \leq g \circ A$ impliziert. Zur Translationsinvarianz rechnen wir

$$J(T_a f) = I((T_a f) \circ A) = I(T_{A^{-1}a}(f \circ A)) = I(f \circ A) = J(f).$$

□

Mit $O(n)$ bezeichnen wir die Menge der orthogonalen $n \times n$ -Matrizen, also der Matrizen mit $A^t = A^{-1}$. Die dadurch beschriebenen Abbildungen erhalten bekanntlich Längen und Winkel; das Bild einer Menge unter einer orthogonalen Transformation ist daher zur Menge selbst kongruent.

KOROLLAR 8.16 (Bewegungsinvarianz). Sei $A \in O(n)$ und $f \in C_c(\mathbb{R}^n)$. Dann gilt

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

BEWEIS. Bezeichnen wir $I(f) = \int_{\mathbb{R}^n} f(x) dx$ und $J(f) = I(f \circ A)$, so sind I und J nach Proposition 8.15 lineare, monotone, translationsinvariante Funktionale und stimmen daher nach Satz 8.11 bis auf eine Konstante überein, $J = cI$. Zur Bestimmung dieser Konstanten betrachten wir eine Funktion $g \in C_c(\mathbb{R}^n)$ mit $g \circ A = g$; zum Beispiel können wir eine beliebige Funktion wählen, die nur von $|x|$ abhängt (denn $|Ax| = |x|$ wegen der Orthogonalität). Für diese Funktion ist $I(g) = J(g)$, und sofern g so gewählt ist, daß $I(g) \neq 0$, folgt $c = 1$ und damit $I(f) = J(f)$ für alle $f \in C_c(\mathbb{R}^n)$. □

Die folgende Proposition haben wir im Grunde bereits in Beispiel 8.5 verwendet, als wir das Volumen einer Kugel vom Radius r berechnet haben (dort war $\lambda_1 = \dots = \lambda_n = r$):

PROPOSITION 8.17. Seien $\lambda_1, \dots, \lambda_n > 0$ und $f \in C_c(\mathbb{R}^n)$. Dann gilt

$$\int_{\mathbb{R}^n} f(x) dx = \lambda_1 \cdot \dots \cdot \lambda_n \int_{\mathbb{R}^n} f(\lambda_1 x_1, \dots, \lambda_n x_n) dx.$$

BEWEIS. Dies folgt durch n -malige Anwendung der Substitution $\tilde{x}_k = \lambda_k x_k$. □

Das folgende Resultat aus der Linearen Algebra zitieren wir ohne Beweis:

LEMMA 8.18 (Singulärwertzerlegung). Sei $A \in \mathbb{R}^{n \times n}$ invertierbar. Dann existieren orthogonale Matrizen $U_1, U_2 \in O(n)$ und eine Diagonalmatrix D mit positiven Einträgen, sodaß $A = U_1 D U_2$.

Die Diagonaleinträge $\lambda_1, \dots, \lambda_n$ von D sind durch A (bis auf die Reihenfolge) eindeutig bestimmt und heißen *Singulärwerte* von A .

SATZ 8.19 (Transformationsformel für lineare Abbildungen). Sei $A \in \mathbb{R}^{n \times n}$ invertierbar und $f \in C_c(\mathbb{R}^n)$. Dann gilt

$$|\det A| \int_{\mathbb{R}^n} f(Ax) dx = \int_{\mathbb{R}^n} f(y) dy.$$

BEWEIS. Sei $A = U_1 D U_2$ die Singulärwertzerlegung von A . Nach Determinantenproduktsatz und wegen $|\det U_1| = |\det U_2| = 1$, $D > 0$ gilt $|\det A| = \det D = \lambda_1 \cdot \dots \cdot \lambda_n$. Mit $I(f) := \int_{\mathbb{R}^n} f(x) dx$ haben wir mit Korollar 8.16 und Proposition 8.17:

$$I(f) = I(f \circ U_1) = \det DI(f \circ U_1 \circ D) = \det DI(f \circ U_1 \circ D \circ U_2) = |\det A| I(f \circ A).$$

□

Man kann dies als eine Art Substitutionsformel lesen, indem man $y = Ax$ und $dy = |\det A| dx$ substituiert.

BEISPIEL 8.20 (Volumen eines Ellipsoids). Ein *Ellipsoid* in $\mathbb{R}^3$ wird durch die Gleichung

$$\left(\frac{x}{a}\right)^2 + \left(\frac{y}{b}\right)^2 + \left(\frac{z}{c}\right)^2 = 1$$

beschrieben, wobei $a, b, c > 0$ die *Halbachsen* sind. Wir können das Volumen des Ellipsoids berechnen, indem wir nach z auflösen und das Integral der Funktion

$$f(x, y) = \begin{cases} c\sqrt{1 - \left(\frac{x}{a}\right)^2 - \left(\frac{y}{b}\right)^2} & \text{falls } \left(\frac{x}{a}\right)^2 + \left(\frac{y}{b}\right)^2 \leq 1 \\ 0 & \text{sonst} \end{cases}$$

berechnen. Wir bemerken dazu, daß $f = c\tilde{f} \circ A$ ist, wobei $\tilde{f}$ die Funktion aus Beispiel 8.5 ist (deren Graph die obere Halbkugel ist) und

$$A = \begin{pmatrix} \frac{1}{a} & 0 \\ 0 & \frac{1}{b} \end{pmatrix}.$$

Die Transformationsformel liefert dann

$$\int_{\mathbb{R}^n} \tilde{f}(x, y) dx dy = c \int_{\mathbb{R}^n} f \circ A(x, y) dx dy = c(\det A)^{-1} \int_{\mathbb{R}^n} f(x, y) dx dy = \frac{2}{3} \pi abc,$$

wobei wir das Resultat aus Beispiel 8.5 verwendet haben. Das Volumen des Ellipsoids beträgt also $\frac{4}{3} \pi abc$.

8.2.4. Die allgemeine Transformationsformel. Der Nutzen der Transformationsformel bleibt natürlich sehr eingeschränkt, solange nur lineare Transformationen zugelassen sind. Deshalb verallgemeinern wir nun Satz 8.19 auf potentiell nichtlineare Transformationen.

Sei $\Omega \subset \mathbb{R}^n$ offen. Analog zu $C_c(\mathbb{R}^n)$ definieren wir $C_c(\Omega)$ als die Menge aller stetigen Funktionen $f : \Omega \rightarrow \mathbb{R}$, die kompakten Träger $K \subset \Omega$ haben. Das Integral von f über Ω ist dann einfach als das Integral derjenigen Funktion über $\mathbb{R}^n$ definiert, die man durch Fortsetzung von f durch null erhält.

Eine Abbildung $\phi : \Omega \rightarrow \phi(\Omega) \subset \mathbb{R}^n$ heißt *C^1 -Diffeomorphismus*, wenn sie bijektiv ist und wenn sowohl ϕ als auch die Umkehrfunktion ϕ^{-1} stetig differenzierbar sind.

SATZ 8.21 (Transformationsformel). Sei $\phi : \Omega \rightarrow \phi(\Omega)$ ein C^1 -Diffeomorphismus. Dann gilt für jedes $f \in C_c(\phi(\Omega))$

$$\int_{\Omega} f(\phi(x)) |\det D\phi(x)| dx = \int_{\phi(\Omega)} f(y) dy.$$

BEWEIS. *Schritt 1.* In diesem Schritt treffen wir lediglich einige Vorbereitungen.

Für $x \in \mathbb{R}^n$ definieren wir die *Maximumsnorm* $|x|_{\infty} := \max\{|x_1|, \dots, |x_n|\}$ (man überzeuge sich, daß dies tatsächlich eine Norm ist). Die Maximumsnorm ist zur euklidischen Norm $|\cdot|$ äquivalent⁵ in dem Sinne, daß

$$\frac{1}{\sqrt{n}} |x| \leq |x|_{\infty} \leq |x|,$$

was man sich leicht überlegt.

Mit $Q(a, \epsilon) := \{x \in \mathbb{R}^n : |x - a|_{\infty} \leq \epsilon\}$ bezeichnen wir den abgeschlossenen Würfel⁶ um $a \in \mathbb{R}^n$ mit Seitenlänge 2ϵ .

Sei $f \in C_c(\phi(\Omega))$ gegeben mit kompaktem Träger $L \subset \phi(V)$ und $K := \phi^{-1}(L)$. Da ϕ^{-1} stetig ist, ist K ebenfalls kompakt (Satz 6.15). Wir zeigen, daß $d(K, \partial\Omega) := \inf\{|x - y| : x \in K, y \in \partial\Omega\}$ positiv ist: Angenommen, dies wäre nicht der Fall, so gäbe es Folgen $(x^k)^{k \in \mathbb{N}} \subset K$ und $(y^k)^{k \in \mathbb{N}} \subset \partial\Omega$, sodaß $|x^k - y^k| \rightarrow 0$. Da K kompakt ist, konvergiert eine Teilfolge $(x^{k_l})^{l \in \mathbb{N}}$ gegen ein $x \in K$; ebenso konvergiert eine weitere Teilfolge $(y^{k_{l_j}})^{j \in \mathbb{N}}$ gegen ein $y \in \partial\Omega$, denn $\partial\Omega$ ist abgeschlossen und $\partial\Omega \cap \overline{B_R(0)}$ ist daher kompakt, wobei $R > 0$ so groß gewählt ist, daß $|y^k| < R$ für alle $k \in \mathbb{N}$ (eine solche Schranke existiert, da sonst $|x^k - y^k|$ nicht gegen null konvergieren könnte – die Folge $(x^k)^{k \in \mathbb{N}}$ ist nämlich beschränkt).

⁵Allgemein sind in endlichdimensionalen Vektorräumen alle Normen äquivalent (Übung).

⁶Aus abstrakter Sicht handelt es sich eigentlich um die *Kugel* um a mit Radius ϵ bzgl. der Norm $|\cdot|_{\infty}$.

Wegen $\lim_{k \rightarrow \infty} |x^k - y^k| = 0$ gilt aber $x = y$, also $K \cap \partial\Omega \neq \emptyset$, im Widerspruch zu $K \subset \Omega$ und $\Omega \cap \partial\Omega = \emptyset$ (da Ω offen ist). Ebenso zeigt man $d(L, \partial\phi(\Omega)) > 0$.

Es existiert daher ein $\epsilon_1 > 0$, sodaß

$$\begin{aligned} K &\subset K' := \{x \in \mathbb{R}^n : |x - y|_\infty \leq 2\epsilon_1 \text{ für ein } y \in K\} \subset \Omega, \\ L &\subset L' := \{x \in \mathbb{R}^n : |x - y|_\infty \leq 2\epsilon_1 \text{ für ein } y \in L\} \subset \Omega. \end{aligned}$$

Beachte dabei, daß K' und L' wieder beschränkt und abgeschlossen, also kompakt sind. Außerdem impliziert die Wahl dieser Mengen, daß

$$Q(a, \epsilon_1) \subset K' \quad \forall a \in K, \quad Q(b, \epsilon_1) \subset L' \quad \forall b \in L.$$

Man veranschauliche sich die Situation an einer Zeichnung.

Da $D\phi$ und $D\phi^{-1}$ stetig sind, sind sie auf den Kompakta K' bzw. L' beschränkt, also existiert $C > 0$ mit

$$|D\phi(a)\xi|_\infty \leq C|\xi|_\infty, \quad |D\phi^{-1}(b)\xi|_\infty \leq C|\xi|_\infty \quad \forall a \in K', b \in L', \xi \in \mathbb{R}^n. \quad (8.7)$$

Mit dem höherdimensionalen Mittelwertsatz (Satz 7.12) erhalten wir für alle $a, x \in \Omega$, deren Verbindungsstrecke in Ω liegt,

$$\phi(x) - \phi(a) = \left(\int_0^1 D\phi(a + t(x - a)) dt \right) \cdot (x - a), \quad (8.8)$$

und somit

$$\phi(Q(a, \epsilon)) \subset Q(\phi(a), C\epsilon), \quad \phi^{-1}(Q(b, \epsilon)) \subset Q(\phi^{-1}(b), C\epsilon) \quad (8.9)$$

für alle $a \in K, b \in L, \epsilon \leq \epsilon_1$.

Schritt 2. Für jedes $a \in \Omega$ definieren wir die lineare Approximation $\lambda_a : \mathbb{R}^n \rightarrow \mathbb{R}^n$ von ϕ am Punkt a als

$$\lambda_a(x) = \phi(a) + D\phi(a)(x - a).$$

Aus (8.8) folgt

$$\phi(x) - \lambda_a(x) = \left(\int_0^1 [D\phi(a + t(x - a)) - D\phi(a)] dt \right) \cdot (x - a). \quad (8.10)$$

Wir behaupten, daß eine monotone nichtnegative Funktion $\omega_1 : [0, \epsilon_1] \rightarrow \mathbb{R}$ mit $\lim_{\epsilon \rightarrow 0} \omega_1(\epsilon) = 0$ existiert⁷, sodaß

$$|D\phi(x) - D\phi(y)| \leq \omega_1(|x - y|_\infty) \quad \forall x, y \in K'. \quad (8.11)$$

In der Tat, für $\delta \geq 0$ setze

$$\omega_1(\delta) := \sup\{|D\phi(x) - D\phi(y)| : x, y \in K', |x - y|_\infty \leq \delta\}.$$

Da $D\phi$ als stetige Funktion auf dem Kompaktum K' beschränkt ist, ist ω_1 wohldefiniert, und per definitionem gilt (8.11). Für den behaupteten Limes beachte, daß es aufgrund der gleichmäßigen Stetigkeit von $D\phi$ in K' zu jedem $\epsilon > 0$ ein $\delta > 0$ gibt, sodaß $|D\phi(x) - D\phi(y)| < \epsilon$ falls $|x - y|_\infty < \delta$. Für solche δ ist daher $\omega_1(\delta) < \epsilon$, und es folgt $\lim_{\epsilon \rightarrow 0} \omega_1(\epsilon) = 0$.

Mit (8.10) erhalten wir insbesondere

$$|\phi(x) - \lambda_a(x)|_\infty \leq \omega_1(|x - a|_\infty)|x - a|_\infty \quad (8.12)$$

für alle $a \in K$ und $x \in \Omega$ mit $|x - a|_\infty \leq \epsilon_1$, d.h. $x \in Q(a, \epsilon_1)$.

Schritt 3. Wir erinnern uns an die Funktionen ψ und Ψ . Wir zeigen zunächst durch Induktion, daß $|\Psi(x) - \Psi(x')| \leq n|x - x'|_\infty$ für alle $x, x' \in \mathbb{R}^n$. Für $n = 1$ ist $\Psi = \psi$ und die

⁷Man bezeichnet eine solche Funktion auch als *Stetigkeitsmodul* für $D\phi$. Im übrigen verwenden wir innerhalb dieses Beweises stets die Matrixnorm $|A| := \sup\{|Ax|_\infty : |x|_\infty \leq 1\}$.

Behauptung folgt einfach aus der Definition 8.4. Ist dagegen für ein bestimmtes n bereits $|\Psi(x) - \Psi(x')| \leq n|x - x'|_\infty$ für alle $x, x' \in \mathbb{R}^n$ gezeigt, dann ist

$$\begin{aligned} & |\Psi(x_1, \dots, x_n, x_{n+1}) - \Psi(x'_1, \dots, x'_n, x'_{n+1})| \\ &= |\psi(x_{n+1})\Psi(x_1, \dots, x_n) - \psi(x'_{n+1})\Psi(x'_1, \dots, x'_n)| \\ &\leq |\psi(x_{n+1})|\|\Psi(x_1, \dots, x_n) - \Psi(x'_1, \dots, x'_n)\| + |\psi(x_{n+1}) - \psi(x'_{n+1})|\|\Psi(x'_1, \dots, x'_n)\| \\ &\leq n|(x_1, \dots, x_n) - (x'_1, x'_2, \dots, x'_n)|_\infty + |x_{n+1} - x'_{n+1}|_\infty \\ &\leq (n+1)|(x_1, \dots, x_{n+1}) - (x'_1, x'_2, \dots, x'_{n+1})|_\infty, \end{aligned}$$

wobei wir $|\psi|, |\Psi| \leq 1$ verwendet haben. Daraus folgt nun unmittelbar

$$|T_b \Psi_\epsilon(x) - T_b \Psi_\epsilon(x')| \leq \frac{n}{\epsilon}|x - x'|_\infty \quad (8.13)$$

für jedes $b \in \mathbb{R}^n$ und $\epsilon > 0$.

Schritt 4. Sei $\epsilon_2 := \frac{\epsilon_1}{C}$. In diesem Beweisschritt zeigen wir: Es existiert ein nichtnegatives monotonen $\omega_2 : [0, \epsilon_2] \rightarrow \mathbb{R}$ mit $\lim_{\epsilon \rightarrow 0} \omega_2(\epsilon) = 0$, sodaß für alle $b \in \mathbb{R}^n$ und $0 < \epsilon \leq \epsilon_2$ die Funktion $h := T_b \Psi_\epsilon$ folgende Abschätzung erfüllt:

$$\left| \int_{\Omega} h(\phi(x)) |\det D\phi(x)| dx - \int_{\phi(\Omega)} h(y) dy \right| \leq \omega_2(\epsilon) \epsilon^n. \quad (8.14)$$

Sei dazu $a := \phi^{-1}(b)$. Aus (8.12) und (8.13) folgt

$$|h(\phi(x)) - h(\lambda_a(x))| \leq \frac{n}{\epsilon} \omega_1(|x - a|_\infty) |x - a|_\infty \quad (8.15)$$

für alle $x \in \Omega$ mit $|x - a|_\infty \leq \epsilon_1$. Nach (8.9) ist aber

$$\text{supp}(h \circ \phi) = \text{supp}(T_b \Psi_\epsilon \circ \phi) \subset Q(a, C\epsilon) \subset Q(a, \epsilon_1) \subset K',$$

und in ähnlicher Weise haben wir wegen der Definition von λ_a und (8.7)

$$\text{supp}(h \circ \lambda_a) = \text{supp}(T_b \Psi_\epsilon \circ \lambda_a) \subset Q(a, C\epsilon) \subset Q(a, \epsilon_1) \subset K'.$$

Insbesondere ist $|x - a|_\infty \leq C\epsilon$ auf dem Träger von $|h \circ \phi - h \circ \lambda_a|$, sodaß aus der Monotonie von ω_1 aus (8.15) folgt

$$|h(\phi(x)) - h(\lambda_a(x))| \leq \frac{n}{\epsilon} \omega_1(C\epsilon) C\epsilon = nC\omega_1(C\epsilon)$$

für alle $x \in \Omega$.

Die Funktion $x \mapsto |\det D\phi(x)|$ ist auf K' (und damit auf dem Träger von $|h \circ \phi - h \circ \lambda_a|$) gleichmäßig stetig. Wie in Schritt 2 folgt daraus die Existenz einer nichtnegativen monotonen Funktion $\omega_3 : [0, \epsilon_1] \rightarrow \mathbb{R}$ mit $\lim_{\epsilon \rightarrow 0} \omega_3(\epsilon) = 0$, sodaß

$$||\det D\phi(x)| - |\det D\phi(a)|| \leq \omega_3(|x - a|_\infty)$$

für alle $x \in Q(a, \epsilon_1)$. Daraus ergibt sich die Abschätzung

$$\begin{aligned} & |h(\phi(x))|\det D\phi(x) - h(\lambda_a(x))|\det D\phi(a)|| \\ &\leq |\det D\phi(x)| |h(\phi(x)) - h(\lambda_a(x))| + |h(\lambda_a(x))| ||\det D\phi(x)| - |\det D\phi(a)|| \\ &\leq MnC\omega_1(C\epsilon) + \omega_3(C\epsilon) =: \tilde{\omega}(\epsilon), \end{aligned}$$

wo wir $M := \max_{x \in K'} |\det D\phi(x)|$ gesetzt haben und $\tilde{\omega} : [0, \epsilon_2] \rightarrow \mathbb{R}$ wieder eine monotone nichtnegative Funktion mit $\lim_{\epsilon \rightarrow 0} \tilde{\omega}(\epsilon) = 0$ ist.

Integration über diese Ungleichung ergibt schließlich

$$\begin{aligned} & \left| \int_{\Omega} h(\phi(x)) |\det D\phi(x)| dx - \int_{\Omega} h(\lambda_a(x)) |\det D\phi(a)| dx \right| \\ &\leq \int_{Q(a, C\epsilon)} \tilde{\omega}(\epsilon) dx = \tilde{\omega}(\epsilon) (2C\epsilon)^n. \end{aligned}$$

Nach Satz 8.19 ist aber

$$\int_{\Omega} h(\lambda_a(x)) |\det D\phi(a)| dx = \int_{\phi(\Omega)} h(y) dy,$$

was (8.14) mit $\omega_2 := (2C)^n \tilde{\omega}$ beweist.

Schritt 5. Nach Lemma 8.14 konvergiert die Funktionenfolge

$$f_{\epsilon}(y) := \sum_{\alpha \in \mathbb{Z}^n} f(\epsilon\alpha) T_{\epsilon\alpha} \Psi_{\epsilon}(y)$$

mit $\epsilon \rightarrow 0$ gleichmäßig gegen f . Außerdem ist $\text{supp } f_{\epsilon} \subset L'$, sofern $\epsilon \leq \epsilon_1$. Daher ist nach Lemma 8.12

$$\lim_{\epsilon \rightarrow 0} \int_{\phi(\Omega)} f_{\epsilon}(y) dy = \int_{\phi(\Omega)} f(y) dy.$$

Mit der Notation

$$g(x) := f(\phi(x)) |\det D\phi(x)|$$

und

$$g_{\epsilon}(x) := f_{\epsilon}(\phi(x)) |\det D\phi(x)| = \sum_{\alpha \in \mathbb{Z}^n} f(\epsilon\alpha) T_{\epsilon\alpha} \Psi_{\epsilon}(\phi(x)) |\det D\phi(x)| dx$$

haben wir die gleichmäßige Konvergenz $g_{\epsilon} \rightarrow g$ (da $f_{\epsilon} \rightarrow f$ gleichmäßig), und daher gilt

$$\lim_{\epsilon \rightarrow 0} \int_{\Omega} g_{\epsilon}(x) dx = \int_{\Omega} g(x) dx.$$

Zu zeigen ist $\int_{\Omega} g(x) dx = \int_{\phi(\Omega)} f(y) dy$, also genügt es zu beweisen, daß

$$\Delta_{\epsilon} := \int_{\Omega} g_{\epsilon}(x) dx - \int_{\phi(\Omega)} f_{\epsilon}(y) dy$$

gegen null konvergiert, wenn $\epsilon \rightarrow 0$.

Mit

$$A_{\epsilon\alpha} := \int_{\Omega} T_{\epsilon\alpha} \Psi_{\epsilon}(\phi(x)) |\det D\phi(x)| dx - \int_{\phi(\Omega)} T_{\epsilon\alpha} \Psi_{\epsilon}(y) dy$$

gilt $\Delta_{\epsilon} = \sum_{\alpha \in \mathbb{Z}^n} f(\epsilon\alpha) A_{\epsilon\alpha}$. Gemäß (8.14) ist $|A_{\epsilon\alpha}| \leq \omega_2(\epsilon) \epsilon^n$.

Da der Träger L von f kompakt ist, gibt es einen achsenparallelen Würfel der Seitenlänge s , der L enthält, und der nicht mehr als $\left(\frac{s}{\epsilon} + 1\right)^n$ Punkte der Form $\epsilon\alpha$, $\alpha \in \mathbb{Z}^n$ enthält. Daher ist die Zahl der nichtverschwindenden Summanden in $\sum_{\alpha \in \mathbb{Z}^n} f(\epsilon\alpha) A_{\epsilon\alpha}$ durch $\left(\frac{s}{\epsilon} + 1\right)^n$ nach oben abgeschätzt, und wir erhalten

$$|\Delta_{\epsilon}| \leq \left(\frac{s}{\epsilon} + 1\right)^n \sup_{y \in L} |f(y)| \omega_2(\epsilon) \epsilon^n = (s + \epsilon)^n \sup_{y \in L} |f(y)| \omega_2(\epsilon) \rightarrow 0$$

mit $\epsilon \rightarrow 0$, und die Transformationsformel ist bewiesen. $\square$

8.3. Integration stetiger Funktionen auf kompakten Mengen

Eine Schwäche unserer bisherigen Integrationstheorie stetiger Funktionen ist die Anforderung, eine Funktion müsse kompakten Träger in $\mathbb{R}^n$ haben. Betrachte aber folgendes Problem: Gegeben eine kompakte Menge $K \subset \mathbb{R}^n$, wie kann ihr Volumen berechnet werden? In einer Dimension kann die Länge eines kompakten Intervalls $[a, b]$ dargestellt werden als $\int_{\mathbb{R}} \chi_{[a, b]}(x) dx$, wobei $\chi_{[a, b]} : \mathbb{R} \rightarrow \mathbb{R}$ die *charakteristische Funktion* des Intervalls $[a, b]$ ist, definiert durch

$$\chi_{[a, b]}(x) := \begin{cases} 1 & \text{falls } x \in [a, b], \\ 0 & \text{sonst.} \end{cases}$$

Analog würden wir für $K \subset \mathbb{R}^n$ das Volumen gerne als $\int_K \chi_K(x) dx$ berechnen, wobei

$$\chi_K(x) := \begin{cases} 1 & \text{falls } x \in K, \\ 0 & \text{sonst.} \end{cases} \quad (8.16)$$

Das Problem ist nun, daß χ_K nicht in $C_c(\mathbb{R}^n)$ ist (sie ist nämlich unstetig in ∂K). Wir werden daher in diesem Abschnitt Integrale der Form $\int_K f(x) dx$ einführen, wobei $f: \mathbb{R}^n \rightarrow \mathbb{R}$ stetig ist.

8.3.1. Definition des Integrals stetiger Funktionen auf kompakten Mengen.

DEFINITION 8.22. Sei $f: \mathbb{R}^n \rightarrow \mathbb{R}$ eine Funktion dergestalt, daß eine Folge $(f_k)_{k \in \mathbb{N}} \subset C_c(\mathbb{R}^n)$ existiert, sodaß $f_k \geq f_{k+1}$ für alle $k \in \mathbb{N}$, und $f_k \rightarrow f$ punktweise für $k \rightarrow \infty$. Dann existiert

$$\int_{\mathbb{R}^n} f(x) dx := \lim_{k \rightarrow \infty} \int_{\mathbb{R}^n} f_k(x) dx$$

als reelle Zahl oder $-\infty$.

In der Tat, die Existenz des Limes folgt aus der Monotonie der Integrale $\int_{\mathbb{R}^n} f_k(x) dx$ (da die Folge (f_k) monoton fällt) und Korollar 2.30.

BEMERKUNG 8.23. Es stellt sich heraus, daß die Funktionen, die auf die in dieser Definition verlangte Weise approximiert werden können, genau die *oberhalbstetigen* Funktionen sind, deren Positivteil kompakten Träger in $\mathbb{R}^n$ hat (siehe dazu [10], §4). Eine Funktion heißt oberhalbstetig im Punkt $x^0 \in \mathbb{R}^n$, wenn $\limsup_{x \rightarrow x^0} f(x) \leq f(x^0)$.

Wir müssen zeigen, daß diese Definition unabhängig von der Wahl der approximierenden Folge ist. Dies wird aus dem Satz von DINI folgen:

SATZ 8.24 (DINI). Sei $f \in C_c(\mathbb{R}^n)$ und $(f_k)_{k \in \mathbb{N}} \subset C_c(\mathbb{R}^n)$ eine Folge, sodaß $f_k \geq f_{k+1}$ für alle $k \in \mathbb{N}$, und $f_k \rightarrow f$ punktweise für $k \rightarrow \infty$. Dann gilt

$$\int_{\mathbb{R}^n} f(x) dx = \lim_{k \rightarrow \infty} \int_{\mathbb{R}^n} f_k(x) dx.$$

BEWEIS. Aufgrund der Monotonie gilt $\text{supp}(f_k) \subset K := \text{supp}(f_0) \cup \text{supp}(f)$ für alle $k \in \mathbb{N}$, denn: Ist $f_k(x) \neq 0$, so ist wegen $f_0(x) \geq f_k(x) \geq f(x)$ mindestens einer der Werte $f_0(x)$ oder $f(x)$ ungleich null. Als endliche Vereinigung kompakter Mengen ist K wieder kompakt.

Setze nun $g_k := f_k - f$, dann ist die Folge (g_k) monoton fallend und nichtnegativ, also $g_k \geq g_{k+1} \geq 0$ für alle $k \in \mathbb{N}$. Beachte, daß $\text{supp}(g_k) \subset K$ für alle $k \in \mathbb{N}$. Zudem haben wir $g_k \rightarrow 0$ punktweise, das heißt, zu jedem $\epsilon > 0$ und $x \in K$ existiert $N(x) \in \mathbb{N}$, sodaß

$$g_k(x) < \frac{\epsilon}{2}$$

für alle $k \geq N(x)$. Weiterhin ist $g_{N(x)}$ stetig, also existiert $\delta(x) > 0$, sodaß aus $|x - y| < \delta(x)$ folgt

$$|g_{N(x)}(x) - g_{N(x)}(y)| < \frac{\epsilon}{2}.$$

Wir folgern $g_{N(x)}(y) < \epsilon$ für alle $y \in B_{\delta(x)}(x)$, und wegen der Monotonie folgt dann auch $g_k < \epsilon$ für alle $k \geq N(x)$ und $y \in B_{\delta(x)}(x)$.

Nun ist $K \subset \bigcup_{x \in K} B_{\delta(x)}(x)$, und da K kompakt ist, gibt es endlich viele $x^1, \dots, x^m$, sodaß $K \subset \bigcup_{j=1}^m B_{\delta(x^j)}(x^j)$. Sei $N := \max_{j=1, \dots, m} N(x^j)$, so gilt

$$g_k(x) < \epsilon$$

für alle $x \in K$ und alle $k \geq N$, und daher konvergiert $g_k \rightarrow 0$ sogar gleichmäßig. Aus der Definition von g_k und Lemma 8.12 folgt nun sofort die Behauptung. $\square$

Mit dem Satz von DINI können wir nun auf folgende Weise zeigen, daß das Integral in Definition 8.22 unabhängig von der gewählten Approximation ist. Sei also $f : \mathbb{R}^n \rightarrow \mathbb{R}$ beliebig und $(f_k)_{k \in \mathbb{N}}, (g_k)_{k \in \mathbb{N}}$ zwei Folgen in $C_c(\mathbb{R}^n)$, die in monoton fallender Weise punktweise gegen f konvergieren. Es genügt zu zeigen, daß für jedes $j \in \mathbb{N}$ gilt

$$\int_{\mathbb{R}^n} f_j(x) dx \geq \lim_{k \rightarrow \infty} \int_{\mathbb{R}^n} g_k(x) dx, \quad (8.17)$$

denn dann gilt dies auch im Limes $j \rightarrow 0$, und die andere Ungleichung folgt durch Vertauschung der Rollen von f_j und g_k .

Um 8.17 zu zeigen, sei $h_k := \max\{f_j, g_k\}$. Da $f_j \geq f$ und $g_k \rightarrow f$ für $k \rightarrow \infty$, folgt $h_k \rightarrow f_j$ in monoton fallender Weise, und da $f_j \in C_c(\mathbb{R}^n)$, folgt aus dem Satz von DINI

$$\int_{\mathbb{R}^n} f_j(x) dx = \lim_{k \rightarrow \infty} \int_{\mathbb{R}^n} h_k(x) dx \geq \lim_{k \rightarrow \infty} \int_{\mathbb{R}^n} g_k(x) dx,$$

also (8.17).

Man beachte ferner: Ist $f \in C_c(\mathbb{R}^n)$, so kann man in Definition 8.22 einfach $f_k = f$ setzen und erhält somit, daß Definition 8.22 konsistent mit dem vorherigen Integralbegriff auf $C_c(\mathbb{R}^n)$ ist.

Wir verzichten hier auf die in Bemerkung 8.23 erwähnte Charakterisierung derjenigen Funktionen, die gemäß Definition 8.22 integrierbar sind, und zeigen nur, daß eine interessante Klasse von Funktionen unter die Definition fällt:

SATZ 8.25. Sei $K \subset \mathbb{R}^n$ kompakt und $f : \mathbb{R}^n \rightarrow \mathbb{R}$ nichtnegativ und stetig. Dann existiert eine monoton fallende Folge in $C_c(\mathbb{R}^n)$, die punktweise gegen $f\chi_K$ konvergiert, wobei χ_K die charakteristische Funktion von K wie in (8.16) ist.

BEWEIS. Wir definieren die *Abstandsfunktion* zu K durch $d_K : \mathbb{R}^n \rightarrow \mathbb{R}$, $d_K(x) := \min\{|x - y| : y \in K\}$. Diese Funktion ist wohldefiniert (da zu jedem $x \in \mathbb{R}^n$ die stetige Funktion $y \mapsto |x - y|$ ihr Minimum auf dem Kompaktum K annimmt). Sie ist aber auch (gleichmäßig) stetig: Sei $\epsilon > 0$ und seien $x, x' \in \mathbb{R}^n$ mit $|x - x'| < \epsilon$. O.B.d.A. nehmen wir an $d_K(x) \geq d_K(x')$. Sei außerdem $y \in K$ ein Punkt mit $|x' - y| = d_K(x')$. Dann gilt

$$|d_K(x) - d_K(x')| = d_K(x) - d_K(x') \leq |x - y| - |x' - y| \leq |x - x'| < \epsilon.$$

Sei außerdem für $k \in \mathbb{N}$ die Funktion $\eta_k : \mathbb{R}_0^+ \rightarrow \mathbb{R}$ gegeben durch

$$\eta_k(t) := \begin{cases} 1 - kt & \text{falls } 0 \leq t \leq \frac{1}{k}, \\ 0 & \text{falls } t > \frac{1}{k}. \end{cases}$$

Dann ist η_k stetig, und die Funktion

$$f_k(x) := f(x)\eta_k(d_K(x))$$

ist als Produkt bzw. Verknüpfung stetiger Funktionen wieder stetig. Wegen $f \geq 0$ und $\eta_k \geq \eta_{k+1}$ ist außerdem $f_k \geq f_{k+1}$ für alle $k \in \mathbb{N}$. Schließlich bemerken wir $f_k = f$ in K und $\lim_{k \rightarrow \infty} f_k(x) = 0$ für jedes $x \notin K$, sodaß in der Tat $(f_k)_{k \in \mathbb{N}}$ eine Folge wie gewünscht ist. $\square$

Dieser Satz erlaubt uns also, für eine stetige nichtnegative Funktion $\mathbb{R}^n \rightarrow \mathbb{R}$ das Integral über beliebige kompakte Mengen zu definieren.

BEMERKUNG 8.26. Sei $K \subset \mathbb{R}^n$ kompakt. Es ist nicht schwer zu zeigen (Übung), daß jede stetige Funktion $f : K \rightarrow \mathbb{R}$ eine stetige Fortsetzung $\tilde{f} : \mathbb{R}^n \rightarrow \mathbb{R}$ besitzt, d.h. $\tilde{f}$ ist auf ganz $\mathbb{R}^n$ stetig und $\tilde{f}|_K = f$. Daraus folgt, daß man das Integral $\int_K f(x) dx$ wie oben für jede stetige nichtnegative Funktion $f : K \rightarrow \mathbb{R}$ definieren kann; f muß dazu also nicht auf ganz $\mathbb{R}^n$ definiert sein.⁸

⁸Für eine große Klasse topologischer Räume (insbesondere für alle metrischen Räume) kann man zeigen, daß jede stetige Funktion auf einer abgeschlossenen Menge eine stetige Fortsetzung auf den ganzen

Nun beseitigen wir noch die Einschränkung $f \geq 0$. Sei dazu $f: \mathbb{R}^n \rightarrow \mathbb{R}$ stetig (aber nicht zwingend nichtnegativ). Dann sind *Positiv-* und *Negativteil* von f definiert als $f^\pm: \mathbb{R}^n \rightarrow \mathbb{R}$,

$$f^+(x) := \begin{cases} f(x) & \text{falls } f(x) \geq 0, \\ 0 & \text{falls } f(x) < 0; \end{cases}$$

$$f^-(x) := \begin{cases} -f(x) & \text{falls } f(x) \leq 0, \\ 0 & \text{falls } f(x) > 0. \end{cases}$$

Offenbar ist dann $f = f^+ - f^-$, und f^+ und f^- sind wieder stetig. Dies motiviert folgende Definition:

DEFINITION 8.27. Sei $K \subset \mathbb{R}^n$ kompakt und $f: \mathbb{R}^n \rightarrow \mathbb{R}$ stetig. Dann definieren wir

$$\int_K f(x) dx := \int_K f^+(x) dx - \int_K f^-(x) dx.$$

Auch diese Definition ist konsistent mit den bisherigen Integraldefinitionen⁹. Wie oben bemerken wir, daß es ausreicht, wenn f nur auf K definiert ist.

Als Übung zeige man, daß für jedes kompakte $K \subset \mathbb{R}^n$ das Funktional $I: C(K) \rightarrow \mathbb{R}$, $I(f) = \int_K f(x) dx$ linear und monoton ist.

Wir nennen eine kompakte Menge $K \subset \mathbb{R}^n$ *Nullmenge*, wenn $\int_K 1 dx = 0$. Zur Übung zeige man, daß etwa $\{0\} \subset \mathbb{R}$ oder $[-1, 1] \times \{0\} \subset \mathbb{R}^2$ Nullmengen sind. Die leere Menge ist selbstverständlich auch eine Nullmenge.

PROPOSITION 8.28. Seien $K_1, K_2 \subset \mathbb{R}^n$ kompakt, sodaß $K_1 \cap K_2$ eine Nullmenge ist. Dann gilt für jedes stetige $f: K_1 \cup K_2 \rightarrow \mathbb{R}$

$$\int_{K_1 \cup K_2} f(x) dx = \int_{K_1} f(x) dx + \int_{K_2} f(x) dx.$$

BEWEIS. Nach Bemerkung 8.26 dürfen wir annehmen, daß f sogar auf ganz $\mathbb{R}^n$ definiert ist.

Wir nehmen außerdem an $f \geq 0$, da die Aussage für allgemeines f einfach durch Zerlegung in Positiv- und Negativteil folgt.

Mit der Notation wie im Beweis von Satz 8.25 setzen wir

$$f_k(x) := f(x)\eta_k(d_{K_1}(x)), \quad g_k(x) := f(x)\eta_k(d_{K_2}(x))$$

für $k \in \mathbb{N}$, sodaß $\int_{\mathbb{R}^n} f_k dx \rightarrow \int_{K_1} f(x) dx$ und $\int_{\mathbb{R}^n} g_k dx \rightarrow \int_{K_2} f(x) dx$ mit $k \rightarrow \infty$.

Wir behaupten, daß $h_k := f_k + g_k$ eine monoton fallende Folge definiert, die punktweise gegen die Funktion

$$\tilde{f}(x) := \begin{cases} f(x) & \text{falls } x \in (K_1 \setminus K_2) \cup (K_2 \setminus K_1), \\ 2f(x) & \text{falls } x \in K_1 \cap K_2, \\ 0 & \text{sonst} \end{cases}$$

konvergiert. Die Monotonie folgt sofort aus der von f_k und g_k . Für die Konvergenz beachte $f_k(x) = f(x)$ für $x \in K_1$, $g_k(x) = f(x)$ für $x \in K_2$, $f_k(x) \rightarrow 0$ für $x \in K_1^c$ und $g_k(x) \rightarrow 0$ für $x \in K_2^c$.

Es folgt nach Definition 8.22 $\int_{\mathbb{R}^n} (f_k + g_k) dx \rightarrow \int_{K_1 \cup K_2} \tilde{f}(x) dx$, sodaß zu zeigen bleibt

$$\int_{K_1 \cup K_2} \tilde{f}(x) dx = \int_{K_1 \cup K_2} f(x) dx.$$

Raum besitzt. Dieses (nichttriviale) Resultat aus der elementaren Topologie ist als Satz von Tietze-Urysohn bekannt.

⁹Insbesondere überlege man sich zur Übung: Ist $K = \Pi_{k=1}^n [a_k, b_k]$, so stimmt das hier definierte Integral mit dem in Abschnitt 8.1.2 definierten Mehrfachintegral überein.

In der Tat, unter Beachtung von $f \geq 0$ gilt:

$$\begin{aligned} & \left| \int_{K_1 \cup K_2} \tilde{f}(x) dx - \int_{K_1 \cup K_2} f(x) dx \right| \\ &= \int_{K_1 \cup K_2} (\tilde{f}(x) - f(x)) dx \\ &= \int_{K_1 \cap K_2} f(x) dx \\ &\leq \sup_{x \in K_1 \cap K_2} |f(x)| \int_{K_1 \cap K_2} 1 dx = 0, \end{aligned}$$

da nach Voraussetzung $K_1 \cap K_2$ eine Nullmenge ist. $\square$

8.3.2. Rotationssymmetrische Funktionen.

DEFINITION 8.29. Seien ρ, R reelle Zahlen mit $0 \leq \rho < R$ und $A_{\rho, R} := \{x \in \mathbb{R}^n : \rho \leq |x| \leq R\}$. Eine Funktion $f : A_{\rho, R} \rightarrow \mathbb{R}$ heißt *rotationssymmetrisch*, wenn es eine Funktion $h : [\rho, R] \rightarrow \mathbb{R}$ gibt, sodaß $f(x) = h(|x|)$.

Für $n = 2$ entsteht der Graph von f durch Rotation des Graphen von h um die vertikale Achse.

SATZ 8.30 (Integration rotationssymmetrischer Funktionen). Sei $f : A_{\rho, R} \rightarrow \mathbb{R}$ stetig und rotationssymmetrisch mit $f(x) = h(|x|)$. Dann gilt

$$\int_{A_{\rho, R}} f(x) dx = n\omega_n \int_{\rho}^R r^{n-1} h(r) dr,$$

wobei ω_n das Volumen der n -dimensionalen Einheitskugel ist.

BEMERKUNG 8.31. Für $n = 2$ oder $n = 3$ folgt die Aussage im wesentlichen aus Übungsaufgabe 2, Blatt 11.

BEWEIS. Betrachte die äquidistante Zerlegung des Intervalls $[\rho, R]$ in N Teilintervalle:

$$r_k = \rho + \frac{k}{N}(R - \rho),$$

und setze

$$A_k := A_{r_{k-1}, r_k}, \quad k = 1, \dots, N.$$

Nun gilt $\int_{\mathbb{R}^n} \chi_{A_k}(x) dx = \omega_n(r_k^n - r_{k-1}^n)$, denn $A_k \cup \overline{B_{r_{k-1}}} = \overline{B_{r_k}}$, der Schnitt dieser beiden Mengen ist die Nullmenge (!) $\{x \in \mathbb{R}^n : |x| = r_{k-1}\}$, und das Volumen der Kugeln beträgt bekanntlich¹⁰ $\omega_n r_{k-1}^n$ bzw. $\omega_n r_k^n$; die Formel für das Volumen von A_k folgt dann aus Proposition 8.28.

Nach dem (eindimensionalen) Mittelwertsatz und wegen $\frac{d}{dr} r^n = n r^{n-1}$ existiert zu jedem $k = 1, \dots, N$ ein $\xi_k \in [r_{k-1}, r_k]$, sodaß

$$\int_{\mathbb{R}^n} \chi_{A_k}(x) dx = n\omega_n \xi_k^{n-1} (r_k - r_{k-1}).$$

Wir setzen

$$f_N := \sum_{k=1}^N h(\xi_k) \chi_{A_k}$$

¹⁰Bei „bekanntlich“ ist bekanntlich stets Vorsicht geboten! In diesem Falle ist prima facie nicht klar, ob unsere Volumenberechnung der Kugel als Integral über $\sqrt{1 - x_1^2 - \dots - x_{n-1}^2}$ mit dem Integral $\int_{\mathbb{R}^n} \chi_{B_r(0)} dx$ übereinstimmt. Wir zeigen dies in Kürze mithilfe des Satzes von FUBINI.

und schreiben kurzerhand¹¹

$$\int_{\mathbb{R}^n} f_N(x) dx := \sum_{k=1}^N h(\xi_k) \int_{\mathbb{R}^n} \chi_{A_k}(x) dx = n\omega_n \sum_{k=1}^N h(\xi_k) \xi_k^{n-1} (r_k - r_{k-1}).$$

Dieser Ausdruck ist allerdings eine RIEMANN-Summe (siehe Satz 4.27) für das Integral $n\omega_n \int_{\rho}^R h(r)r^{n-1}dr$, also haben wir die Konvergenz

$$\lim_{N \rightarrow \infty} \int_{\mathbb{R}^n} f_N(x) dx = n\omega_n \int_{\rho}^R h(r)r^{n-1}dr. \quad (8.18)$$

Andererseits gilt aufgrund der gleichmäßigen Stetigkeit von f auf $A_{\rho,R}$, daß zu jedem $\epsilon > 0$ ein $N_0 \in \mathbb{N}$ existiert, sodaß

$$|f(x) - f_N(x)| < \epsilon \quad \forall x \in A_{\rho,R}, \forall N \geq N_0.$$

Nach Proposition 8.28 gilt

$$\int_{\rho}^R f(x) dx = \sum_{k=1}^N \int_{A_k} f(x) dx,$$

und wegen der Monotonie des Integrals gilt für jedes k und für $N \geq N_0$

$$\int_{A_k} f(x) dx - \epsilon \int_{\mathbb{R}^n} \chi_{A_k}(x) dx \leq \int_{A_k} f_N(x) dx \leq \int_{A_k} f(x) dx + \epsilon \int_{\mathbb{R}^n} \chi_{A_k}(x) dx. \quad (8.19)$$

Da aber $\sum_{k=1}^N \int_{\mathbb{R}^n} \chi_{A_k}(x) dx = \omega_n(R^n - \rho^n)$, was insbesondere nicht von N abhängt, folgt aus (8.19) nach Summation über alle k

$$\lim_{N \rightarrow \infty} \int_{\mathbb{R}^n} f_N(x) dx = \int_{A_{\rho,R}} f(x) dx.$$

Zusammen mit (8.18) folgt die Behauptung. $\square$

8.3.3. Volumenberechnungen.

8.3.3.1. *Ein paar Werkzeuge.* Ein Schlüssel zur Berechnung vieler Volumina ist der folgende Satz, der – wie bei Mehrfachintegralen (Satz 8.4) – die Vertauschung der Integrationsreihenfolge erlaubt. Genauer gilt:

SATZ 8.32 (FUBINI). Sei $f : \mathbb{R}^n \rightarrow \mathbb{R}$ eine Funktion, die (wie in Definition 8.22) monotoner punktwiser Limes einer Folge $(f_k)_{k \in \mathbb{N}} \subset C_c(\mathbb{R}^n)$ ist. Sei $1 \leq m \leq n$ und schreibe $x = (\tilde{x}, x')$ für einen Vektor $x = (x_1, \dots, x_n) \in \mathbb{R}^n$, wobei $\tilde{x} = (x_1, \dots, x_m) \in \mathbb{R}^m$ und $x' = (x_{m+1}, \dots, x_n) \in \mathbb{R}^{n-m}$.

Dann ist die Funktion

$$F : \mathbb{R}^{n-m} \rightarrow \mathbb{R}, \quad F(x') := \int_{\mathbb{R}^m} f(\tilde{x}, x') d\tilde{x}$$

ebenfalls monotoner punktwiser Limes von Funktionen $(g_k)_{k \in \mathbb{N}} \subset C_c(\mathbb{R}^{n-m})$, und es gilt

$$\int_{\mathbb{R}^{n-m}} F(x') dx' = \int_{\mathbb{R}^n} f(x) dx. \quad (8.20)$$

BEWEIS. Sei $(f_k)_{k \in \mathbb{N}} \subset C_c(\mathbb{R}^n)$ eine Folge, die monoton gegen f konvergiert. Wir setzen für jedes $k \in \mathbb{N}$ und jedes $x' \in \mathbb{R}^{n-m}$

$$g_k(x') := \int_{\mathbb{R}^m} f_k(\tilde{x}, x') d\tilde{x}.$$

Beachte dazu, daß für fest gewähltes $x' \in \mathbb{R}^{n-m}$ die Funktion $\tilde{x} \mapsto f_k(\tilde{x}, x')$ in $C_c(\mathbb{R}^m)$ ist, also ist g_k wohldefiniert. Nach Satz 8.1 ist g_k wieder stetig, und sein Träger ist kompakt, da er in der beschränkten Menge $\{x' \in \mathbb{R}^{n-m} : \exists \tilde{x} \in \mathbb{R}^k \quad (\tilde{x}, x') \in \text{supp}(f_k)\}$ enthalten ist.

¹¹Weil wir nämlich das Integral über unstetige Funktionen nicht definiert haben.

Aufgrund der Monotonie des Integrals und der Folge (f_k) ist außerdem (g_k) monoton. Nun gilt aber für jedes $x' \in \mathbb{R}^{m-n}$

$$g_k(x') = \int_{\mathbb{R}^m} f_k(\tilde{x}, x') d\tilde{x} \rightarrow \int_{\mathbb{R}^m} f(\tilde{x}, x') d\tilde{x} = F(x'),$$

wobei wir für den Grenzübergang Definition 8.22 für festes x' auf die Folge $(f_k(\cdot, x'))_{k \in \mathbb{N}}$ angewendet haben, die ja monoton gegen $f(\cdot, x')$ konvergiert. Damit ist bereits gezeigt, daß F monotoner punktweiser Limes der Funktionen $(g_k)_{k \in \mathbb{N}}$ ist.

Für (8.20) rechnen wir, die monotone Konvergenz $g_k \rightarrow F$ sowie Definition 8.22 ausnützend,

$$\begin{aligned} \int_{\mathbb{R}^{n-m}} F(x') dx' &= \lim_{k \rightarrow \infty} \int_{\mathbb{R}^{n-m}} g_k(x') dx' \\ &= \lim_{k \rightarrow \infty} \int_{\mathbb{R}^{n-m}} \int_{\mathbb{R}^m} f_k(\tilde{x}, x') d\tilde{x} dx' \\ &= \lim_{k \rightarrow \infty} \int_{\mathbb{R}^n} f_k(x) dx \\ &= \int_{\mathbb{R}^n} f(x) dx, \end{aligned}$$

wobei wir im vorletzten Schritt die Definition des Integrals auf $C_c(\mathbb{R}^n)$ als Mehrfachintegral verwendet haben. Damit ist alles gezeigt. $\square$

BEMERKUNG 8.33. Mit einem sehr ähnlichen Beweis (man muß in der letzten Rechnung im Integral über f_k nur die Integrationsreihenfolge gemäß Satz 8.4 die Integrationsreihenfolge vertauschen) kann man zeigen, daß für Funktionen, die monotone Limites von $C_c(\mathbb{R}^n)$ -Folgen sind, die Integrationsreihenfolge beliebig gewählt werden darf.

Sei $K \subset \mathbb{R}^n$ kompakt, so bezeichnen wir mit $\text{Vol}(K) := \int_K 1 dx = \int_{\mathbb{R}^n} \chi_K(x) dx$ das n -dimensionale *Volumen* von K . Für $n = 1$ spricht man von *Länge*, für $n = 2$ von *Fläche*.

BEMERKUNG 8.34. Bei der Berechnung des Volumens einer Kugel haben wir die Halbkugel als Graphen einer Funktion dargestellt und diese Funktion integriert. Wie paßt das zusammen mit der gerade gegebenen Definition des Volumens? Mit dem Satz von FUBINI können wir folgendes allgemeine Prinzip zeigen:

Sei $K \subset \mathbb{R}^n$ kompakt und $f : K \rightarrow \mathbb{R}$ stetig und nichtnegativ. Sei

$$K_f := \{(x, y) \in \mathbb{R}^{n+1} : x \in K, y \in [0, f(x)]\}.$$

Dann gilt $\text{Vol}(K_f) = \int_K f(x) dx$. In der Tat, es ist $\chi_{K_f}(x, y) = \chi_K(x) \chi_{[0, f(x)]}(y)$ und daher nach FUBINI und der anschließenden Bemerkung

$$\begin{aligned} \text{Vol}(K_f) &= \int_{\mathbb{R}^n} \int_{\mathbb{R}} \chi_K(x) \chi_{[0, f(x)]}(y) dx dy \\ &= \int_{\mathbb{R}^n} \chi_K(x) \left(\int_{\mathbb{R}} \chi_{[0, f(x)]}(y) dy \right) dx \\ &= \int_{\mathbb{R}^n} f(x) \chi_K(x) dx \\ &= \int_K f(x) dx. \end{aligned}$$

Diese Art der Volumenberechnung haben wir nicht nur für das Kugelvolumen, sondern auch beim Beweis von Satz 8.30 verwendet. Nun ist gezeigt, daß dies gerechtfertigt war.

KOROLLAR 8.35 (Prinzip von CAVALIERI). Sei $K \subset \mathbb{R}^n$ kompakt und schreibe für festes $x' \in \mathbb{R}$

$$K_{x'} := \{\tilde{x} \in \mathbb{R}^{n-1} : (\tilde{x}, x') \in K\}.$$

Dann gilt

$$\text{Vol}(K) = \int_{\mathbb{R}} \text{Vol}(K_{x'}) dx'.$$

BEWEIS. Wende den Satz von FUBINI auf die charakteristische Funktion χ_K an, wobei $m = n - 1$ gewählt wird. $\square$

Das Prinzip besagt also, daß man ein Volumen ‚scheibchenweise‘ berechnen darf. Beachte, daß in dieser Formel Vol links das n -dimensionale und rechts das $n - 1$ -dimensionale Volumen bezeichnet.

KOROLLAR 8.36 (Volumen kartesischer Produkte). Sei $1 \leq m \leq n$ und seien $K_1 \subset \mathbb{R}^m$, $K_2 \subset \mathbb{R}^{n-m}$ kompakt. Dann gilt

$$\text{Vol}(K_1 \times K_2) = \text{Vol}(K_1) \text{Vol}(K_2).$$

BEWEIS. Wie im Satz von FUBINI schreiben wir $x = (\tilde{x}, x')$ für einen Vektor $x = (x_1, \dots, x_n) \in \mathbb{R}^n$, wobei $\tilde{x} = (x_1, \dots, x_m) \in \mathbb{R}^m$ und $x' = (x_{m+1}, \dots, x_n) \in \mathbb{R}^{n-m}$. Offenbar gilt $\chi_K(x) = \chi_{K_1}(\tilde{x})\chi_{K_2}(x')$ und daher mit dem Satz von FUBINI

$$\begin{aligned} \text{Vol}(K_1 \times K_2) &= \int_{K_1 \times K_2} \chi_K(x) dx \\ &= \int_{K_1 \times K_2} \chi_{K_1}(\tilde{x})\chi_{K_2}(x') dx \\ &= \int_{K_2} \left(\int_{K_1} \chi_{K_1}(\tilde{x})\chi_{K_2}(x') d\tilde{x} \right) dx' \\ &= \int_{K_2} \left(\int_{K_1} \chi_{K_1}(\tilde{x}) d\tilde{x} \right) \chi_{K_2}(x') dx' \\ &= \int_{K_2} \text{Vol}(K_1) \chi_{K_2}(x') dx' \\ &= \text{Vol}(K_1) \text{Vol}(K_2). \end{aligned}$$

$\square$

Schließlich erinnern wir uns an die Transformationsformel:

PROPOSITION 8.37. Sei $K \subset \mathbb{R}^n$ kompakt und $f : K \rightarrow \mathbb{R}$ stetig. Ist $A \in \mathbb{R}^{n \times n}$ invertierbar und $b \in \mathbb{R}^n$, so gilt¹²

$$|\det A| \int_{A^{-1}(K-b)} f(Ax+b) dx = \int_K f(x) dx.$$

BEWEIS. Nach Satz 8.25 und Definition 8.27 sind beide Integrale wohldefiniert als Limes von Integralen über $C_c(\mathbb{R}^n)$ -Funktionen. Für diese gilt die Aussage offenbar wegen der Translationsinvarianz des Integrals und Satz 8.19. Im Limes bleibt die Gleichheit bestehen. $\square$

Ein wichtiger Spezialfall ist die *Homothetie* $A = r \text{Id}$ mit $r > 0$, deren Determinante r^n beträgt, sodaß aus der Proposition $\text{Vol}(rK) = r^n \text{Vol}(K)$ folgt, was inzwischen nicht mehr überraschen dürfte. Hier haben wir wie üblich $rK := \{rx : x \in K\}$ geschrieben.

¹²Wir schreiben $A^{-1}(K-b) := \{A^{-1}(x-b) : x \in K\}$.

8.3.3.2. *Beispiele.*

- (1) Für einen Quader $Q = \prod_{k=1}^n [a_k, b_k]$ erhalten wir mit Korollar 8.36 (oder auch unmittelbar über das Mehrfachintegral):

$$\text{Vol}(Q) = \prod_{k=1}^n (b_k - a_k).$$

- (2) Sei $B \subset \mathbb{R}^{n-1}$ kompakt und $h \geq 0$. Dann heißt $Z := B \times [0, h] \subset \mathbb{R}^n$ *Zylinder* mit Basis B und Höhe h . Unmittelbar aus Korollar 8.36 folgt

$$\text{Vol}(Z) = h \text{Vol}(B).$$

- (3) Sei wieder $B \subset \mathbb{R}^{n-1}$ kompakt und $h \geq 0$. Dann heißt

$$K := \{((1-\lambda)\xi, \lambda h) : \xi \in B, \lambda \in [0, 1]\} \subset \mathbb{R}^n$$

Kegel mit Basis B und Höhe h . Zur Volumenberechnung ziehen wir das Prinzip von CAVALIERI heran, indem wir bemerken, daß für $x' \in \mathbb{R}$ gilt $K_{x'} := \{\tilde{x} \in \mathbb{R}^{n-1} : (\tilde{x}, x') \in K\} = \emptyset$ falls $x' < 0$ oder $x' > h$, und

$$K_{x'} = (1-\lambda)B := \{(1-\lambda)\tilde{x} : \tilde{x} \in B\},$$

sofern $x' = \lambda h$ für ein $\lambda \in [0, 1]$. Gemäß Bemerkung nach Proposition 8.37 ist $\text{Vol}(K_{x'}) = (1-\lambda)^{n-1} \text{Vol}(B)$, wenn $x' = \lambda h$ (beachte $B \subset \mathbb{R}^{n-1}$). Nach Prinzip von CAVALIERI gilt daher

$$\text{Vol}(K) = \int_0^h \text{Vol}(K_{x'}) dx' = h \text{Vol}(B) \int_0^1 (1-\lambda)^{n-1} d\lambda = h \text{Vol}(B) \frac{1}{n}.$$

Insbesondere ergibt sich für $n = 3$ und $B = B_r(0)$ die bekannte Formel $\text{Vol}(K) = \frac{1}{3}\pi r^2 h$.

- (4) Als *Parallelepipiped* (oder *Parallelotop*) bezeichnet man eine Menge der Form

$$K := \left\{ \sum_{k=1}^n \lambda_k v_k : \lambda_k \in [0, 1] \right\},$$

wobei $(v_k)_{k=1, \dots, n}$ eine Basis des $\mathbb{R}^n$ bilden. Bezeichne $A \in \mathbb{R}^{n \times n}$ die Matrix, deren k -te Spalte der Vektor v_k ist. A ist die darstellende Matrix (bzgl. der Standardbasis) derjenigen linearen Abbildung, die den Standardbasisvektor e_k auf v_k abbildet. Daher bildet A den Einheitswürfel $Q := [0, 1]^n$ auf das Parallelepipiped K ab, und nach Proposition 8.37 ist deshalb

$$\text{Vol}(K) = \int_K 1 dx = \int_{AQ} 1 dx = |\det A| \int_Q 1 dx = |\det A|.$$

Dies ist genau die geometrische Interpretation der Determinante: Ihr Betrag gibt die Volumenänderung einer linearen Abbildung an. (Ihr Vorzeichen gibt an, ob die Abbildung die Orientierung erhält oder umkehrt.) Sind $(v_k)_{k=1, \dots, n}$ nicht linear unabhängig, ist $\det A = 0$, und das Volumen des Parallelotops ist null, da es in einem Untervektorraum einer Dimension $< n$ enthalten ist.

- (5) Ein zweidimensionaler *Torus* ist gegeben durch

$$K := \{(x, y, z) \in \mathbb{R}^3 : (\sqrt{x^2 + y^2} - R)^2 + z^2 = \rho^2\}, \quad (8.21)$$

wo $0 < \rho < R$ fest gewählt sind.¹³ Wir wollen das vom Torus eingeschlossene Volumen berechnen und bemerken dazu, daß sich die obere Hälfte des Torus als

¹³Reine Mathematiker*innen definieren den Torus gerne als $\mathbb{S}^1 \times \mathbb{S}^1$, wobei $\mathbb{S}^1$ den Einheitskreis bezeichnet. (Reine Mathematiker*innen definieren den Einheitskreis nicht als $\{x \in \mathbb{R}^2 : |x| = 1\}$, sondern als den Quotientenraum $\mathbb{R}/\mathbb{Z}$; entsprechend kann man den Torus auch als $\mathbb{R}^2/\mathbb{Z}^2$ oder allgemeiner als $\mathbb{R}^n/\mathbb{Z}^n$ definieren.) Alle diese Definitionen sind topologisch äquivalent: Die durch (8.21) bestimmte Teilmenge des $\mathbb{R}^3$ ist *homöomorph* zu $\mathbb{S}^1 \times \mathbb{S}^1$, d.h. es existiert zwischen beiden Mengen eine stetige Bijektion, deren Inverse wieder stetig ist. Vgl. dazu den Begriff des C^1 -Diffeomorphismus.

Graph der Funktion

$$f(x, y) = \sqrt{\rho^2 - (\sqrt{x^2 + y^2} - R)^2}$$

darstellen läßt, sofern $R - \rho \leq \sqrt{x^2 + y^2} \leq R + \rho$. Diese Funktion ist rotationssymmetrisch, sodaß ihr Integral nach Satz 8.30 als

$$2\omega_2 \int_{R-\rho}^{R+\rho} r \sqrt{\rho^2 - (r-R)^2} dr$$

berechnet werden kann. Mit der Substitution $r' = r - R$ und mit $\omega_2 = \pi$ erhalten wir daraus

$$2\pi \int_{-\rho}^{\rho} (r' + R) \sqrt{\rho^2 - (r')^2} dr' = 2\pi \int_{-\rho}^{\rho} r' \sqrt{\rho^2 - (r')^2} dr' + 2\pi R \int_{-\rho}^{\rho} \sqrt{\rho^2 - (r')^2} dr'.$$

Das erste Integral auf der rechten Seite ist null, denn der Integrand ist ungerade und wird über ein symmetrisches Intervall um null integriert. Das zweite Integral ist genau die halbe Fläche einer Kreisscheibe mit Radius ρ und beträgt daher $\frac{1}{2}\pi\rho^2$. Multiplikation mit 2 (da der Graph von f nur den halben Torus beschreibt) liefert schließlich den Wert $2\pi^2 R \rho^2$ für das Volumen des Torus.

8.4. Der GAUSSSCHE Integralsatz

Ein Beweis des GAUSSschen Integralsatzes ist aus Zeitgründen leider nicht mehr möglich. Hier soll lediglich die Aussage dieses wichtigen Satzes diskutiert werden. Im einfachen Fall eines achsenparallelen Quaders gelingt der Beweis allerdings mühelos. Eine vollständige Behandlung des Integralsatzes und verwandter Themen (Integration auf Mannigfaltigkeiten etc.) wird nächstes Semester in Analysis III angeboten.

8.4.1. Der GAUSSSCHE Integralsatz auf einem Quader. Sei $Q = [0, 1]^n \subset \mathbb{R}^n$ der Einheitswürfel (natürlich sind auch allgemeine achsenparallele Quader der Form $\prod_{k=1}^n [a_k, b_k]$ zulässig). Sei $f: Q \rightarrow \mathbb{R}^n$ ein stetig differenzierbares Vektorfeld. Wir erinnern uns an den Divergenzoperator

$$\operatorname{div} f(x) := \sum_{k=1}^n \partial_k f_k(x),$$

sodaß $\operatorname{div} f$ eine stetige Funktion $Q \rightarrow \mathbb{R}$ ist. Für die folgende Rechnung beachten wir, daß eine Funktion auf einem Quader einfach sukzessive und in beliebiger Reihenfolge integriert werden kann (Abschnitt 8.1.2). Daher ist

$$\begin{aligned} \int_Q \operatorname{div} f(x) dx &= \sum_{k=1}^n \int_0^1 \int_0^1 \cdots \int_0^1 \partial_k f_k(x_1, \dots, x_n) dx_1 \cdots dx_n \\ &= \sum_{k=1}^n \int_0^1 \int_0^1 \cdots \left(\int_0^1 \partial_k f_k(x_1, \dots, x_n) dx_k \right) dx_1 \cdots dx_n \\ &= \sum_{k=1}^n \int_0^1 \cdots \int_0^1 [f_k(x_1, \dots, 1, x_{k+1}, \dots, x_n) - f_k(x_1, \dots, 0, x_{k+1}, \dots, x_n)] dx_1 \cdots dx_n. \end{aligned} \tag{8.22}$$

Im letzten Schritt haben wir partielle Integration verwendet. Man beachte, daß Punkte der Form $(x_1, \dots, 1, x_{k+1}, \dots, x_n)$ bzw. $(x_1, \dots, 0, x_{k+1}, \dots, x_n)$ sämtlich auf dem Rand von Q liegen. Das Integral von $\operatorname{div} f$ hängt also nur von den Werten von f auf ∂Q ab.

Genauer können wir folgendes feststellen: An einem Punkt $(x_1, \dots, 1, x_{k+1}, \dots, x_n)$, für den $x_j \in (0, 1)$ für alle $j \neq k$, ist der k -te Einheitsvektor e_k der eindeutig bestimmte Vektor mit Einheitslänge, der senkrecht auf ∂Q steht und bzgl. Q nach außen zeigt. Ebenso ist

$-e_k$ der Vektor, der im Punkt $(x_1, \dots, 0, x_{k+1}, \dots, x_n)$ senkrecht auf ∂Q steht und nach außen zeigt. Wenn wir also

$$\nu(x) := \begin{cases} e_k & \text{falls } x = (x_1, \dots, 1, x_{k+1}, \dots, x_n) \text{ und } x_j \in (0, 1) \text{ für } j \neq k, \\ -e_k & \text{falls } x = (x_1, \dots, 0, x_{k+1}, \dots, x_n) \text{ und } x_j \in (0, 1) \text{ für } j \neq k \end{cases}$$

setzen, so ist ν ein auf einer Teilmenge von ∂Q definiertes Vektorfeld, das überall in seinem Definitionsbereich senkrecht auf ∂Q steht und nach außen zeigt; man nennt ν deshalb das *äußere Normalenfeld* auf ∂Q .

Daß ν auf den ‚Kanten‘ und ‚Ecken‘ des Würfels, also auf den Mengen, auf denen mehr als eine Koordinate 0 oder 1 ist, nicht definiert ist, stört uns nicht, da die Randpunkte bei der Integration in (8.22) keine Rolle spielen¹⁴. Wir können nun (8.22) eleganter formulieren, indem wir feststellen, daß an einem Punkt der Form $x = (x_1, \dots, 1, x_{k+1}, \dots, x_n)$ gilt

$$f_k(x) = f(x) \cdot e_k = (f \cdot \nu)(x)$$

und analog an einem Punkt der Form $x = (x_1, \dots, 0, x_{k+1}, \dots, x_n)$

$$-f_k(x) = f(x) \cdot (-e_k) = (f \cdot \nu)(x).$$

In der Summe $k = 1, \dots, n$ durchläuft die Integration in (8.22) alle Seitenflächen des Würfels, also können wir schreiben

$$\int_Q \operatorname{div} f(x) dx = \int_{\partial Q} (f \cdot \nu)(x) dS(x), \quad (8.23)$$

wobei die Schreibweise $dS(x)$ andeuten soll, daß bezüglich des $n-1$ -dimensionalen Oberflächenmaßes integriert wird. Im vorliegenden Falle handelt es sich um eine Kurzschreibweise für die in (8.22) auftretenden Mehrfachintegrale der Form $dx_1 \cdots dx_n$, bei denen dx_k ausgelassen wird.

Die Formel (8.23) ist als GAUSSscher Integralsatz bekannt. Man beachte den Spezialfall $n = 1$: In diesem Falle ist $f : [0, 1] \rightarrow \mathbb{R}$ eine stetig differenzierbare Funktion, $\operatorname{div} f = f'$, und der Rand des Intervalls $[0, 1]$ besteht nur aus den beiden Punkten 0 und 1, an denen $\nu(0) = -1$ und $\nu(1) = 1$. Daher reduziert sich der Satz von GAUSS zu

$$\int_0^1 f'(x) dx = f(1) - f(0),$$

also zum Hauptsatz der Differential- und Integralrechnung. Umgekehrt kann der Satz von GAUSS daher als höherdimensionale Verallgemeinerung des Hauptsatzes angesehen werden.

8.4.2. Der Integralsatz auf allgemeineren Gebieten.

8.4.2.1. *Untermannigfaltigkeiten.* Wir beschränken uns hier der Einfachheit halber auf zweidimensionale Untermannigfaltigkeiten des $\mathbb{R}^3$.

DEFINITION 8.38. Eine Teilmenge $M \subset \mathbb{R}^3$ heißt zweidimensionale *Untermannigfaltigkeit*, wenn zu jedem $x \in M$ eine Umgebung $U \subset \mathbb{R}^3$ und eine stetig differenzierbare Funktion $f : U \rightarrow \mathbb{R}$ existiert, sodaß

- (1) $M \cap U = \{y \in U : f(y) = 0\}$,
- (2) $\nabla f(x) \neq 0$.

Eine Untermannigfaltigkeit ist also lokal als Nullstellenmenge einer Funktion definiert. Nach dem Satz über implizite Funktionen kann die Gleichung $f(y) = 0$ nach einer der Koordinaten y_k ($k = 1, 2, 3$) aufgelöst werden, o.B.d.A. nach y_3 : Dann ist die Untermannigfaltigkeit lokal dargestellt als Graph einer stetig differenzierbaren Funktion $y_3 = y_3(y_1, y_2)$.

Anschaulich ist eine solche Untermannigfaltigkeit eine *zweidimensionale gekrümmte Fläche* in $\mathbb{R}^3$. Als Beispiel betrachte die Einheitskugel $\mathbb{S}^2 = \{x \in \mathbb{R}^3 : |x| = 1\}$; sie ist eine

¹⁴Genauer gesagt bilden diese Kanten und Ecken $n-1$ -dimensionale Nullmengen.

Untermannigfaltigkeit, denn sie ist die Nullstellenmenge der Funktion $f(x) = |x|^2 - 1$, und es ist $\nabla f(x) = 2x \neq 0$ für jedes $x \in \mathbb{S}^2$.

Sei M eine Untermannigfaltigkeit von $\mathbb{R}^3$ und $\gamma : (-1, 1) \rightarrow M$ eine stetig differenzierbare Kurve mit $\gamma(0) = x^0 \in M$. Dann heißt der Vektor $\gamma'(0) \in \mathbb{R}^3$ *Tangentenvektor* an M im Punkt x^0 , und der von allen Tangentenvektoren aufgespannte Untervektorraum $T_{x^0}M$ heißt *Tangentenraum* von M in x^0 . Der Tangentenraum an eine zweidimensionale Untermannigfaltigkeit in einem Punkt ist stets zweidimensional: In der Tat, ist M in einer Umgebung von $x^0 \in M$ die Nullstellenmenge einer Funktion f , so steht $\nabla f(x^0) \neq 0$ senkrecht auf jeder Kurve in M (siehe die Diskussion nach Satz 7.14), also ist die Dimension von $T_{x^0}M$ höchstens zwei. Sie ist aber auch mindestens gleich zwei, denn nach dem Satz über implizite Funktionen ist M lokal um x^0 der Graph einer Funktion $x_3 = g(x_1, x_2)$, und dann sind $t \mapsto (x_1 + t, x_2, g(x_1 + t, x_2))$ und $t \mapsto (x_1, x_2 + t, g(x_1, x_2 + t))$ zwei Kurven in M , deren Ableitungen bei $t = 0$ die linear unabhängigen Vektoren $(1, 0, \partial_1 g(x_1, x_2))$ bzw. $(0, 1, \partial_2 g(x_1, x_2))$ sind.

Zu einem zweidimensionalen Untervektorraum von $\mathbb{R}^3$ gibt es stets genau zwei Normalenvektoren mit Einheitslänge, ν und $-\nu$. Eine Mannigfaltigkeit M heißt *orientierbar*, wenn es eine stetige Abbildung $\nu : M \rightarrow \mathbb{S}^2$ gibt, sodaß für jedes $x \in M$ der Vektor $\nu(x)$ senkrecht auf dem Tangentenraum $T_x M$ steht.¹⁵

Sei $\Omega \subset \mathbb{R}^3$ beschränkt und offen, dann ist der Abschluß $\bar{\Omega} = \Omega \cup \partial\Omega$ kompakt. Angenommen, $\partial\Omega$ ist eine zweidimensionale Untermannigfaltigkeit (was für ‚vernünftige‘ Gebiete stets der Fall sein wird), so kann man das Integral einer Funktion über Ω als dasjenige über $\bar{\Omega}$ definieren, und man kann zeigen (vgl. [10], §15), daß $\partial\Omega$ orientierbar ist, und das Normalenvektorfeld ν kann auf $\partial\Omega$ in eindeutiger Weise so gewählt werden, daß es nach außen zeigt, das heißt es gibt ein $t^0 > 0$ mit $x + t\nu(x) \notin \Omega$ für alle $t \in (0, t^0)$.

In der Formel (8.23) sind damit die Terme $\int_{\Omega} \operatorname{div} f(x) dx$ und $(f \cdot \nu)(x)$ (für $x \in \partial\Omega$) wohldefiniert. Was ist schließlich mit der Integration über $\partial\Omega$ bezüglich des Oberflächenmaßes $dS(x)$? Das ist eine recht aufwendig zu beantwortende Frage, der in Analysis III nachgegangen werden wird. Wir begnügen uns hier mit einer Erläuterung in einer Dimension.

Analog dem Falle zweidimensionaler Untermannigfaltigkeiten in $\mathbb{R}^3$ ist eine eindimensionale Untermannigfaltigkeit in $\mathbb{R}^2$ lokal der Graph einer stetig differenzierbaren Funktion. Sei $[a, b] \in \mathbb{R}$ und $f : [a, b] \rightarrow \mathbb{R}$ stetig differenzierbar, und wir betrachten den Graphen $\Gamma = \{(x, f(x)) : x \in [a, b]\}$. Wir können uns also Γ als Teil einer Mannigfaltigkeit vorstellen, auf der wir integrieren wollen. Was ist dafür ein sinnvoller Integralbegriff?

Um $\int_{\Gamma} g(t) dS(t)$ für eine stetige Funktion $g : \Gamma \rightarrow \mathbb{R}$ zu definieren, betrachten wir den einfachsten Fall: $g \equiv 1$. Dann sollte $\int_{\Gamma} 1 dS(t)$ die Länge der Kurve Γ angeben. Parametrisieren wir Γ durch $s \mapsto (s, f(s))$, so ist nach Beispiel 8.7 die Länge von Γ durch das Integral von $|(s, f(s))'| = |(1, f'(s))|$ gegeben, also

$$\int_{\Gamma} 1 dS(t) = \int_a^b |(1, f'(s))| ds = \int_a^b \sqrt{1 + f'(s)^2} ds.$$

Es tragen also ‚steile‘ Abschnitte der Kurve mehr zu deren Länge bei als ‚flache‘, was unmittelbar der geometrischen Anschauung entspricht.

Dies motiviert folgende Definition des Kurvenintegrals über eine beliebige Funktion $g : \Gamma \rightarrow \mathbb{R}$:

$$\int_{\Gamma} g(t) dS(t) := \int_a^b g(s, f(s)) \sqrt{1 + f'(s)^2} ds.$$

In ähnlicher Weise kann man das Integral über eine zweidimensionale Fläche definieren, sofern diese als Graph einer Funktion dargestellt wird. Da eine Untermannigfaltigkeit M

¹⁵Das bekannteste Beispiel einer nichtorientierbaren Mannigfaltigkeit ist das *Möbiusband*; es wird empfohlen, sich den Wikipediaeintrag anzusehen.

lokal als Graph dargestellt werden kann, definiert man das Oberflächenintegral über M , indem man die einzelnen Graphen miteinander ‚verklebt‘; dies geschieht mittels einer Teilung der Eins, wie wir sie in Abschnitt 8.2.2 kennengelernt haben.

Man kann dann, analog zu (8.23), zeigen:

SATZ 8.39 (GAUSSscher Integralsatz). Sei $\Omega \subset \mathbb{R}^n$ beschränkt und offen und so, daß $\partial\Omega$ eine $(n-1)$ -dimensionale Untermannigfaltigkeit ist. Sei $f : \overline{\Omega} \rightarrow \mathbb{R}^n$ ein stetig differenzierbares Vektorfeld und $\nu : \partial\Omega \rightarrow \mathbb{S}^{n-1}$ das äußere Normalenvektorfeld an $\partial\Omega$. Dann gilt

$$\int_{\Omega} \operatorname{div} f(x) dx = \int_{\partial\Omega} (f \cdot \nu)(x) dS(x).$$

8.4.3. Anwendung: inkompressible Strömungen. Wir betrachten den Fluß einer Flüssigkeit durch den Raum $\mathbb{R}^3$. Die Flüssigkeit soll inkompressibel (als nicht zusammenpreßbar) und homogen sein (d.h. die Massendichte ist konstant). Wasser bei konstanter Temperatur ist ein gutes Beispiel.

Fixiere einen Zeitpunkt, zu dem die Strömung untersucht werden soll. An jedem Punkt $x \in \mathbb{R}^3$ des Raums beschreibe der Vektor $v(x) \in \mathbb{R}^3$ die Strömungsgeschwindigkeit des Fluids. Dann ist $v : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ das *Strömungsfeld*. Wir wollen mithilfe des GAUSSschen Integralsatzes untersuchen, was aus dem Prinzip der Massenerhaltung für das Strömungsfeld folgt.

Die Massenerhaltung besagt, daß nirgends in $\mathbb{R}^3$ Flüssigkeit zu- oder abfließt. Sei $\Omega \subset \mathbb{R}^3$ eine beschränkte offene Teilmenge mit glattem Rand, dann besagt das Prinzip der Massenerhaltung also, daß genausoviel Flüssigkeit in das Gebiet Ω ein- wie ausfließt. An jedem Randpunkt $x \in \partial\Omega$ ist der infinitesimale Ab- bzw. Zufluß nach Ω gegeben durch $(v \cdot \nu)(x)$, also fordern wir insgesamt

$$\int_{\partial\Omega} (v \cdot \nu)(x) dS(x) = 0.$$

Nach GAUSS ist dies gleichbedeutend mit

$$\int_{\Omega} \operatorname{div} f(x) dx = 0,$$

und da $\Omega \subset \mathbb{R}^3$ beliebig gewählt war, folgt $\operatorname{div} f \equiv 0$.

Das Strömungsfeld einer inkompressiblen Flüssigkeit ist also stets divergenzfrei.

Literaturverzeichnis

- [1] M. AIGNER und G. ZIEGLER. Das BUCH der Beweise. *Springer-Verlag, Berlin/Heidelberg*, 5. Aufl. 2018.
- [2] H. AMANN und J. ESCHER. Analysis I. Grundstudium Mathematik. *Birkhäuser Verlag, Basel*, 3. Aufl. 2006.
- [3] H. AMANN und J. ESCHER. Analysis II. Grundstudium Mathematik. *Birkhäuser Verlag, Basel*, 2. Aufl. 2008.
- [4] G. CANTOR. Beiträge zur Begründung einer transfiniten Mengenlehre. *Math. Ann.* **46** (1895), Nr. 4, 481–512.
- [5] R. COURANT und H. ROBBINS. What is mathematics? An elementary approach to ideas and methods. Bearb. v. I. STEWART. *Oxford Univ. Press, New York*, 2. Aufl. 1996.
- [6] O. DEISER. Einführung in die Mengenlehre. Die Mengenlehre Georg Cantors und ihre Axiomatisierung durch Ernst Zermelo. Springer-Lehrbuch. *Springer-Verlag, Berlin/Heidelberg*, 3. Aufl. 2010.
- [7] H.-D. EBBINGHAUS, J. FLUM und W. THOMAS. Einführung in die mathematische Logik. *Springer Spektrum, Berlin/Heidelberg*, 6. Aufl. 2018.
- [8] O. FORSTER. Analysis 1. Differential- und Integralrechnung einer Veränderlichen. Grundkurs Mathematik. *Springer Spektrum, Berlin/Heidelberg*, 12. Aufl. 2015.
- [9] O. FORSTER. Analysis 2. Differentialrechnung im $\mathbb{R}^n$, gewöhnliche Differentialgleichungen. Grundkurs Mathematik. *Springer Spektrum, Berlin/Heidelberg*, 11. Aufl. 2017.
- [10] O. FORSTER. Analysis 3. Integralrechnung im $\mathbb{R}^n$ mit Anwendungen. Aufbaukurs Mathematik. *Friedr. Vieweg & Sohn Verlagsgesellschaft mbH, Braunschweig/Wiesbaden*, 3. Aufl. 1999.
- [11] T. GOWERS. Mathematics: a very short introduction (= Very short introductions **66**). *Oxford Univ. Press, Oxford*, 2002.
- [12] K. JÄNICH. Topologie. *Springer-Verlag, Berlin/Heidelberg*, 8. Aufl. 2008.
- [13] T. WILHOLT. Logik und Argumentation. Materialien zu einführenden Vorlesungen über formale Logik und Argumentationstheorie. http://www.philos.uni-hannover.de/fileadmin/institut_fuer_philosophie/Personen/Wilholt/Logik.pdf, letzter Aufruf am 7.10.2018.