

ANALYSIS
in 3 Zügen

Sommersemester 1996

H. Dinges

Inhaltsverzeichnis

Vorwort zum Manuskript	1
Thesen zur Grundvorlesung Analysis	4
A Analysis I	8
A.1 Zur Analysis um 1800	8
1.1 Die Erfindung der elementaren Funktionen	8
1.2 Heutiges Schulwissen	9
1.3 Die Erfindung der Differential- und Integralrechnung	13
1.4 Die Lage um 1800	19
1.5 Zum Anliegen der Vorlesung	20
A.2 Die komplexen Zahlen; Möbiustransformationen	21
A.3 Polynome, gebrochenrationale Funktionen, Partialbruchzerlegung	32
A.4 Ansätze zur Integrationstheorie, Stammfunktionen	39
A.5 Variablensubstitution. Spezielle Kurven. Zu Newtons Ansatz.	51
A.6 Elementare Lösungen von Differentialgleichungen	65
A.7 Die Gammafunktion	78
A.8 Die Fehlerfunktion	89
A.9 Frühe Ansätze zur Theorie der Fourier-Reihen	102
B Analysis I	1
B.1 Das Programm der Arithmetisierung der Analysis	1
B.2 Wurzelziehen	10
B.3 Konstruktion der Exponentialfunktion	18
B.4 Konstruktion der trigonometrischen Funktionen	27
B.5 Konvergente Potenzreihen	33
B.6 Formale Potenzreihen	39
B.7 Totalstetige Funktionen; die Taylorsche Formel	48
B.8 Totale Differenzierbarkeit	59
B.9 „Dirichlets Funktionsbegriff“	71
C Analysis I	1
C.1 Zahlen aus der Sicht der Mengenlehre	1
1.1 Aufbau des Zahlensystems	1
1.2 Cantors Mengendefinition	3
1.3 Spezielle Mengen von Zahlen	4
1.4 Cartesische Produkte	5
1.5 Abbildungen und ihre Graphen	5
1.6 Abzählbare und überabzählbare Mengen	7
C.2 Aussagen	18
C.3 Metrische Räume	27
C.4 Der Begriff der Stetigkeit	40
C.5 Die Eigenschaft von Bolzano–Weierstraß: Sätze über stetige Funktionen	48
C.6 Der Satz von Picard–Lindelöf	59
C.7 Konvexe Funktionen	65
C.8 Normierte Vektorräume; die p -Normen	78

C.9 Fourier-Reihen aus funktionalanalytischer Sicht	85
---	----

Dieses Skript finden Sie im Internet unter der Adresse
<http://ismi.math.uni-frankfurt.de/dinges/lecturenotes/ANALYSISin3Zuegen.pdf>

Vorwort zum Manuskript Analysis I (in drei Zügen A, B, C)

Anlaß

Analysis I ist weithin eine Pflichtvorlesung für Studienanfänger in mehreren Studiengängen: Mathematik–Diplom, Physik–Diplom, Informatik–Diplom, Lehramt an Gymnasien. Ein Quäntchen Analysis wird an vielen Orten auch von den Studierenden der „niederen“ Lehrämter verlangt. Es ist nun offenbar nicht dasselbe, ob man nur das Abiturwissen befestigen will, ob man einen Einblick in die Denkweisen der Mathematik geben und damit Neugierde auf die Mathematik wecken will, oder ob es gilt, die technischen Grundlagen für einen nichtmathematischen Studiengang bereitzustellen. Die bekannten Kompromisse lassen m.E. bedauerlich viele Wünsche nach Differenzierung und nach Transparenz offen. Mir scheint, daß neue Arrangements gefordert sind, wenn die Analysis I nicht in Scholastik erstarren soll.

Es ergab sich im Sommersemester 1996 in Frankfurt, daß gleichzeitig für eine kleine Schar von Studienbeginnern (Mathematik, Informatik) und für ein größeres Kontingent von Studenten der „niederen“ Lehrämter im 4. Fachsemester eine Analysis anzubieten war. Die einen sollten ans Studium herangeführt werden, die anderen waren mit dem Endpunkt ihrer „wissenschaftlichen Grundlagen der Schulmathematik“ zu versorgen. Die einen hatten als Sommeranfänger 6 Stunden (+ 2 Stunden Übungen) Zeit, die anderen nur 4 Stunden (+ 2 Stunden Übungen). In dieser Lage kam ich auf die Idee, die Vorlesung in drei Zügen anzulegen. Das Abstraktionsniveau sollte systematisch differenziert werden, während sich die Gegenstände überschneiden und (aus der Sicht des Mathematikers) ergänzen sollten.

Aufbau

Der klassische Kalkül („Calculus“) wurde am Montag behandelt; dort sollten unter Hintansetzung der Skrupel der Mathematiker des 19. Jahrhunderts klassische Bestände vorgestellt werden. Am Mittwoch wurden im B–Zug die Bemühungen des 19. Jahrhunderts um den „allgemeinen“ Funktionsbegriff und die „Arithmetisierung der Analysis“ vorgestellt. Im C–Zug wurden die modernen Ausdrucksweisen der naiven „Mengenlehre und Aussagenlogik“ eingeführt und die Linien der weiteren mathematischen Abstraktion vorbereitet.

Lehrerfolg

Eine formelle Evaluation der Lehre fand nicht statt. Ich kann auf einige Studenten verweisen, die es sehr anregend fanden, daß viele Türen aufgestoßen wurden; es hat sich gezeigt, daß sie durchaus Mut für tieferes Studieren geschöpft haben. Viel weniger weiß ich über diejenigen, die enttäuscht waren; möglicherweise haben sie sich deshalb unbefriedigt von der Mathematik abgewandt, weil allzuvielen offenen Enden produziert oder angedeutet worden sind. Durch sog. Merkblätter versuchte ich denen entgegenzukommen, die bei jedem Thema gleich nervös fragen, wie denn die entsprechenden Prüfungsfragen aussehen könnten. Ich weiß wenig über den Effekt; die Prüfungen in den Diplomstudiengängen finden nach guter alter akademischer Tradition erst dann statt, wenn sich die Eindrücke aus den Anfängervorlesungen mit weiteren Studienerfahrungen amalgamiert haben. Die Lehramtsstudenten auf der anderen Seite

müssen nach anderen Traditionen ihren Anspruch auf einen Übungsschein durch eine Klausur in der letzten Vorlesungswoche untermauern. Die Leistungen waren im SS 96 so ernüchternd wie eh und je; die Übungsscheine konnten wie üblich in der Regel eigentlich nur die Mitarbeit in den Übungsstunden bestätigen. Ein didaktischer Erfolg ist also nicht zu melden.

Grund der Publikation

Trotz all der offenen Fragen haben wir uns (besonders zu erwähnen ist die Geduld und die L^AT_EX-Kompetenz von Frau M. Schmidt) die Mühe gemacht, die Veranstaltung zu dokumentieren. Auf jeden Fall hoffe ich, daß die Ausführungen für weitere Gelgenheiten als Steinbruch nützlich werden. Die Vielfalt der didaktischen Randüberlegungen sollte anregend sein. In besonderem Maße hoffe ich, daß die bewußte Differenzierung der Abstraktions- bzw. Präzisionsstufen nützliche Bezugspunkte liefern wird, wenn es gilt, mit den Verantwortlichen für die verschiedenen Studiengänge neu auszuhandeln, welche Art von Analysis nun wirklich gefordert werden soll.

Offene Fragen

Was steckt dahinter, wenn die Lehrerausbilder manchmal sagen, die Analysis müsse für den Gebrauch in der Schule möglichst anschaulich und mit einem Minimum an Formelkram behandelt werden, wobei aber bei der typisch mathematischen Präzision keine Abstriche gemacht werden dürften? Ist es wahr, daß die Physikstudenten von der Analysis vor allem erhoffen, daß sie Sicherheit gewinnen im Umgang mit konkreten Funktionen und konkreten Differentialgleichungen? Ist für die Informatiker die Analysis der Platz, wo sie sich an Mengen, Aussagen, Abbildungen, Quantoren usw. gewöhnen sollten und vielleicht noch ein wenig hören sollten über Polynome und diejenigen Funktionen, die bei der Beschreibung von Komplexität auftreten? Ist es wahr, daß es für das eigentliche Mathematikstudium schon im ersten Semester geboten ist, daß die Theorie lückenlos aufgebaut wird und die Studierenden das Beweisen lernen, wobei die Gegenstände von sekundärer Bedeutung sind?

Stoff-Fülle

Die Analysis ist so reich, daß es nicht schwer fällt, in der Vorlesung ein Feuerwerk von Themen abzubrennen, welches die mathematischen Talente begeistert. Das Feuerwerk braucht aber Zeit, wenn es sich nicht in der Zurschaustellung mathematischer Virtuosität erschöpfen soll. Diese Zeit können sich nicht alle Studierenden gönnen. Im Hinblick auf die zu verkürzenden Studienzeiten wird immer lauter eine „Entrümpelung“ gefordert.

Die Schwierigkeit für die Lehrenden liegt bei der Beschränkung des Lehrstoffs und bei der Auswahl dessen, was eingeübt werden soll. Ausdünnung des Lehrstoffs kann nicht gemeint sein; denn der Bedarf an mathematischer Kompetenz ist keinesfalls geringer geworden. Es sollte eine Orientierung liefern, wenn die Gefahr entsteht, bei der Lektüre ausführlicher Lehrbücher den Faden zu verlieren. Es muß bedacht werden, welche Lehrgegenstände und Methoden langfristig (auf das gesamte Studium hin gesehen und vielleicht sogar darüber hinaus) den besten Ertrag erbringen. Eine sinnvolle Evaluation ist gewiß nicht einfach.

Umfang des Skriptums

Es fällt natürlich schwer, traditionelle Lehrstoffe abzuwählen, die in so vielen schönen Büchern methodisch gut aufgearbeitet sind und zum festen Wissensbestand der älteren Semester gehören. Es ist viel leichter, ausführlich zu sein. Der Leser meines Vorlesungsmanuskripts sollte nicht glauben, daß alles was da steht, in der Vorlesung so gesagt wurde. Es ist klar, daß ein Großteil der hier abgehandelten Gegenstände für viele Studenten zunächst einmal oder vielleicht sogar endgültig, entbehrlich ist. Ein Überfliegen der Passagen, die ihn nicht (oder noch nicht) ansprechen, kann aber nicht schaden, meine ich. Die Gliederung sollte das bewußte Auswählen dessen, was er gründlich lernen will, erleichtern. Der Computer erledigt allerlei, mathematisches Verständnis kann er jedoch schwerlich ersetzen. Mathematische Kompetenz braucht kohärentes Wissen auf einem abstrakten Niveau. Die Lehrveranstaltungen dürfen daher nicht als Vorbereitungskurse für die nächste Klausur angelegt werden. In meinem Manuskript wird auch deshalb so viel über Altmodisches geredet, weil ich benennen wollte, was ich den Studenten, die keine entsprechenden Spezialkenntnisse brauchen, mit Bedacht ersparen möchte.

Fazit

Die Mathematiker haben im Laufe der Jahrzehnte recht verschiedene Erwartungen in die Analysis I gesetzt. Es ist problematisch, das alles zu einem Einheitsbrei zu verkochen; und dann je nach Studienrichtung dem einen mehr, dem anderen weniger davon zu verabreichen.

Thesen zur Grundvorlesung Analysis für Mathematiker

Es ist das höchst anspruchsvolle Anliegen der Anfängervorlesungen, die Mathematik als Wissenschaft vorzustellen. Die Anfänger müssen in die Lage versetzt werden, zu entscheiden, ob sie sich auf die Anstrengungen des Mathematikstudiums einlassen wollen und können.

Eine Anfängerveranstaltung an einer Universität muß die Ziele des wissenschaftlichen Denkens zum Vorschein bringen; sie darf sich nicht darauf beschränken, das technische Rüstzeug für die Bewältigung von Spezialvorlesungen bereitzustellen. Anders als ein Grundkurs an einer Fachhochschule, welcher der Vorbereitung auf den Gebrauch der Mathematik als Hilfswissenschaft dient und unverzüglich punktuell durch eine Leistungsprüfung abgeschlossen werden kann, muß die Einführung über die elementaren Techniken hinausweisen. Das Ziel des Grundstudiums sollte es sein, daß sich die Kenntnisse bis zur Vordiplomsprüfung abrunden und dort in der angemessenen Breite zur Sprache gebracht werden können. Da die Beschreibung der Probleme, an welchen die aktuelle mathematische Forschung arbeitet, im allg. höchst voraussetzungsreich ist, sind Kontexte der Anfängervorlesungen wahrscheinlich besser als viele spezielle Kontexte aktueller Forschung geeignet, die für die Mathematik charakteristischen Modellierungs-, Abstraktions- und Verallgemeinerungsprozesse zum Vorschein zu bringen. Die Gegenstände der Anfängerausbildung dürfen allerdings nicht als Preziosen aus alten Vorratskammern präsentiert werden; sie müssen als Konstruktionen zur Bewältigung von anregenden Problemen verstanden werden.

Im Folgenden stelle ich einige Themen heraus, die ich anders als vielerorts üblich behandeln möchte. Diejenigen, die meinen, ich wüßte nicht, daß der Teufel häufig im Detail steckt, verweise ich auf meine Skripten zur Analysis.

1. (Die historische Dimension)

Die zentralen Themen der Analysis stehen seit fast 300 Jahren fest: Konvergenz und Stetigkeit sowie Differentiation und Integration. Die Anfängervorlesung hat es aber nicht mit totem Lernstoff zu tun; es muß um die Mathematisierungsprozesse gehen. Den Studierenden sollte nicht zuletzt auch eine Ahnung vermittelt werden, wie sich das Interesse an den Themen der Analysis im Laufe der Zeit gewandelt hat. Im 18. Jahrhundert (bei Leibniz, Bernoulli, Euler usw.) stand die Konstruktion spezieller Funktionen (zur Lösung von Problemen aus Mechanik und Astronomie) im Vordergrund; das 19. Jahrhundert versuchte den „allgemeinen“ Funktionsbegriff zu fassen und zu begründen (Cauchy, Dirichlet, Weierstraß usw.); und im 20. Jahrhundert erscheint der Begriff der reellen Funktionen einer reellen Veränderlichen als Spezialfall des Begriffs einer Abbildung einer abstrakten Menge in eine abstrakte Menge (Cantor u.a.). In der Analysis des 20. Jahrhunderts geht es dann in erster Linie um Räume von Funktionen (Lebesgue, Riesz, Banach usw.)

Das 19. Jahrhundert hatte mit fehlgeleiteten intuitiven Vorstellungen von den zur Betrachtung zuzulassenden Funktionen zu kämpfen; und man erfand zu diesem Zweck die „mathematische Strenge“. Über dem Training dieser (geschichtslos und problemunabhängig vorgestellten) mathematischen Strenge vernachlässigten die heutigen Anfängerveranstaltungen häufig die Konstruktion; (sowohl die Konstruktion und Anwendung von speziellen Funktio-

nen im Sinne des 18. Jahrhunderts, als auch die Konstruktion von Funktionenräumen im Sinne des 20. Jahrhunderts). Es sollten endlich wieder die mathematischen Ideen gegenüber den mathematischen Beweisen in den Vordergrund gerückt werden.

2. (Funktionen und Abbildungen)

Mengen und Abbildungen von Mengen verbinden sich dem Anfänger nicht ohne weiteres mit dem in der Schule behandelten Funktionsbegriff. Um die Vorstellungen wirksam zusammenzubringen, werden bei uns die Polynome und die rational gebrochenen Funktionen von vorneherein auch schon im Komplexen diskutiert. Z.B.: Es bereitet auf den Begriff der reellwertigen Funktion von mehreren Variablen vor, wenn man (bei einem Beweisansatz zum Fundamentalsatz der Algebra) den Betrag eines Polynoms im Komplexen ins Auge faßt. Bei einem weiteren Beweisansatz bietet es sich an, schon einmal von Kurven (als Abbildungen eines Intervalls in den Raum $\mathbb{C}$ oder in einen Raum $\mathbb{R}^d$) zu sprechen. Die linear gebrochenen Funktionen scheinen besonders gut geeignet, den Abbildungs- und Funktionsbegriff nebeneinander zu stellen: als Möbiustransformationen werden sie hintereinandergeschaltet, bei der Partialbruchzerlegung werden sie als Funktionen addiert.

3. (Konvergenz und Stetigkeit)

Quantoren sollen bei uns nicht nur als Kürzel innerhalb der Umgangssprache verwendet werden. Wir wollen nicht auf die Kraft des formellen Umgangs mit den Quantoren verzichten, wenn wir Begriffe diskutieren wie

- (i) Stetigkeit in einem Punkt
- (ii) gleichmäßige Stetigkeit einer Funktion (über einem Bereich)
- (iii) gleichgradige Stetigkeit einer Funktionenschar (in einem Punkt).

Wenn die Konvergenz von Zahlenfolgen geklärt ist (im metrischen Sinn wie auch im anordnungstheoretischen als Gleichheit von oberem und unterem Limes), dann wollen wir keine Zeit verlieren mit der punktweisen Konvergenz von Funktionen. Wir wenden uns sogleich den viel wichtigeren Konvergenzbegriffen zu:

- (i) gleichmäßige Konvergenz auf Kompakten (etwa bei den Potenzreihen)
- (ii) Konvergenz bzgl. einer Norm (etwa bei den Fourierreihen)
- (iii) fastsichere Konvergenz (als fastsichere Gleichheit von oberem und unterem Limes bei der Integrationstheorie in Analysis II).

4. (Fundamentalsatz der Differential- und Integralrechnung)

Wir wollen uns nicht um eine möglichst allgemeine Version des sog. Fundamentalsatzes bemühen, weil alle Betrachtungen zum Definitionsbereich der Operationen „Differentiation“ und „Stammfunktionenbildung“ vom Kalkül (im Sinne des 18. Jahrhunderts) ablenken.

Wir beweisen nur den **Satz** : „Sei $(h_n)_n$ eine Nullfolge. Wenn $f(\cdot)$ auf $\mathbb{R}$ **gleichmäßig stetig** ist und $F(\cdot)$ eine Stammfunktion zu $f(\cdot)$ ist, dann strebt die Funktionenfolge $f_n(x) = \frac{1}{h_n} \cdot (F(x + h_n) - F(x))$ **gleichmäßig** gegen $f(\cdot)$.“

(Wenn von $f(\cdot)$ nur die Stetigkeit auf (a, b) vorausgesetzt ist, dann haben wir gleichmäßige Konvergenz auf jedem Kompaktum $\subseteq (a, b)$.)

5. (Gegen das Riemann–Integral)

Das Riemann–Integral hat keinen Platz in unserer Einführung. Es scheint uns einigermaßen abstrus, die Integrationstheorie zu einem Übungsfeld für Konvergenzüberlegungen zu degradieren (vgl. Forster I, § 18). Die adhoc–Methode, mit welcher Riemann die Dirichlet’schen Bedingungen für die Entwickelbarkeit einer Funktion in eine Fourier–Reihe untersucht hat, führt nicht zu einer Beschreibung eines natürlichen Definitionsbereichs des Funktional „Integral“, wie Lebesgue immer wieder betont hat.

Das Riemann–Integral ist auch deswegen didaktisch verfehlt, weil es quer liegt zur Idee der Vervollständigung, die wir für metrische Räume, und nicht nur für die rationalen Zahlen durchführen.

6. (Integrabilität)

Das Auffinden elementarer Stammfunktionen soll (in mäßigem Umfang) geübt werden, ebenso das Lösen elementar lösbarer Differentialgleichungen (getrennte Variable, exakte Differentiale). Integrabilität soll aber nicht problematisiert werden. Wie in der Zeit vor Cauchy soll es uns genügen zu sagen: Für eine positive Funktion $f(\cdot)$ ist das Integral die Fläche unter der Kurve (endlich oder $+\infty$). Wenn für ein $f(\cdot)$ sowohl $f^+(\cdot)$ als auch $f^-(\cdot)$ unendliche Fläche unter der Kurve haben (beispielsweise $\frac{1}{x} \cdot \sin x$), dann gibt es kein Integral; ansonsten

$$\int_a^b f(x) dx = \int_a^b f^+(x) dx - \int_a^b f^-(x) dx .$$

Die Additivität des Integrals wird als Selbstverständlichkeit behandelt.

Anmerkung : Ob es positive Funktionen geben könnte, bei denen fraglich sein könnte, was unter der Fläche unter der Kurve zu verstehen ist, wird als Thema für die Analysis II angekündigt. — Niemand hat jemals eine solche Funktion gesehen; in der konkreten Umgebung der Analysis I ist also kein Anlaß zum Zweifel gegeben. Es ist allerdings vielleicht schon einmal ganz interessant, sich einigermaßen bizarre positive Funktionen mit dem Integral $= 0$ auszudenken („Nullfunktionen“). Funktionen, die sich nur um eine Nullfunktion unterscheiden, haben schließlich dieselbe Stammfunktion. Und das zeigt, daß der „Fundamentalsatz“ nicht dahingehend verallgemeinert werden kann, daß man sagt: Differentiation und Stammfunktionenbildung sind zueinander inverse Operation.

7. (Reihen)

Bedingt konvergente Reihen und uneigentliche Integrale gibt es bei uns nicht. Wenn man eine Reihe mit der Folge der Partialsummen identifiziert, verspielt man die Chance zu einer wirksamen Theorie der Reihen (dominierte Konvergenz, Fatou's Lemma etc.). Die sog. Summierungstheorie gehört in die Geschichte der Analysis des 19. Jahrhunderts oder in entsprechende Spezialvorlesungen. Ausdrücke, wie $\int_{-\infty}^{\infty} \frac{1}{x} \sin x dx$, sind von der Tradition geheiligte nützliche Abkürzungen, passen aber nicht in eine aufgeklärte Integrationstheorie.

8. (Differenzierbarkeit und der Satz von Taylor)

Jeder, der Analysis konkret betreibt, weiß, daß der sog. Dirichlet'sche Funktionsbegriff nicht das Thema ist. Es gibt nirgends einen Grund, über den Begriff der borelmeßbaren Funktion hinauszugehen; und erst in der Integralrechnung wird man bis dahin vorstoßen wollen. Welche Funktionen man im Auge haben sollte, wenn man elementare Differentialrechnung lehrt, ist eine didaktische Frage. Liouville wollte bekanntlich noch 1893 gewisse stetige Funktionen, die Weierstraß vorgestellt hat, als Monstrositäten aus der „gesunden“ Analysis“ verbannen. Für Vorurteile in die eine oder in die andere Richtung ist bis heute viel Platz geblieben.

Hier meine Thesen:

- (i) Totale Differenzierbarkeit einer Funktion (von mehreren Variablen) in einem Punkt ist zweifellos ein guter Begriff.
- (ii) Stetige Differenzierbarkeit ist ebenfalls ein nützlicher Begriff.
- (iii) Unfruchtbar scheint mir der Begriff der (in jedem Punkt eines Intervalls) differenzierbaren Funktion zu sein. Es gibt keine ansehnliche Theorie der Funktionen, die diese und keine weiteren Regularitätseigenschaften haben. (Über den Wert des Mittelwertsatzes der Differentialrechnung kann man streiten.) Mathematiker sind bekanntlich frei, jeden in sich nicht widersprüchlichen Begriff einzuführen; ob ein Professor das Recht hat, die Studierenden mit unfruchtbaren Begriffsbildungen zu plagen, ist eine andere Frage.
- (iv) Ein wunderschöner Begriff ist natürlich der Begriff der komplexen Differenzierbarkeit in jedem Punkt; aber nur deswegen, weil er die Holomorphie impliziert. Es ist wohl eine Geschmacksfrage, ob man davon schon in der Analysis I sprechen soll.
- (v) Die Funktionen, die als Stammfunktionen auftreten können, heißen bekanntlich totalstetige Funktionen. Der Begriff der Totalstetigkeit ist über jeden Zweifel erhaben; allerdings wird man seine Tugenden erst in der Analysis II ermessen können.
(Beispiele: $|x|$ ist totalstetig, $x^2 \cdot \sin \frac{1}{x^2}$ ist nicht totalstetig, wiewohl in jedem Punkt differenzierbar.)

Den Taylorschen Satz formulieren und beweisen wir als einen Satz über Funktionen, deren $(n - 1)$ -te Ableitung totalstetig ist (und im Zentralpunkt darüberhinaus stetig differenzierbar ist).

Kapitel A

ANALYSIS I

A.1 Zur Analysis um 1800

Die folgenden historischen Anmerkungen stützen sich vor allem (zum Teil wörtlich) auf D. STRUIK: Abriss der Geschichte der Mathematik, VEB Deutscher Verlag der Wissenschaften Berlin, 2. Auflage 1963. Ausführlichere Auskünfte gibt C.H. EDWARDS, JR.: The Historical Development of the Calculus, Springer Verlag 1979.

1.1 Die Erfindung der elementaren Funktionen

Im Zeitalter der großen astronomischen Theorien von Kopernikus (1473–1543), Tycho Brahe (1546–1601) und Kepler (1571–1630) brauchte man Trigonometrie. Trigonometrische und astronomische Tafeln mit ständig wachsender Genauigkeit erschienen besonders in Deutschland. Eine vollständige Einführung in die Trigonometrie (1464, aber erst 1533 gedruckt) gilt als das Hauptwerk des JOHANNES MÜLLER REGIOMONTANUS, der führenden mathematischen Persönlichkeit des fünfzehnten Jahrhunderts. Da die heutigen bequemen Bezeichnungen noch nicht existierten, mußten die Sätze, wie z.B. der sogenannte Sinussatz für sphärische Dreiecke in Worten ausgedrückt werden.

Auch die Zahldarstellungen waren kompliziert. Zwar hatten persische Mathematiker schon im ersten Teil des 15. Jahrhunderts mit Dezimalbrüchen mit Komma gerechnet. Im Abendland wurden die Dezimalzahlen aber erst durch SIMON STEVIN (1548–1620) systematisch eingeführt.

Durch die allgemeine Einführung der indisch-arabischen Zahlenschreibweise wurde auch der Weg bereitet für die Erfindung der Logarithmen. Mehrere Mathematiker des 16. Jahrhunderts hatten mit der Möglichkeit gespielt, arithmetische und geometrische Reihen zu verknüpfen, hauptsächlich zu dem Zweck, das Arbeiten mit den komplizierten trigonometrischen Tafeln zu erleichtern. Ein bedeutender Beitrag zu diesem Ziel wurde von einem schottischen Gutsbesitzer, JOHN NEPER (oder NAPIER) 1614 geleistet. Neper war mit seinem ersten Ansatz selbst nicht ganz zufrieden und einigte sich dann mit seinem Bewunderer HENRY BRIGGS, Professor am Gresham College in London auf die „Briggs’schen“ Logarithmen zur Basis 10. 1617 vollendeten zwei holländische Feldmesser ADRIAEN VLACQ und EZECHIEL DE DECKER

das von Briggs begonnene Tafelwerk. Die neue Erfindung und die Tafeln wurden sofort von Mathematikern und Astronomen und besonders auch von Kepler begrüßt.

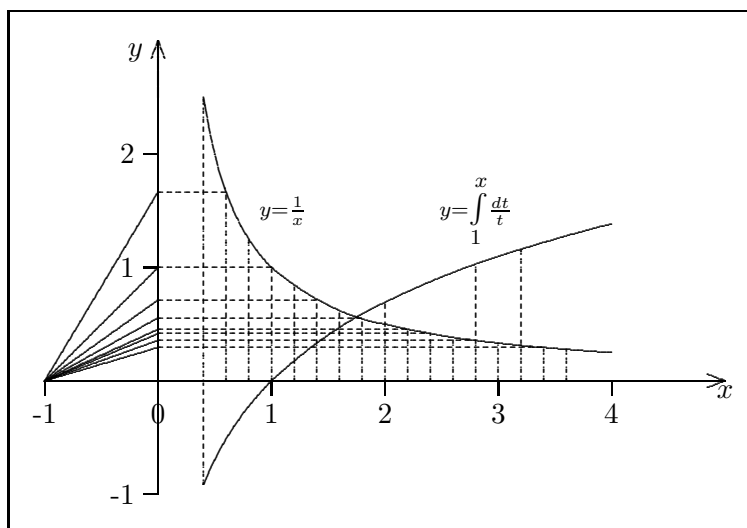
Natürliche Logarithmen auf der Grundlage der Exponentialfunktion erschienen fast gleichzeitig mit den Briggs'schen Logarithmen, aber ihre fundamentale Bedeutung wurde nicht eher erkannt, als bis die Infinitesimalrechnung besser verstanden worden war. Die heutigen bequemen Bezeichnungen gehen auf EULER (1748) zurück.

1.2 Heutiges Schulwissen

Die folgende Aussage dürfte von der Schule her bekannt sein:

$$\int_x^y \frac{1}{u} du = \ln y - \ln x = \ln \left(\frac{y}{x} \right) .$$

Die Fläche unter der Kurve $\frac{1}{u}$ über dem Intervall ist durch die Logarithmusfunktion gegeben. Diese ist auf jedem Taschenrechner, dem modernen Ersatz für die Logarithmentafel, allgemein verfügbar.



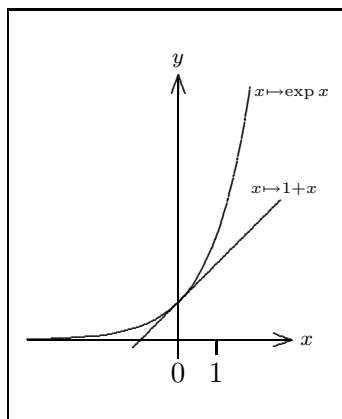
Ebenso bekannt wie der natürliche Logarithmus dürfte seine Umkehrfunktion, die Exponentialfunktion $x \mapsto e^x = \exp(x)$ sein.

Es gilt

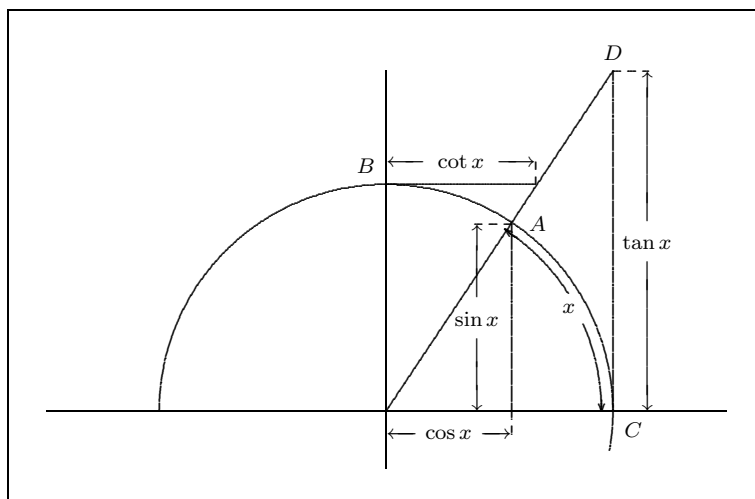
$$\begin{aligned} \ln(\exp(x)) &= x \text{ für alle } x \in \mathbb{R} \\ \exp(\ln y) &= y \text{ für alle } y \in \mathbb{R}_+ . \end{aligned}$$

Die Exponentialfunktion (die ebenfalls auf jedem Taschenrechner verfügbar ist) kann auf vielerlei Art gekennzeichnet werden, z.B. durch

$$e^0 = 1 , \quad \frac{d}{dx} e^x = e^x \text{ für alle } x .$$



Betrachten wir den Einheitskreis und kennzeichnen wir die Punkte P auf ihm durch die Länge des Kreisbogens vom Punkt $(1, 0)$ bis P , $\varphi \in [0, 2\pi)$.

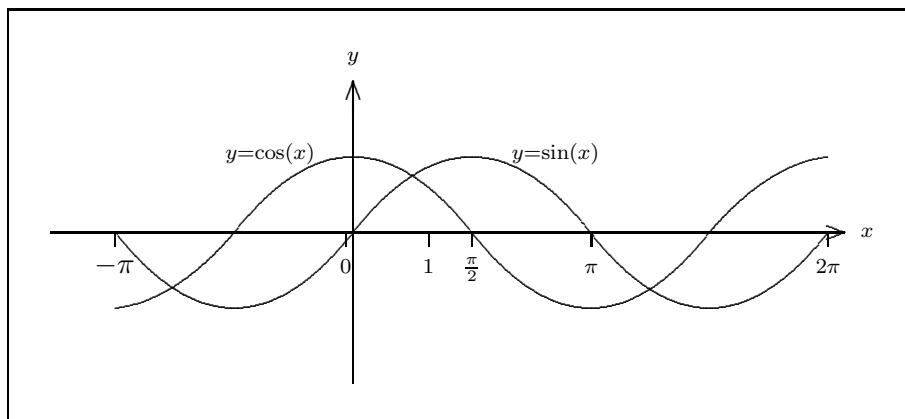


Für die cartesischen Koordinaten x, y gilt dann

$$x = \cos \varphi, \quad y = \sin \varphi, \quad \cos^2 \varphi + \sin^2 \varphi = x^2(P) + y^2(P) = 1 \quad \text{für alle } \varphi.$$

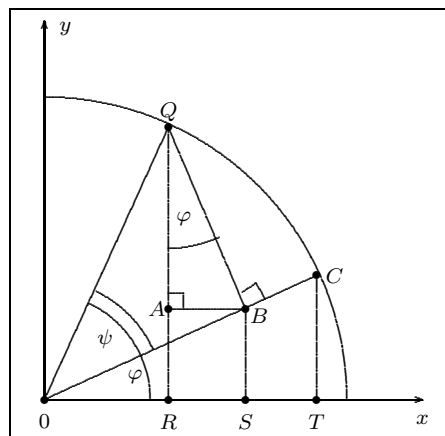
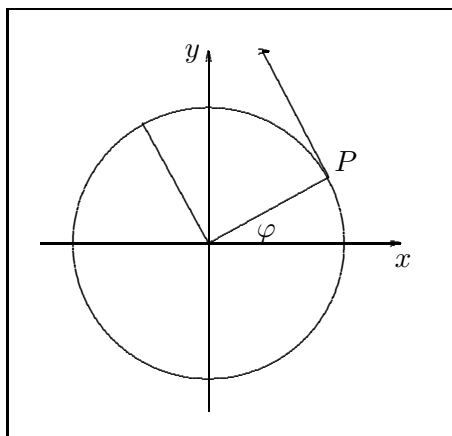
Wenn wir den Einheitskreis mit der Geschwindigkeit 1 durchlaufen, dann gelangen wir zu den Zeitpunkten $\varphi, \varphi + 2\pi, \varphi + 4\pi, \dots$ (und auch zu den Zeitpunkten $\varphi - 2\pi, \varphi - 4\pi, \dots$) in den Punkt P . Für alle $k \in \mathbb{Z}$ haben wir also

$$\cos(\varphi + k \cdot 2\pi) = \cos \varphi, \quad \sin(\varphi + k \cdot 2\pi) = \sin \varphi.$$



Die Geschwindigkeit im Punkt P ist ein Einheitsvektor mit den Komponenten

$$(\dot{x}(P), \dot{y}(P)) = (-\sin \varphi, \cos \varphi) = \left(\cos \left(\varphi + \frac{\pi}{2} \right), \sin \left(\varphi + \frac{\pi}{2} \right) \right)$$



Sinus und Cosinus sind 2π -periodische Funktionen. Der Cosinus ist eine gerade Funktion, der Sinus ein ungerade, d.h. $\cos(-\varphi) = \cos \varphi$, $\sin(-\varphi) = -\sin \varphi$.

Additionstheorem :

$$\sin(\varphi + \psi) = \sin \varphi \cdot \cos \psi + \cos \varphi \cdot \sin \psi$$

$$\cos(\varphi + \psi) = \cos \varphi \cdot \cos \psi - \sin \varphi \cdot \sin \psi$$

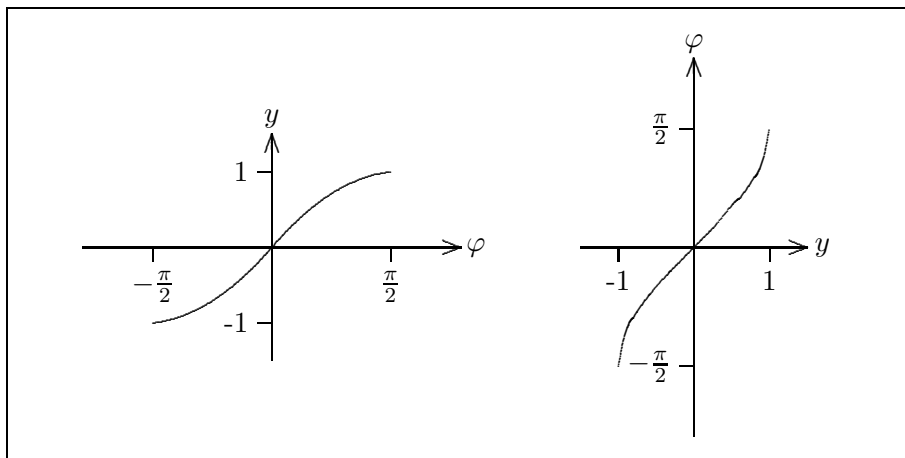
Die Umkehrfunktion von $\sin(\varphi)$ für $\varphi \in \left[-\frac{\pi}{2}, +\frac{\pi}{2}\right]$ heißt der Arcussinus. Der Arcussinus bildet das Einheitsintervall $[-1, +1]$ isoton auf das Intervall $\left[-\frac{\pi}{2}, +\frac{\pi}{2}\right]$ ab

$$\sin(\cdot) : \left[-\frac{\pi}{2}, +\frac{\pi}{2}\right] \longrightarrow [-1, +1]$$

$$\arcsin(\cdot) : [-1, +1] \longrightarrow \left[-\frac{\pi}{2}, +\frac{\pi}{2}\right]$$

$$\arcsin(\sin \varphi) = \varphi \quad \text{für} \quad -\frac{\pi}{2} \leq \varphi \leq +\frac{\pi}{2}$$

$$\sin(\arcsin x) = x \quad \text{für} \quad -1 \leq x \leq +1.$$



Was mancher vielleicht vergessen haben könnte, und daher auch später noch behandelt werden soll, ist der

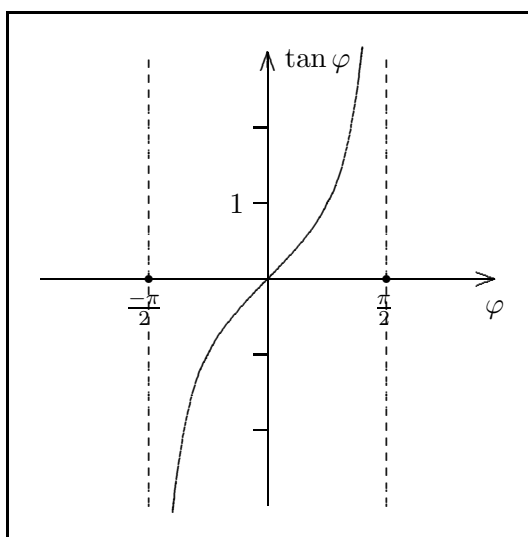
Satz :

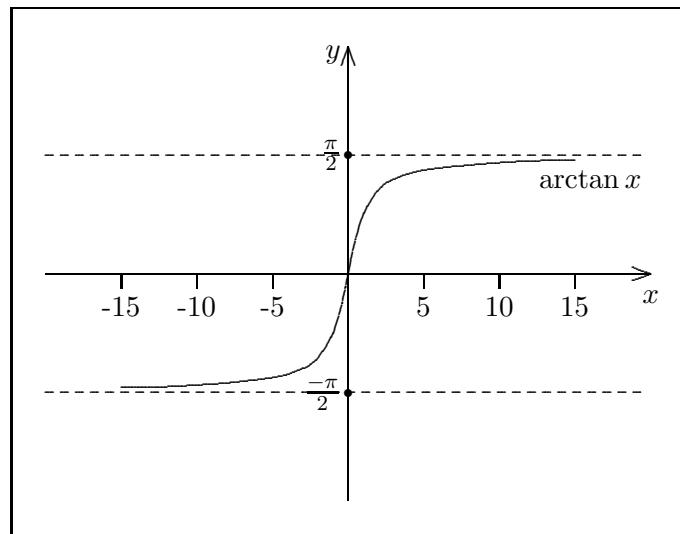
$$\frac{d}{dx} \arcsin x = \frac{1}{\sqrt{1-x^2}} \quad \text{für } |x| < 1 .$$

Beweis mit Hilfe der Kettenregel:

$$\begin{aligned} \sin(\arcsin x) &= x \quad \text{für } |x| < 1 \\ \sin'(\arcsin x) \cdot (\arcsin x)' &= 1 \\ \cos(\arcsin x) &= \sqrt{1 - \sin^2(\arcsin x)} = \sqrt{1 - x^2} . \end{aligned}$$

Einen entsprechenden Satz für den Arcustangens behandeln wir in den Übungen.





Warum müssen Studenten mehr wissen? Welche Fragen schließen sich an das auf der Schule vermittelte Wissen an?

Wir wollen in einem allgemeinerem Rahmen diskutieren, was es mit den Konstruktionen der Integration und Differentiation auf sich hat. Die intuitiven Vorstellungen

$$\begin{aligned}\text{Integral} &= \text{Fläche unter der Kurve} \\ \text{Ableitung} &= \text{Geschwindigkeit der Veränderung} .\end{aligned}$$

sollen auf eine begriffliche Basis gestellt werden, die weitreichende Anwendungen erlaubt.

Es geht nicht darum, über viele weitere elementare Funktionen alle möglichen Einzelheiten zu erfahren. Und es wird auch nicht um die Perfektion der numerischen Berechnung gehen.

1.3 Die Erfindung der Differential- und Integralrechnung zu Ende des 17. Jahrhunderts

Die Differential- und Integralrechnung ist etwa gleichzeitig von Newton und Leibniz erfunden worden. Newton besaß sie zuerst, Leibniz hat sie als erster veröffentlicht. Bei beiden Pionieren steht sie in weitreichenden Zusammenhängen. Die Leibnizsche Darstellungsweise war wesentlich eleganter als die von Newton. Wenige Menschen haben über die Einheit von Form und Inhalt so tief und erfolgreich wie Leibniz nachgedacht. (vgl. Struik S. 121 und S. 125).

Newton gab mit seiner „Philosophiae naturalis principia mathematica“ (1687) der Mechanik eine axiomatische Grundlage. Er zeigte durch strenge mathematische Beweisführung, wie die empirisch aufgestellten Keplerschen Gesetze der Planetenbewegung durch das Gravitationsgesetz erklärt werden konnte. (Kraft = Masse $\times$ Beschleunigung).

Bei Leibniz war die Hauptantriebsfeder seines Lebens die Suche nach einer universellen Methode, mit der man Wissen erlangen, Erfindungen machen und das Wesen der Einheit des Universums verstehen konnte. Die „scientia universalis“, die er aufzubauen versuchte, besaß viele Seiten, und einige von ihnen führten Leibniz zu mathematischen Entdeckungen. Auf

Leibniz gehen viele noch heute übliche mathematische Bezeichnungen zurück (z.B. stammt das Integralzeichen $\int$ (1686) von Leibniz.) Seine Erfindung des „Kalküls“ (Calculus ist noch heute im Englischen die Bezeichnung für die Differential- und Integralrechnung) war das Ergebnis seines Forschens nach einer „lingua universalis“ der Veränderung und insbesondere der Bewegung.

Leibniz fand seinen Kalkül zwischen 1673 und 1676 in Paris unter dem persönlichen Einfluß von Huygens und durch das Studium von Descartes und Pascal. Er wurde dazu durch die Nachricht angeregt, daß Newton im Besitz einer solchen Methode sein sollte. Seine erste Arbeit über den Kalkül veröffentlichte Leibniz 1684 in der Zeitschrift „Acta eruditorum“, die er selbst 1682 gegründet hatte: „Eine neue Methode für Tangenten, die durch gebrochene und irrationale Werte nicht beeinträchtigt wird, und eine merkwürdige Art des Kalküls dafür“ (auf lateinisch natürlich). Die Arbeit enthält die noch heute üblichen Symbole dx, dy und die Differentiationsregeln einschließlich $d(uv) = u \cdot dv + v \cdot du$. Die Regeln der Integralrechnung folgten in einer Buchbesprechung im Jahre 1686.

Nach 1687 war Leibniz eng mit den Brüdern Bernoulli verbunden, die seine Methoden begierig aufnahmen. Noch vor 1700 hatten diese Männer das meiste von dem gefunden, was heute (das sagt Dirk J. Struik, Professor am MIT 1963) den Studenten in der Differential- und Integralrechnung geboten wird, dazu aber wichtige Teile höherer Gebiete einschließlich der Lösung einiger Probleme aus der Variationsrechnung. 1696 erschien das erste Lehrbuch über Differential- und Integralrechnung, die „Analyse des infiniment petits“ des Marquis de l'Hopital, eines Schülers von Johann Bernoulli, ein Werk, das sich auf die Bernoullischen Vorlesungen über den Differentialkalkül stützt. Das erste für den Nichthistoriker wirklich lesbare Werk ist Eulers „Introductio in analysin infinitorum“ (1748). Dieses Buch bereinigte viele strittige Bezeichnungsfragen für immer. Lagrange, Laplace und Gauss übernahmen Eulers Bezeichnungen in allen ihren Werken.

Wir zitieren nun Struik zum 18. Jahrhundert (l.c., S. 131–139):

DAS ACHTZEHNTE JAHRHUNDERT

1. Im achtzehnten Jahrhundert konzentrierte sich die mathematische Produktivität auf die Differential- und Integralrechnung und ihre Anwendung auf die Mechanik. Die bedeutendsten Namen können in der Art eines Stammbaumes angeordnet werden, um ihre geistige Verwandtschaft anzudeuten.

Leibniz (1646–1716);
 die Gebrüder Bernoulli: Jakob (1654–1705), Johann (1667 bis 1748);
 Euler (1707–1783);
 Lagrange (1736–1813);
 Laplace (1749–1827).

In enger Beziehung zu dem Werk dieser Männer stand die Tätigkeit einer Gruppe französischer Mathematiker, zu der insbesondere Clairaut, d'Alembert und Maupertuis gehörten, die ihrerseits mit den Philosophen der Aufklärung verbunden waren. Ihnen sind noch die Schweizer Mathematiker Lambert und Daniel Bernoulli hinzuzurechnen. Die wissenschaftliche Tätigkeit konzentrierte sich gewöhnlich im

Umkreis von Akademien, unter denen die von Paris, Berlin und Petersburg hervorragten. Der Universitätsunterricht spielte demgegenüber eine geringe oder gar keine Rolle. In dieser Zeit wurden einige der führenden Länder Europas von Despoten regiert, die man in beschönigender Manier aufgeklärt genannt hat, von Friedrich II., Katharina der Großen; auch Ludwig XV. und Ludwig XVI. können dazu gerechnet werden. Einen Teil ihres Anspruchs auf Ruhm leiteten diese Despoten aus ihrer Freude daran her, Gelehrte in ihrer Umgebung zu wissen. Diese Freude war eine Art von geistigem Snobismus, der dadurch etwas gemildert wurde, daß sie einiges von der wichtigen Rolle begriffen, welche die Naturwissenschaft und die Angewandte Mathematik bei der Verbesserung der Produktion und der Erhöhung der Schlagkraft der Armee spielen konnten. Man hat beispielsweise gesagt, daß die Vortrefflichkeit der französischen Flotte auf der Tatsache beruhte, daß die Meister des Schiffsbaus beim Bau von Fregatten und Linienschiffen teilweise mathematische Hilfsmittel heranzogen. Eulers Arbeiten sind voll von Anwendungen auf Fragen, die für Heer und Flotte von Bedeutung sind. Die Astronomie spielte weiterhin unter königlicher und kaiserlicher Schirmherrschaft eine hervorragende Rolle bei Anregungen zu mathematischer Forschungsarbeit.

2. Basel in der Schweiz, seit 1263 freie Reichsstadt, war bereits seit langem ein Mittelpunkt der Gelehrsamkeit. Schon zur Zeit des Erasmus war die Basler Universität ein bedeutendes Zentrum. Wie in den Städten Hollands blühten auch in Basel Kunst und Wissenschaft unter der Herrschaft von Kaufmannspratziern. Zu diesen Basler Patriziern gehörte die Kaufmannsfamilie der Bernoullis, die im siebzehnten Jahrhundert aus Antwerpen nach Basel übergesiedelt war, nachdem ihre Heimatstadt von den Spaniern erobert worden war. Seit der zweiten Hälfte des siebzehnten Jahrhunderts bis in unsere Zeit hat diese Familie in jeder Generation Wissenschaftler hervorgebracht. Es ist tatsächlich schwierig, in der ganzen Geschichte der Wissenschaft eine Familie von vergleichbarer Höchstleistung zu finden.

Diese Höchstleistung beginnt mit zwei Mathematikern, Jakob und Johann Bernoulli. Jakob hatte Theologie, Johann Medizin studiert; als aber die Arbeiten von Leibniz in den „Acta eruditorum“ erschienen, faßten beide den Entschluß, Mathematiker zu werden. Sie wurden die ersten bedeutenden Schüler von Leibniz. Im Jahre 1687 übernahm Jakob den Lehrstuhl für Mathematik an der Universität Basel, wo er bis zu seinem Tode (1705) lehrte. 1697 wurde Johann Professor in Groningen; nach dem Tode seines Bruders wurde er dessen Nachfolger auf dem Lehrstuhl in Basel, wo er dreiundvierzig Jahre, bis zu seinem Tode, blieb.

Jakob begann den Briefwechsel mit Leibniz im Jahre 1687. In beständigem Gedankenaustausch mit Leibniz und untereinander — oftmals in heftiger Rivalität miteinander — entdeckten die beiden Brüder nach und nach die in der kühnen Pioniertat von Leibniz enthaltenen Schätze. Die Reihe ihrer Ergebnisse ist lang und enthält nicht nur vieles von dem, was man heute in unseren elementaren Lehrbüchern der Differential- und Integralrechnung findet, sondern auch die Integration von vielen gewöhnlichen Differentialgleichungen. Unter den Leistungen von Jakob sind zu nennen: die Verwendung von Polarkoordinaten, das Studium der Kettenlinie, die schon von Huygens und anderen diskutiert worden war, die Lemniskate (1694) und die logarithmische Spirale. 1690 fand er die sogenannte Isochrone, die von Leibniz 1687 als diejenige Kurve eingeführt worden war, längs der ein Körper mit gleichmäßiger Geschwindigkeit fällt; sie ergab sich als semikubische Parabel. Jakob diskutierte auch isoperimetrische Figuren (1701), die auf ein Problem der Variationsrechnung führten. Die logarithmische Spirale, die die Eigenschaft besitzt, sich bei verschiedenen Transformation zu reproduzieren (ihre Evolute ist wieder eine logarithmische Spirale, so daß beide Fußpunktkurve und Kaustik bezüglich des Pols sind), bereitete Jakob eine derartige Freude, daß er anordnete, diese Kurve mit der Inschrift „eadem mutata resurgo“¹⁾ solle auf seinem Grabstein eingemeißelt werden.

¹⁾ „Ich bleibe dieselbe, auch wenn ich verändert werde.“ Die Spirale auf dem Grabstein sieht jedoch wie eine Archimedische Spirale aus.

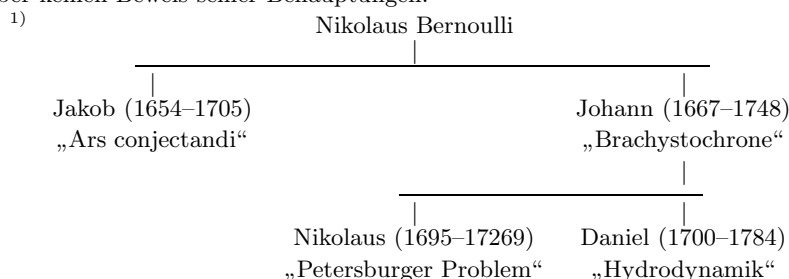
Jakob Bernoulli war auch einer der ersten Bearbeiter der Theorie der Wahrscheinlichkeit, worüber er die „Ars conjectandi“ schrieb, die 1713 posthum veröffentlicht wurde. Im ersten Teil dieses Buches ist die Huygenssche Abhandlung über die Glücksspiele neu abgedruckt; die anderen Teile behandeln Permutationen und Kombinationen und erreichen ihren Höhepunkt im „Bernoullischen Theorem“ über binomiale Verteilungen. Die „Bernoullischen Zahlen“ erscheinen in diesem Buch bei einer Diskussion des Pascalschen Dreiecks.

3. Das Werk von Johann Bernoulli war mit dem seines älteren Bruders eng verbunden, und es ist nicht immer leicht, die Ergebnisse dieser beiden Männer auseinanderzuhalten. Johann wird auf Grund seiner Arbeit zum Problem der Brachystochrone häufig als Entdecker der Variationsrechnung angesehen. Das ist die (ebene) Kurve der kürzesten Fallzeit für einen Massenpunkt, der sich unter dem Einfluß des Schwerfeldes bewegt, und die von Leibniz und den Bernoullis 1697 und in den folgenden Jahren studiert wurde. In dieser Zeit fanden sie die Gleichung der Geodätischen auf einer Fläche.²⁾ Die Lösung des Problems der Brachystochrone ist die Zykloide. Diese Kurve liefert auch die Lösung des Problems der Tautochrone, d.h. derjenigen Kurve, längs der ein Massenpunkt im Gravitationsfeld den tiefsten Punkt nach einer von seinem Startpunkt unabhängigen Zeit erreicht. Huygens hatte diese Eigenschaft der Zykloide entdeckt und sie bei der Konstruktion von tautochronen Pendeluhren verwendet (1673), bei denen die Schwingungsdauer unabhängig von der Amplitude ist.

Unter den Bernoullis, die die Entwicklung der Mathematik beeinflußt haben, befinden sich zwei Söhne von Johann, Nikolaus und vor allem Daniel.¹⁾ Nikolaus wurde nach Petersburg berufen, das erst wenige Jahre vorher von dem Zaren Peter dem Großen gegründet worden war; er blieb dort nur kurze Zeit. Das Problem aus der Theorie der Wahrscheinlichkeit, das er während seines Aufenthaltes in dieser Stadt stellte, ist als Petersburger „Problem“ (oder, dramatischer gesagt, „Paradoxon“) bekannt. Dieser Sohn von Johann starb in jungen Jahren, aber der andere, Daniel, erreichte ein hohes Alter. Bis 1777 war er Professor an der Universität Basel. Daniels reiche wissenschaftliche Tätigkeit war in der Hauptsache der Astronomie, Physik und Hydrodynamik gewidmet. Seine „Hydrodynamica“ erschien 1738, und einer der darin aufgestellten Sätze über den hydraulischen Druck trägt seinen Namen. Im gleichen Jahr stellte er die kinetische Gastheorie auf; mit d'Alembert und Euler entwickelte er die Theorie der schwingenden Saite. Während sein Vater und sein Onkel die Theorie der gewöhnlichen Differentialgleichungen entwickelten, leistete Daniel auf dem Gebiet der partiellen Differentialgleichungen Pionierarbeit.

4. Aus Basel stammte auch der produktivste Mathematiker des achtzehnten Jahrhunderts — wenn nicht aller Zeiten —, Leonhard Euler. Sein Vater studierte Mathematik bei Jakob Bernoulli und Leonhard bei Johann. Als Johanns Sohn Nikolaus 1725 nach Petersburg reiste, folgte ihm der junge Euler

²⁾Newton hatte schon in einer Anmerkung der „Principia“ (II, Satz 35) die Form eines Rotationskörpers diskutiert, der bei einer Bewegung in einer Flüssigkeit den geringsten Widerstand erfährt. Er veröffentlichte aber keinen Beweis seiner Behauptungen.



und blieb bis 1741 an der dortigen Akademie. Von 1741 bis 1766 arbeitete Euler an der Berliner Akademie, die unter der besonderen Schirmherrschaft Friedrichs II. stand; von 1766 bis 1783 war er wieder in Petersburg, nunmehr unter der Ägide der Kaiserin Katharina. Er heiratete zweimal und hatte dreizehn Kinder. Das Leben dieses Akademikers des achzehnten Jahrhunderts war fast ausschließlich den verschiedenen Gebieten der reinen und der angewandten Mathematik gewidmet. Obwohl er 1735 das eine und 1766 das zweite Auge verlor, konnte nichts seine enorme Produktivität unterbrechen. Der erblindete Euler, der über ein phänomenales Gedächtnis verfügte, fuhr fort, seine Entdeckungen zu diktieren. Während seines Lebens erschienen 530 Bücher und Arbeiten; bei seinem Tode hinterließ er viele Manuskripte, die während der nächsten siebenundvierzig Jahre von der Petersburger Akademie veröffentlicht wurden. (Damit erhöhte sich die Anzahl seiner Werke auf 771, aber durch Nachforschungen von Gustav Eneström wurde die Liste auf 886 ergänzt.)

Euler lieferte in allen Gebieten der Mathematik, die zu seiner Zeit existierten, maßgebliche Beiträge. Er veröffentlichte seine Ergebnisse nicht nur in Einzelarbeiten von verschiedener Länge, sondern auch in einer eindrucksvollen Zahl von großen Lehrbüchern, die das während der früheren Zeit gesammelte Material ordneten und geschlossen darstellten. Auf einigen Gebieten war die Eulersche Darstellung so gut wie endgültig. Ein Beispiel dafür ist unsere heutige Trigonometrie mit ihrer Auffassung der trigonometrischen Werte als Verhältniszahlen und ihrer üblichen Schreibweise, die aus Eulers „Introductio in analysin infinitorum“ (1748) stammt. Das überragende Ansehen seiner Lehrbücher bereinigte viele strittige Bezeichnungsfragen der Algebra und Infinitesimalrechnung für immer; Lagrange, Laplace und Gauß kannten Euler und übernahmen seine Bezeichnungen in allen ihren Werken.

Die „Introductio“ von 1748 enthält in ihren zwei Bänden eine Vielzahl von verschiedenen Gegenständen. Sie gibt eine Darstellung der unendlichen Reihen einschließlich derjenigen für e^x , $\sin x$ und $\cos x$ und die Relation $e^{ix} = \cos x + i \sin x$ (die schon von Johann Bernoulli und anderen in ähnlichen Formen entdeckt worden war). Kurven und Flächen werden mit Hilfe ihrer Gleichungen so gewandt untersucht, daß man die „Introductio“ als die erste lehrbuchmäßige Darstellung der analytischen Geometrie betrachten kann. Auch eine algebraische Eliminationstheorie ist darin enthalten. Zu den reizvollsten Teilen dieses Buches gehören das Kapitel über die Zetafunktion und ihre Beziehung zur Theorie der Primzahlen ebenso wie das Kapitel über die „partitio numerorum“. ¹⁾

Ein anderes großes und reichhaltiges Lehrbuch war Eulers Werk „Institutiones calculi differentialis“ (1755), dem weitere drei Bände der „Institutiones calculi integralis“ (1768–1774) folgten. Hier findet man nicht nur unsere elementare Differential- und Integralrechnung, sondern auch eine Theorie der Differentialgleichungen, den Satz von Taylor mit zahlreichen Anwendungen, die Eulersche Summenformel und die Eulersche Γ - und B -Integrale. Der Abschnitt über Differentialgleichungen mit seiner Unterscheidung zwischen „linearen“, „exakten“ und „homogenen“ Gleichungen dient noch heute als Muster unserer einschlägigen elementaren Lehrbücher.

Eulers „Mechanica, sive motus scientia analytice exposita“ (1736) war das erste Lehrbuch, in dem die Newtonsche Dynamik des Massenpunktes mit analytischen Methoden entwickelt wurde. Dieses Lehrbuch enthält die Eulersche Gleichung für einen um einen Punkt rotierenden Körper. Die „Vollständige Anleitung zur Algebra“ (1770), die deutsch geschrieben und einem Diener diktiert wurde, ist zum Vorbild für viele spätere Lehrbücher der Algebra geworden. Es führt bis zur Theorie der kubischen und biquadratischen Gleichungen.

Im Jahre 1744 erschien Eulers „Methodus inveniendi lineas curvas maximi minimive gaudentes“. Das war die erste Darstellung der Variationsrechnung; sie enthielt die Eulersche Differentialgleichung mit vielen Anwendungen einschließlich der Entdeckung, daß das Katenoid und die Schraubenfläche Minimalflächen sind. Viel andere Resultate von Euler finden sich in seinen kleineren Arbeiten, die manchen sogar heute wenig bekannten Edelstein enthalten. Zu den wohlbekannten Entdeckungen

¹⁾Siehe A. Speisers Vorrede zur *Introduction* in: Euler, *Opera* I, 9 (1945).

gehören der Eulersche Polyedersatz, der die Anzal der Ecken (E), Kanten (K) und Flächen (F) eines geschlossenen Polyeders verknüpft ($E+F-K = 2$)²⁾, die Eulersche Gerade im Dreieck, die Kurven konstanter Breite (Euler nannte sie kreisähnliche Kurven) und die Eulersche Konstante C :

$$\lim_{n \rightarrow \infty} \left(1 + \frac{1}{2} + \cdots + \frac{1}{n} - \log n \right) = 0,577216 \dots$$

Einige Arbeit sind mathematischen Unterhaltungen gewidmet (das Königsberger Brückenproblem, Rösselsprünge). Eulers Beiträge zur Zahlentheorie würden allein hinreichen, um ihm eine Nische in der Ruhmeshalle zu sichern; zu seinen Entdeckungen auf diesem Gebiet gehört das quadratische Reziprozitätsgesetz.

Ein großer Teil der Tätigkeit von Euler war der Astronomie gewidmet, wo die Mondtheorie, ein wichtiger Sonderfall des Dreikörperproblems und auch für die Lösung der uralten Frage der Längenbestimmung von Bedeutung, seine besondere Aufmerksamkeit auf sich zog. Die „*Theoria motus plenatarum et cometarum*“ (1774) ist eine Abhandlung über Himmelsmechanik. Verwandt mit dieser Arbeit war Eulers Studium der Anziehung von Ellipsoiden (1738).

Euler hat Bücher über Hydraulik, Schiffsbau und Artillerie geschrieben. 1769–1771 erschienen drei Bände einer „*Dioptrica*“ mit einer Theorie des Strahlenvorlaufs beim Durchgang durch ein Linsensystem. Im Jahre 1739 erschien seine neue Theorie der Musik, von der man gesagt hat, daß sie zu musikalisch für Mathematiker und zu mathematisch für Musiker war. Eulers philosophische Darlegung der wichtigsten Probleme der Naturwissenschaft in seinen „*Briefen an eine deutsche Prinzessin*“ (1760/61 geschrieben) sind ein Musterbeispiel einer populären Darstellung geblieben.

Die enorme Fruchtbarkeit Eulers bildet eine ständige Quelle der Überraschung und Bewunderung für jeden, der versucht hat, sein Werk zu studieren, übrigens eine Aufgabe, die nicht so schwierig ist, wie sie scheint, denn Eulers Latein ist sehr einfach, und seine Bezeichnungen sind fast modern — oder vielleicht sollte man besser sagen, daß die heutigen Bezeichnungen fast die gleichen wie bei Euler sind! Man kann eine lange Liste der bekannten Entdeckungen aufstellen, bei denen Euler die Priorität besitzt, und eine weitere Liste von Ideen, die es noch verdienen, bearbeitet zu werden. Große Mathematiker haben stets ihre Schuld gegenüber Euler betont. Laplace pflegte zu jüngeren Mathematikern zu sagen: „*Lisez Euler, c'est notre maître à tous.*“ (Lest Euler, er ist unser aller Meister.) Und Gauß drückte sich etwas gewichtiger so aus: „Das Studium der Werke Eulers bleibt die beste Schule in den verschiedenen Gebieten der Mathematik und kann durch nichts anderes ersetzt werden.“ Riemann kannte Eulers Werke sehr gut, und einige seiner tiefeschürfenden Werke haben einen Eulerschen Zug. Verleger können kaum etwas besseres tun, als Übersetzungen von einigen Werken Eulers unter Beifügung moderner Kommentare herausbringen.

5. Es ist recht lehrreich, nicht nur einige von Eulers Beiträgen zur Wissenschaft darzulegen, sondern auch einige seiner Schwächen. Unendliche Prozesse wurden im achtzehnten Jahrhundert noch unsorgfältig gehandhabt, und vieles in den Werken der führenden Mathematiker jener Zeit mutet uns wie ein abenteuerlich-enthusiastisches Experimentieren an. Dieses Experimentieren betraf unendliche Reihen, unendliche Produkte und Integrationen sowie den Gebrauch solcher Symbole wie 0 , ∞ , $\sqrt{-1}$. Kann man manchen von Eulers Schlüssen heute zustimmen, so kommen auch andere vor, denen gegenüber wir Vorbehalte machen müssen. Wir stimmen beispielsweise Eulers Feststellung zu, daß $\log n$ unendlich viele Werte besitzt, die alle komplex sind, außer dann, wenn n positiv ist, in welchem Falle einer der Werte reell ist. Euler kam zu dieser Schlußfolgerung in einem Brief an d'Alembert (1747), der behauptet hatte, daß $\log(-1) = 0$ ist. Aber wir können Euler darin nicht folgen, wenn er

²⁾Es war schon Descartes bekannt.

$1 - 3 + 5 - 7 + \dots = 0$ setzt oder wenn er aus

$$n + n^2 + \dots = \frac{n}{1-n} \text{ und } 1 + \frac{1}{n} + \frac{1}{n^2} + \dots = \frac{n}{n-1}$$

auf $\dots + \frac{1}{n^2} + \frac{1}{n} + 1 + n + n^2 + \dots = 0$ schließt.

Doch man muß Zurückhaltung üben und sollte Euler wegen seiner Art, mit unendlichen Reihen umzugehen, nicht zu vorschnell kritisieren; er verwendete einfach nicht immer einige unserer heutigen Konvergenz- oder Divergenzkriterien als Richtschnur für die Gültigkeit der von ihm behandelten Reihen. Viele der von ihm durch unbekümmertes Arbeiten mit Reihen gewonnene Ergebnisse sind durch die moderne Mathematik in strenger Weise umgedeutet und gerechtfertigt worden.

1.4 Die Lage um 1800

Die Mathematiker des 17. und 18. Jahrhunderts waren an Ergebnissen interessiert; sie kümmerten sich nicht allzusehr um die Grundlagen ihrer Arbeit. „Geht voraus, der Glaube wird sich schon einstellen“ soll d’Alembert gesagt haben. (siehe Struik S. 172). Wenn sich Leute wie Euler oder Lagrange doch einmal um Strenge sorgten, waren ihre Argumente nicht immer überzeugend. Die Formeln hatten Vorrang vor den Begriffen.

Struik schreibt (S. 155): „Es ist eine merkwürdige Tatsache, daß einige der führenden Mathematiker gegen Ende des 18. Jahrhunderts dem Gefühl Ausdruck verliehen, der Bereich der Mathematik schiene irgendwie erschöpft zu sein. Die mühevollen anstrengenden Arbeiten von Euler, Lagrange, d’Alembert und anderen hatten schon die meisten wichtigen Sätze geliefert; die großen Standardlehrbücher hatten sie bereits (oder würden es bald tun) im Zusammenhang dargestellt; die wenigen Mathematiker der nächsten Generation würden nur noch geringere Probleme zu lösen haben. Lagrange schrieb 1772 an d’Alembert: „Scheint es Ihnen nicht, daß erhabene Geometrie (die Bezeichnung „Geometrie“ wird im Französischen des 18. Jahrhunderts für die Mathematik insgesamt gebraucht) ein wenig dazu neigt, dekadent zu werden?“ „Sie hat keine andere Stütze als Sie und Herrn Euler.“

Eine Hauptquelle für den „fin de siècle“-Pessimismus sieht Struik in der Tendenz, den Fortschritt der Mathematik allzusehr mit dem der Mechanik und Astronomie gleichzusetzen. Seit den Zeiten des alten Babylon bis zu denen von Euler und Lagrange hatte die Astronomie in der Mathematik die erhabensten Entdeckungen herbeigeführt und angeregt; nunmehr schien diese Entwicklung ihren Gipfel erreicht zu haben. Mathematiker und Mathematikhistoriker sind sich heute wie schon damals in der Einschätzung einig, daß der Übergang vom 18. Jahrhundert ins 19. einen Wendepunkt in der Entwicklung der Wissenschaften darstellt. Diderot, der große Enzyklopädist schrieb in seinem Aufsatz „Zur Interpretation der Natur“:

„Ein großer Umbruch in den Wissenschaften steht bevor. In Anbetracht der gegenwärtigen Bestrebungen der großen Denker, möchte ich behaupten, daß es im nächsten Jahrhundert keine drei großen Mathematiker in Europa geben wird. Diese Wissenschaft wird zum Stillstand kommen an dem Ort, wo die Bernoullis, Euler, Maupertuis, Clairaut, Fontaine, d’Alembert und Lagrange sie hingebracht haben.“

Nicht zur zur Forschung, sondern auch zur Lehre ist um 1800 viel Bemerkenswertes geäußert worden. Der zweiundzwanzigjährige Cauchy verkündete 1811: „Die Arithmetik, die Geometrie, die Algebra, die tanszendente Mathematik (Analysis) sind Wissenschaften, die man als abgeschlossen betrachten kann; es bleibt nur noch übrig, von ihnen nützliche Anwendungen zu machen.“

Aus den Lehrbüchern von S.F. LACROIX (1765–1843) haben Generationen von Studenten der École Polytechnique die Infinitesimalrechnung gelernt. Die Übersetzung ins Englische (1816) brachte in England eine Befreiung von unhaltbaren Lehrtraditionen. Lacroix sagte 1810: „Solche Spitzfindigkeiten, mit denen sich die Griechen abquälten, brauchen wir heute nicht mehr.“

D.J. STRUIK, Professor am MIT, sage 1963 (er war damals schon emeritiert): „Noch vor 1700 hatten die Brüder Bernoulli zusammen mit Leibniz das meiste von dem gefunden, was heute den Studenten der Differential- und Integralrechnung geboten wird.“

1.5 Zum Anliegen der Vorlesung

Die Betrachtungsweisen im A-Zug werden sich an der Mathematikauffassung des 18. Jahrhunderts orientieren. Mancher Student, der kein Mathematiker werden will, könnte aus dem Gesagten den Schluß ziehen, daß das, was im A-Zug geboten wird, (an Inhalt und Methode) eigentlich ausreichen sollte. Der B-Zug beschäftigt sich nämlich mit Fragestellungen, die der reife Cauchy und die anderen Pioniere der strengen Analysis im 19. Jahrhundert entwickelt haben. Der C-zug geht darüber noch wesentlich hinaus; da soll Mathematik des 20. Jahrhunderts vorgestellt werden. Die abstrakte Sichtweise des 20. Jahrhunderts, deren Weg durch G. CANTORS Mengenlehre gebahnt wurde, soll im C-Zug entwickelt werden. Der Umbruch wurde von manchen Mathematikern und Hochschullehrern als höchst dramatisch empfunden. A. SCHÖNFLIES schreibt in *Acta math.* 50 (1927) S. 2: „Es übersteigt nicht das erlaubte Mass, wenn ich sage, dass die Kroneckersche Einstellung den Eindruck hervorbringen musste, als sei Cantor in seiner Eigenschaft als Forscher und Lehrer ein Verderber der Jugend.“ Die Vorzüge der abstrakten Betrachtungsweise wollen wir parallel zu den konkreten Formeln (aus der Zeit vor 1800) und den Anstrengungen um Begründungen im 19. Jahrhundert kennenlernen.

Prüfungsrelevant sind vor allem die Gegenstände des Mittwochszugs; nach der herrschenden Meinung sollen nämlich die Studenten in der Analysis vor allem den Strengebegriff des 19. Jahrhunderts aufnehmen; die Mathematik des 18. Jahrhunderts wird üblicherweise im wesentlichen als Beispielmateriale geschätzt. Einen knappen Überblick über das Prüfungsrelevante finden Sie im Buch von Forster (entstanden aus dem „Regensburger Trichter“; den Schrei nach effektiver Lehre gab es auch schon in den 60er Jahren; und es hat sich in den vergangenen Jahren gezeigt, daß Forsters Buch diesen Forderungen in hervorragendem Maß gerecht wird.) Die Reihenfolge der Themen wird bei uns ganz anders sein als bei Forster (1 und 2). Differentiation kommt bei uns relativ spät. Integration wird zunächst ohne die Fesseln der Strenge des 19. Jahrhunderts im Montagzug vorgestellt, aber mit Anleihen bei der reifen Integrationstheorie des 20. Jahrhunderts, die im Wintersemester von Grund auf entwickelt werden soll.

A.2 Die komplexen Zahlen; Möbiustransformationen

Man sagt manchmal, die „imaginäre Einheit“ i sei die Quadratwurzel aus -1 . Das erklärt aber sehr wenig. Zunächst einmal gibt es für eine negative Zahl keine Quadratwurzel; und wenn es eine gäbe, dann würde $-i$ den Namen ebenso verdienen. Die komplexen Zahlen kann man nur aus ihren Eigenschaften heraus verstehen und nicht durch Reduktion auf Bekanntes. Die komplexen Zahlen sind formale Ausdrücke

$$z = a + ib \quad \text{mit} \quad a, b \text{ reell.}$$

Wenn man $a + i \cdot 0$ mit a identifiziert, dann kann man die reellen Zahlen als spezielle komplexe Zahlen ansehen; die Rechenregeln, die wir definieren werden, sind damit verträglich. Die imaginäre Einheit ist die spezielle komplexe Zahl $i = 0 + i \cdot 1$.

Man definiert für komplexe Zahlen Addition und Multiplikation wie folgt:

Definition Für $z = a + ib$, $w = c + id$ definiert man

$$\begin{aligned} z + w &= (a + c) + i(b + d) \\ z \cdot w &= (ac - bd) + i(ad + bc). \end{aligned}$$

Man kann nun nachrechnen, daß diese Operationen die Menge $\mathbb{C}$ der komplexen Zahlen zu einem Körper machen.

Mit $0 = 0 + i0$ und $1 = 1 + i0$ erhält man für alle z

$$z + 0 = z, \quad z \cdot 1 = z, \quad z \cdot 0 = 0.$$

Man hat somit ein bzgl. der Addition neutrales Element 0 und ein bzgl. der Multiplikation neutrales Element.

Es ist eine reine Rechenaufgabe, die Kommutativ-, Assoziativ- und Distributivgesetze nachzuweisen. Man kann subtrahieren: Wenn $z = a + ib$, dann gilt für $-z = (-a) + i(-b)$ $z + (-z) = 0$.

Für jedes $z \neq 0$ gibt es eine multiplikative Inverse, nämlich

$$\frac{1}{z} = \frac{a}{a^2 + b^2} + i \cdot \frac{(-b)}{a^2 + b^2}.$$

Die folgende Abkürzungen erweisen sich als zweckmäßig.

Für $z = a + ib$ schreibt man

$$\bar{z} = a - ib, \quad |z| = \sqrt{a^2 + b^2} = \sqrt{z \cdot \bar{z}}.$$

Es gelten die Rechenregeln

$$\begin{aligned} \overline{(z + w)} &= \bar{z} + \bar{w}, & \overline{z \cdot w} &= \bar{z} \cdot \bar{w} \\ |z + w| &\leq |z| + |w|, & |z \cdot w| &= |z| \cdot |w|. \end{aligned}$$

Beweis

1) Zu zeigen ist

$$\sqrt{(a+c)^2 + (b+d)^2} \leq \sqrt{a^2 + b^2} + \sqrt{c^2 + d^2},$$

oder äquivalent dazu

$$\begin{aligned} (a+c)^2 + (b+d)^2 &\leq a^2 + b^2 + c^2 + d^2 + 2 \cdot \sqrt{(a^2 + b^2)(c^2 + d^2)} \\ ac + bd &\leq \sqrt{(a^2 + b^2)(c^2 + d^2)} \\ a^2 \cdot c^2 + b^2 \cdot d^2 + 2acbd &\leq a^2 \cdot c^2 + b^2 \cdot c^2 + a^2 \cdot d^2 + b^2 \cdot d^2 \end{aligned}$$

Dies aber ergibt sich aus $(ad - bc)^2 \geq 0$.

$$2) |z \cdot w|^2 = z \cdot w \cdot \bar{z} \cdot \bar{w} = (z \cdot \bar{z}) \cdot (w \cdot \bar{w}) = |z|^2 \cdot |w|^2. \quad \blacksquare$$

Betrachten wir die $z \in \mathbb{C}$ mit $|z| = 1$. Bekanntlich gibt es zu jedem Paar reeller Zahlen (a, b) mit $a^2 + b^2 = 1$ genau ein $\varphi \in (-\pi, +\pi]$ mit

$$a = \cos \varphi, \quad b = \sin \varphi.$$

Satz Zu jedem $z \neq 0$ gibt es ein bis auf ganzzahliges Vielfaches von 2π eindeutig bestimmtes φ , so daß

$$z = |z|(\cos \varphi + i \sin \varphi).$$

Man schreibt zur Abkürzung

$$\cos \varphi + i \sin \varphi = e^{i\varphi}.$$

Bemerke $e^{i\pi} + 1 = 0$; eine merkwürdige Beziehung zwischen den berühmtesten Zahlen: $0, 1, i, e, \pi$.

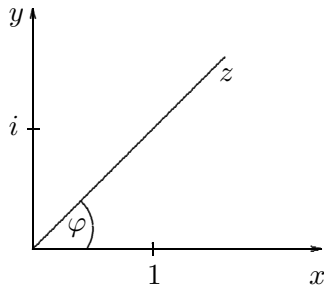
Satz Wenn $z = |z| \cdot e^{i\varphi}$, $w = |w| \cdot e^{i\psi}$, dann gilt

$$z \cdot w = |z \cdot w| e^{i(\varphi + \psi)}.$$

Beweis durch die Additionstheoreme für die trigonometrischen Funktionen.

Bemerkung $(\cos \varphi + i \sin \varphi)^n = \cos(n\varphi) + i \sin(n\varphi)$.

Es war die Idee von Gauß, die komplexen Zahlen $z = x + iy$ als Punkte in der Ebene darzustellen,



wo x und y die kartesischen Koordinaten sind und $|z|$, φ die Polarkoordinaten

$$z = x + iy = |z| \cdot e^{i\varphi}.$$

Die Addition ist die übliche Vektoraddition. Die Multiplikation addiert die Winkel und multipliziert die Abstände vom Ursprung.

Diese geometrische Repräsentation erklärt nicht sehr viel. Sie ist kaum geeignet, plausibel zu machen, daß

$$(\mathbb{C}, 0, 1, +, \cdot)$$

ein Körper ist. Das formale Nachrechnen der Rechengesetze wird durch die geometrische Interpretation nicht überflüssig.

Die geometrische Darstellung ist aber für andere Zwecke nützlich, z.B. dann, wenn man nach den Lösungen der Gleichung

$$z^n = a = |a|e^{i\varphi} \quad (\text{für ein } a \neq 0)$$

fragt. Es gibt offenbar genau n verschiedene Lösungen

$$z_j = \sqrt[n]{|a|} \cdot e^{i\left(\frac{\varphi}{n} + j \cdot \frac{2\pi}{n}\right)} \quad j = 1, 2, \dots, n.$$

Die Lösungen der Gleichung $z^n = 1$ heißen die n -ten *Einheitswurzeln*:

$$\text{Für } n = 3: \quad \left(-\frac{1}{2} + i \frac{1}{2}\sqrt{3}\right)^3 = \left(-\frac{1}{2} - i \frac{1}{2}\sqrt{3}\right)^3 = 1^3 = 1$$

$$\text{Für } n = 4: \quad i^4 = (-1)^4 = (-i)^4 = 1^4 = 1$$

Wir studieren nun eine Abbildung, die durch die Operationen mit komplexen Zahlen vermittelt werden

$$1) \quad z \mapsto \bar{z} \quad (\text{„Spiegelung an der reellen Achse“})$$

$$2) \quad z \mapsto \frac{1}{\bar{z}} \quad (\text{„Spiegelung am Einheitskreis“})$$

Die Punkte auf dem Einheitskreis sind Fixpunkte. Jeder Strahl vom Ursprung aus wird auf sich abgebildet. Man fügt gerne den Punkt ∞ zu $\mathbb{C}$ hinzu $\overline{\mathbb{C}} = \mathbb{C} \cup \{\infty\}$ und sagt, daß der

unendlichferne Punkt bei der Spiegelung am Einheitskreis in den Nullpunkt abgebildet wird, während der Nullpunkt auf den unendlichfernen Punkt ∞ abgebildet wird.

$$3) \quad z \longmapsto \frac{1}{z}.$$

Diese Abbildung ergibt sich durch Hintereinanderausführung der Spiegelungen 1) und 2) (in beliebiger Reihenfolge)

$$4) \quad z \longmapsto z + b.$$

Es handelt sich um eine Verschiebung; der unendlichferne Punkt bleibt fest.

$$5) \quad z \longmapsto az.$$

Es handelt sich um eine Drehstreckung; der unendlichferne Punkt bleibt fest.

Definition Wenn a, b, c, d komplexe Zahlen sind mit $ad - bc \neq 0$, dann heißt die Abbildung

$$z \longmapsto \frac{az + b}{cz + d}$$

eine Möbiustransformation.

Man macht die Möbiustransformationen zu Abbildungen von $\overline{\mathbb{C}}$ auf $\overline{\mathbb{C}}$ durch die folgenden Festsetzungen:

Wenn $c = 0$, dann wird ∞ auf ∞ abgebildet.

Wenn $c \neq 0$, dann wird ∞ auf $\frac{a}{c}$ und der Punkt $-\frac{d}{c}$ auf ∞ abgebildet.

Bemerkung Wenn man die Koeffizienten a, b, c, d mit einem gemeinsamen Faktor multipliziert, verändert man die Möbiustransformation nicht. Man kann $ad - bc = 1$ annehmen, wenn man das praktisch findet.

Satz Die Möbiustransformationen sind umkehrbar eindeutige Abbildungen von $\overline{\mathbb{C}}$ auf sich. Die Umkehrabbildung ist ebenfalls eine Möbiustransformation. Das Hintereinanderschalten von Möbiustransformationen liefert wieder Möbiustransformationen.

Kurz: Die Gesamtheit der Möbiustransformationen bildet eine Gruppe.

Beweis

$$1) \quad w = \frac{az + b}{cz + d} \iff w(cz + d) = az + b \iff (a - cw)z = b + dw$$

$$z = \frac{dw - b}{-cw + a} \text{ ist die inverse Möbiustransformation.}$$

(Der Punkt $+\infty$ ist gesondert zu untersuchen.)

$$2) \quad \frac{a \cdot \frac{\alpha z + \beta}{\gamma z + \delta} + b}{c \cdot \frac{\alpha z + \beta}{\gamma z + \delta} + d} = \frac{(a\alpha + b\gamma)z + (a\beta + b\delta)}{(c\alpha + d\gamma)z + (c\beta + d\delta)}$$

Hinweis Vergleiche mit dieser Kompositionsregel die Regel für die Matrizenmultiplikation

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix} \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix} = \begin{pmatrix} a\alpha + b\gamma & a\beta + b\delta \\ c\alpha + d\gamma & c\beta + d\delta \end{pmatrix}$$

Bekanntlich ist die Multiplikation von (nichtsingulären) 2×2 -Matrizen nicht kommutativ.

Satz Jede Möbiustransformation läßt sich durch Zusammensetzung von Abbildungen der folgenden Typen gewinnen

- i) $z \mapsto \alpha z$ („Drehstreckung“)
- ii) $z \mapsto z + \beta$ („Verschiebung“)
- iii) $z \mapsto \frac{1}{z}$

Beweis Für $c \neq 0$ schreiben wir

$$\frac{az + b}{cz + d} = \frac{a}{c} - \frac{ad - bc}{c} \cdot \frac{1}{cz + d}.$$

Für $c = 0$ ist die Aussage trivial.

Aufgaben zu den Möbiustransformationen

Aufgabe 1 : (Spezielle Transformationen)

(„Abbildung auf“ steht hier kurz für „surjektive Abbildung auf“)

- a) Die Abbildung $z \mapsto \frac{1}{z}$ bildet das Innere des Einheitskreises $\{z : |z| < 1\}$ auf das Äußere des Einheitskreises $\{z : |z| > 1\}$ ab. Es gibt genau zwei Fixpunkte, nämlich $+1$ und -1 .
- b) Die „Cayley-Transformation“

$$z \mapsto \frac{z - i}{z + i}$$

bildet die reelle Achse auf den Einheitskreis ab; die obere Halbebene wird auf das Innere des Einheitskreises abgebildet, die untere auf das Äußere.

- c) Für z_0 mit $|z_0| < 1$ und $\varphi \in \mathbb{R}$ bildet

$$z \mapsto e^{i\varphi} \cdot \frac{z - z_0}{\bar{z}_0 \cdot z - 1}$$

das Innere des Einheitskreises auf sich ab.

Jede Möbiustransformation, welche das Innere des Einheitskreises auf sich abbildet, ist von dieser Gestalt.

Aufgabe 2 : Jede Möbiustransformation, die nicht die Identität ist, hat genau einen oder genau zwei Fixpunkte.

(Bemerke : Bei den Drehstreckungen und den Verschiebungen ist ∞ ein Fixpunkt.)

Aufgabe 3 :

a) Für $b \in \mathbb{R}$, $\alpha \in \mathbb{C}$ mit $\alpha \neq 0$ nennen wir

$$\{z : \bar{\alpha}z + \alpha\bar{z} + b = 0\} \cup \{\infty\}$$

eine Gerade in der komplexen Ebene oder auch einen Kreis durch den unendlich fernen Punkt.

Für $\alpha \in \mathbb{C}$, a, b reell mit $D := \alpha\bar{\alpha} - ab > 0$ nennen wir

$$\{z : az\bar{z} + \bar{\alpha}z + \alpha\bar{z} + b = 0\}$$

einen Kreis in der komplexen Ebene.

Machen Sie das durch die Ersetzung $z = x + iy$ plausibel.

Welche Rolle spielt D ?

b) Zeigen Sie: Jede Möbiustransformation bildet Kreise in Kreise ab.

Hinweis Es genügt, die Behauptung für Drehstreckungen, Verschiebungen und die Abbildung $z \mapsto \frac{1}{z}$ nachzuweisen.

Aufgabe 4 : Für paarweise verschiedene z_1, z_2, z_3, z_4 definieren wir das Doppelverhältnis

$$DV(z_1, z_2, z_3, z_4) := \frac{z_4 - z_3}{z_4 - z_1} \cdot \frac{z_2 - z_1}{z_2 - z_3}.$$

Seien w_1, w_2, w_3, w_4 die Bilder bzgl. einer Möbiustransformation. Zeigen Sie

$$DV(w_1, w_2, w_3, w_4) = DV(z_1, z_2, z_3, z_4).$$

Hinweis Es genügt, die Behauptung für Drehstreckungen, Verschiebungen und $z \mapsto \frac{1}{z}$ nachzuweisen.

Anhang : Zur Geschichte des Zahlbegriffs

Die Zahlen und das Rechnen mit Zahlen haben eine sehr lange Geschichte. Die abendländische Mathematik beginnt in Griechenland. (Ich stütze mich auf DIRK J. STRUIK: Abriss der Geschichte der Mathematik, Berlin 1963)

„Der traditionelle Vater der griechischen Mathematik ist der Kaufmann Thales von Milet, der der Legende nach in der ersten Hälfte des sechsten Jahrhunderts Babylon und Ägypten besuchte. Das Studium der Mathematik in der griechischen Frühzeit verfolgt ein hauptsächliches Ziel, nämlich die Gewinnung einer aus einem Vernunftgebäude ableitbaren Einsicht in die Stellung des Menschen innerhalb des Kosmos. Die Mathematik diente dazu, Ordnung im Chaos zu schaffen, Ideen in logischen Ketten anzuordnen und fundamentale Prinzipien zu entdecken. Sie war die am stärksten verstandesmäßig bestimmte unter allen Wissenschaften, und obwohl kaum ein Zweifel daran bestehen kann, daß die griechischen Kaufleute auf ihren Handelsreisen die orientalische Mathematik kennenlernten, so fanden sie doch bald heraus, daß die Völker des Orients die meiste Arbeit zur verstandesmäßigen Durchdringung ungetan gelassen hatten.“ (Struik S. 33)

„Zum erstenmal in der Geschichte untersucht eine Gruppe von kritischen Menschen, die Sophisten, die weniger als irgendeine andere frühere Gruppe von Gelehrten durch den Einfluß der Tradition gehemmt war, Probleme mathematischer Natur mehr im Geiste des Verstehens als in dem der Nützlichkeit.“

Aus dieser Zeit (mehr als ein Jahrhundert vor EUKLID) ist eine Schrift des ionischen Philosophen Hippokrates von Chios erhalten. Das Fragment, welches einen hohen Grad der Vollkommenheit des mathematischen Denkens beweist, behandelt ein Problem, welches eine unmittelbare Beziehung zum Problem der Quadratur des Kreises hat. Das Problem der Quadratur des Kreises, das Problem der Dreiteilung des Winkels und das Problem der Verdoppelung des Würfels werden die „drei berühmten mathematischen Probleme des Altertums“ genannt.

„Wahrscheinlich außerhalb der Gruppe der Sophisten, die in gewissem Umfange mit der demokratischen Bewegung verbunden war, stand eine andere Gruppe von mathematisch interessierten Philosophen, die Beziehungen zu aristokratischen Strömungen besaß. Sie nannten sich selbst Pythagoreer nach dem ziemlich mythischen Gründer der Schule, Pythagoras, der vermutlich Wissenschaftler und aristokratischer Politiker war. Während die Sophisten nachdrücklich die Realität der Veränderung vertraten, verlegten die Pythagoreer das Schwergewicht ihrer Studien auf die unveränderlichen Elemente in Natur und Gesellschaft. In ihrer Suche nach den ewigen Gesetzen des Kosmos studierten sie Geometrie, Arithmetik, Astronomie und Musik (das „Quadrivium“). ... Ihre Arithmetik war eine in hohem Maße spekulative Wissenschaft, die mit der gleichzeitigen babylonischen Rechentechnik wenig zu tun hatte. Die Zahlen wurden in Klassen eingeteilt, in ungerade, gerade,

geradzahlig oft gerade, ungeradzahlig of ungerade, Primzahlen und zusammengesetzte Zahlen, vollkommene, befreundete, Dreieck-, Quadrat-, Fünfeckzahlen usw. Die Pythagoreer untersuchten ihre Eigenschaften, wobei sie ihr besonderes Merkmal des Zahlenmystizismus hinzufügten und sie zum Mittelpunkt einer kosmischen Philosophie machten, die alle Beziehungen auf Zahlenbeziehungen zurückzuführen versuchte („Alles ist Zahl“).“ (Struik S. 38)

„Was nun den Satz des Pythagoras betrifft, so schrieben die Pythagoreer seine Entdeckung ihrem Meister zu, von dem angenommen wurde, er haben den Göttern als Zeichen seiner Dankbarkeit hundert Ochsen geopfert. Wie wir bereits sahen, war der Satz schon im Babylon Hammurabis bekannt, aber der erste allgemeine Beweis kann sehr wohl der Schule der Pythagoreer entstammen.

Die wichtigste den Pythagoreern zugeschriebene Entdeckung war aber die Entdeckung des Irrationalen mit Hilfe von inkommensurablen Strecken. ...“ (siehe unten bei der Geschichte von $\sqrt{2}$).

Sie kam noch zu einer anderen Schwierigkeit hinzu, die aus Erörterungen über die Realität der Veränderungen entstanden war, d.h. aus Erörterungen, die die Philosophen seit jener Zeit bis in unsere Tage beschäftigt haben.

Diese Schwierigkeit wird ZENO VON ELEA (um 450), einem Schüler von Parmenides, zugeschrieben. Zeno war konservativer Philosoph, nach dessen Lehre der menschliche Verstand nur das absolute, unveränderliche Sein der Dinge erkennen kann, während alle Veränderungen nur Schein sind. Sie gewann mathematische Bedeutung, als im Zusammenhang mit solchen Fragen wie der Bestimmung des Pyramidenvolumens unendliche Prozesse studiert werden mußten. Hier gerieten die Zenonischen Paradoxien in Widerspruch mit einigen älteren intuitiven Begriffen bezüglich des unendlich Kleinen und unendlich Großen. ... Die Kritik des Zeno zweifelte die Zulässigkeit dieser Begriffe an, und seine vier Paradoxien erregten ein solches Aufsehen, daß man die Auswirkungen davon noch heute beobachten kann. Sie sind durch Aristoteles überliefert worden und als Paradoxien des Achilles, des fliegenden Pfeils, der (unendlich oft wiederholten) Halbierung und des Stadions bekannt. Sie sind so formuliert, daß sie Widersprüche in den Begriffen der Bewegung und der Zeit hervorkehren; es wird dabei kein Versuch unternommen, die Widersprüche aufzulösen. ...

Die Überlegungen von Zeno machten klar, daß eine endliche Strecke in unendlich viele kleine Stecken zerlegt werden kann, von denen jede eine endliche Länge besitzt. Sie zeigten außerdem, daß es schwierig zu erklären ist, was man eigentlich damit meint, wenn man sagt, daß eine Linie aus Punkten „zusammengesetzt“ ist. Es ist sehr wahrscheinlich, daß Zeno selbst gar keine Vorstellung von den mathematischen Konsequenzen seiner Überlegung hatte. Probleme, die zu seinen Paradoxien führen, sind nämlich im Verlauf von philosophischen und theologischen Diskussionen regelmäßig aufgetaucht. Sie sind bekannt als jene Fragen, die die Beziehungen zwischen dem potentiellen und dem aktuellen Unendlich

betreffen., (Struik S. 41)

„Die scholastischen Schriftsteller des Mittelalters, insbesondere THOMAS VON AQUINO übernahmen Aristoteles' Satz „infinitum actu non datur“ („Es gibt keine aktuelle Unendlichkeit“), betrachten vielmehr jedes Kontinuum als potentiell unbegrenzt teilbar. Folglich gab es für sie keine kleinste Strecke, da jeder ihrer Teile wieder die Eigenschaften einer Strecke besitzt. Somit konnte ein Punkt nicht Teil einer Strecke sein, da er unteilbar ist: „ex indivisibilibus non potest comparari aliquod continuum“ („Ein Kontinuum kann nicht aus Indivisibilen bestehen.“) Ein Punkt konnte durch Bewegung eine Strecke erzeugen. Solche Spekulationen beeinflussten die Erfinder Infinitesimalrechnung im siebzehnten Jahrhundert und die Philosophen des Unendlichen im neunzehnten.“ (Struik S. 94)

Die Geschichte der imaginären Zahlen beginnt im 16. Jahrhundert. LUCA DE PACIOLI beendet sein Buch „Summa de arithmetica“ (1494), in dem er den Wissensstand der mittelalterlichen Rechenmeister über Arithmetik, Algebra und Trigonometrie zusammenfaßt und kommentiert, mit der Bemerkung, daß die Lösung der Gleichungen $x^3 + mx = n$ und $x^3 + n = mx$ beim derzeitigen Stand der Wissenschaft genauso unmöglich sei wie die Quadratur des Kreises. Struik schreibt dazu (S. 97)

„An dieser Stelle setzte die Arbeit der Mathematiker an der Universität Bologna ein. Diese Universität war um die Wende des fünfzehnten Jahrhunderts eine der größten und berühmtesten in Europa. Allein ihre astronomische Fakultät hatte zeitweise sechzehn Lektoren. Aus allen Teilen Europas strömten die Studenten herbei, um die Vorlesungen zu hören – und zu den öffentlichen Disputationen, die ebenfalls die Aufmerksamkeit großer, sportlich eingestellter Hörermassen auf sich lenkten. Unter diesen Studenten befanden sich zeitweise Pacioli, Albrecht Dürer und Kopernikus. Charakteristisch für dieses neue Zeitalter war der Wunsch, nicht nur das klassische Erbe aufzunehmen, sondern zugleich Neues zu schaffen und über die von den Klassikern abgesteckten Grenzen hinaus vorzudringen. Die Erfindung der Buchdruckerkunst und die Entdeckung Amerikas waren Beispiele für derartige Möglichkeiten. War es möglich, in der Mathematik Neues zu finden? Griechen und Orientalen hatten ihren Scharfsinn an der Lösung der kubischen Gleichung versucht, aber sie hatten nur einige Spezialfälle numerisch lösen können. Die Mathematiker in Bologna versuchten nun, die allgemeine Lösung zu finden.

Dieses kubischen Gleichungen konnten sämtlich auf drei Typen reduziert werden:

$$x^3 + px = q, \quad x^3 = px + q, \quad x^3 + q = px$$

wobei p und q positive Zahlen waren.“

Die genaue Geschichte der Entdeckung braucht uns hier nicht zu interessieren. Die Lösung wird heute als Cardanische Formel bezeichnet, die im Falle $x^3 + px = q$ folgende Form hat

$$x = \sqrt[3]{\sqrt{\frac{p^3}{27} + \frac{q^2}{4}} + \frac{9}{2}} - \sqrt[3]{\sqrt{\frac{p^3}{27} + \frac{q^2}{4}} - \frac{9}{2}}.$$

Wie man sieht, führte die Lösung Größen der Form $x = \sqrt[3]{a + \sqrt{b}}$ ein, die von den Euklidischen $\sqrt{a + \sqrt{b}}$ verschieden waren. Struik schreibt:

Es ist eine merkwürdige Tatsache, daß die erste Einführung der komplexen Zahlen in der Theorie der kubischen Gleichungen geschah, und zwar gerade in dem Falle, wo es klar war, daß reelle Lösungen existieren (wenn auch in unerkennbarer Form) und nicht in der Theorie der quadratischen Gleichungen, wo sie in unseren heutigen Lehrbüchern eingeführt werden.“

RAFFAEL BOMBELLI, der letzte der großen Bologneser Mathematiker, führte in seiner „Algebra“ (1572) eine konsequente Theorie der rein imaginären Zahl ein, was ihm erlaubte, den irreduziblen Fall zu behandeln. Bombellis Buch wurde viel gelesen; Leibniz benutzte es und Euler zitiert es in seiner eigenen Algebra. Von da an verloren die komplexen Zahlen etwas von ihrem übernatürlichen Charakter. Uneingeschränkte Anerkennung als Zahlen fanden sie aber erst im 19. Jahrhundert. Einen ganz wesentlichen Anteil hatte dabei C.F. GAUSS, der die komplexen Zahlen durch Punkte in der Ebene darstellte (Gaußsche Zahlenebene) und 1799 den sog. Fundamentalsatz der Algebra bewies.

Die „Entdeckung“ der komplexen Zahlen im Geiste der griechischen Mathematik hatte zunächst keinen entscheidenden Einfluß auf die Entwicklung der Wissenschaften. Wichtiger waren zwei „Erfindungen“ im Geiste der „Rechenhaftigkeit“ des fünfzehnten und sechzehnten Jahrhunderts. Den Ausdruck „Rechenhaftigkeit“ hat W. SOMBART 1913 in seinem Buch „Der Bourgeois“ geprägt. Er weist auf eine Bereitschaft zum Rechnen hin, auf den Glauben an die Nützlichkeit der Beschäftigung und Arithmetik. Handel, Navigation, Astronomie und Feldmessung hatte die Bevölkerung der wachsenden Städte für das Zählen und das Berechnen interessiert. Die großen Erfindungen waren die Dezimalzahlen und die Logarithmen.

Das 17. und 18. Jahrhundert können wir in unserer kurzen Geschichte des Zahlbegriffs überspringen. Der Zahlbegriff wurde in dieser Zeit nicht thematisiert. Ein neues Nachdenken über die Zahlen auf verschiedenen Ebenen wurde von C.F. GAUSS (1777–1855) initiiert. Zum einen ergaben sich neue Einblicke in die schon von den Griechen gestellten Fragen nach der Konstruierbarkeit von Figuren der Geometrie mit Zirkel und Lineal. Zum anderen lieferte die Dissertation den ersten strengen Beweis für den sogenannten Fundamentalsatz der Algebra, also für den Satz, daß jede algebraische Gleichung vom Grad n wenigstens eine und damit n Wurzeln besitzt. Der Satz selbst geht auf Albert Girard, den Herausgeber der Werke von Stevin („Invention nouvelle en algèbre“, 1629) zurück. d'Alembert hatte im Jahre 1746 versucht, einen Beweis dafür zu geben. Gauß liebte diesen Satz, gab später zwei weitere Beweise und kam 1849 auf seinen ersten Beweis zurück. Der dritte Beweis (1816) verwendete komplexe Integrale und zeigte die frühzeitige Beherrschung der Theorie der komplexen Zahlen durch Gauß.

Die „Disquisitiones arithmeticae“ stellen alle von den Vorgängern von Gauß stammenden Meisterleistungen in der Zahlentheorie zusammen und bereicherten sie in einem derartigen Umfange, daß man manchmal den Beginn der modernen Zahlentheorie von der Veröffentlichung dieses Buches ab rechnet.“ (Struik S.161)

A.3 Polynome, gebrochenrationale Funktionen, Partialbruchzerlegung

Für einen *Algebraiker* ist ein Polynom (vom Grade $\leq n$ in einer Unbestimmten) ein *formaler Ausdruck* der Gestalt

$$a_0 + a_1x + a_2x^2 + \dots + a_nx^n .$$

Die a_i heißen die Koeffizienten, x ist die Unbestimmte. Man kann ein Polynom vom Grade $\leq n$ auch als ein Polynom vom Grade $\leq n+1$ auffassen, indem man $a_{n+1} = 0$ setzt. Das größte n mit $a_n \neq 0$ heißt der (genaue) Grad des Polynoms.

Für den Algebraiker sind die Polynome *Rechengrößen*. Man kann sie addieren und multiplizieren. Für

$$p = a_0 + a_1x + \dots + a_nx^n , \quad q = b_0 + b_1x + \dots + b_nx^n$$

setzt man

$$p + q = (a_0 + b_0) + (a_1 + b_1)x + \dots + (a_n + b_n)x^n$$

$$p \cdot q = c_0 + c_1x + \dots + c_{2n}x^n$$

mit $c_0 = a_0b_0$, $c_1 = a_0b_1 + a_1b_0$, $c_2 = a_0b_2 + a_1b_1 + a_2b_0$, $\dots$.

Solches Rechnen ist wohldefiniert, wenn man die Koeffizienten addieren und multiplizieren kann, also z.B. wenn es sich um ganze Zahlen, rationale Zahlen, reelle Zahlen oder komplexe Zahlen handelt. Die Rechenregeln für die Koeffizienten liefern entsprechende Rechenregeln für die Polynome: Kommutativität, Assoziativität von Addition und Multiplikation sowie Distributivität.

Übrigens kennen auch die Algebraiker den Begriff der Ableitung eines Polynoms; es ist eine formale Operation, die nichts mit Grenzwerten zu tun hat. Die Regeln sind die bekannten. $(\sum a_n x^n)' = \sum a_n \cdot n \cdot x^{n-1}$.

Die *Analytiker* beschäftigen sich vorwiegend mit Polynomen mit reellen oder komplexen Koeffizienten. Sie fassen die Polynome als Funktionen oder als Abbildungen auf, im komplexen Fall als Abbildungen von $\mathbb{C}$ nach $\mathbb{C}$

$$\mathbb{C} \ni z \longmapsto p(z) = a_0 + a_1z + \dots + a_nz^n \in \mathbb{C} .$$

Geschichtliches Die Dissertation von Gauß lieferte 1799 den ersten strengen Beweis für den sogenannten Fundamentalsatz der Algebra, also für den Satz, daß jedes komplexe Polynom vom Grade n wenigstens eine Nullstelle besitzt. Der Satz selbst geht auf ALBERT GIRARD, den Herausgeber der Werke von Stevin („Invention nouvelle en algèbre“, 1629) zurück. d’Alembert hatte im Jahre 1746 versucht, einen Beweis dafür zu geben. Gauß liebte

diesen Satz, gab später zwei weitere Beweise und kam 1849 auf seinen ersten Beweis zurück. Der dritte Beweis (1816) verwendete komplexe Integrale. — Nach dem heutigen Verständnis von Algebra ist der Fundamentalsatz kein Satz aus der Algebra sondern ein Satz aus der komplexen Analysis. Der Name kommt daher, daß die Suche nach den Wurzeln algebraischer Gleichungen

$$a_0 + a_1 z + \dots + a_n z^n = 0$$

in der Renaissance als die große Herausforderung an die Algebra angesehen worden war. Das Buch „Summa de Arithmetica“ (1494) des Franziskanermönchs LUCA DE PACIOLI, eines der ersten gedruckten mathematischen Bücher überhaupt, endet mit der Bemerkung, daß die Lösung der Gleichungen

$$x^3 + mx = n \quad \text{und} \quad x^3 + n = mx$$

beim derzeitigen Stand der Wissenschaft genauso unmöglich sei wie die Quadratur des Kreises. Das Problem fand immer wieder die Aufmerksamkeit der Mathematiker. Schon die Griechen und die Orientalen hatten ihren Scharfsinn an der Lösung der kubischen Gleichung versucht, aber sie hatten nur einige Spezialfälle numerisch lösen können. Die Mathematiker in Bologna versuchten nun um 1500, die allgemeine Lösung zu finden. Und sie hatten einige Erfolge (siehe Struik, S. 96 ff). Die Anstrengungen gaben nicht sofort, sondern in langen Inkubationszeiten Anstöße zur Herausbildung des Begriffs der komplexen Zahl. Das Geheimnis, das die komplexen Zahlen über Jahrhunderte umgeben hatte, wurde dann endgültig 1831 von Gauß dadurch beseitigt, daß er die komplexen Zahlen als Punkte einer Ebene darstellte (siehe Struik S. 163).

Horners Schema Eine wichtige Manipulation mit Polynomen wurde 1819 von WILLIAM G. HORNER veröffentlicht; Horner wußte wohl nicht, daß er auf eine für eintausend Jahre chinesischer Mathematik typische Methode gestoßen war (siehe Struik S. 83). Das Hornersche Schema steht schon bei den alten Chinesen im Zusammenhang mit der Methode der sukzessiven Approximation der numerischen Lösung von Gleichungen höheren Grades. Dieses Ziel lassen wir aber hier beiseite; wir behandeln nur die Manipulation selbst.

Aufgabe Wir wollen den Wert des Polynoms

$$p(x) = a_4 x^4 + a_3 x^3 + a_2 x^2 + a_1 x + a_0$$

im Punkte x_0 berechnen.

1) Jemand, der nicht weiter nachdenkt, würde vielleicht folgendermaßen vorgehen:

Berechne zunächst $x_0 \cdot x_0$, dann $(x_0 \cdot x_0) \cdot x_0$, dann $(x_0 \cdot x_0 \cdot x_0) \cdot x_0$. Berechne sodann $a_1 x_0, a_2 x_0^2, a_3 x_0^3, a_4 x_0^4$. Addiere sodann

$$a_4 x_0^4 + a_3 x_0^3 + a_2 x_0^2 + a_1 x_0 + a_0 = p(x_0) .$$

Es sind 3+4 Multiplikationen erfolgt und 4 Additionen.

2) Mit weniger Aufwand kommt man aus, wenn man bedenkt

$$p(x) = \{[(a_4x + a_3)x + a_2]x + a_1\} + a_0 .$$

Man braucht nur 4 Multiplikationen und 4 Additionen.

3) Wir wollen nun etwas mehr. Wir suchen die Darstellung

$$p(x) = b_4(x - x_0)^4 + b_3(x - x_0)^3 + b_2(x - x_0)^2 + b_1(x - x_0) + b_0 ,$$

also nicht nur den Wert $b_0 = p(x_0)$. (Nebenbei bemerkt: Die Zahl $\frac{1}{k!} b_k$ ist die k -te Ableitung von $p(\cdot)$ im Punkte x_0 .)

$$\begin{aligned} p(x) &= \left\{ \left[(a_4x + a_3^{(0)})x + a_2^{(0)} \right] x + a_1^{(0)} \right\} x + a_0^{(0)} \\ &= \left\{ \left[(a_4x + a_3^{(1)})x + a_2^{(1)} \right] x + a_1^{(1)} \right\} (x - x_0) + b_0 \end{aligned}$$

mit

$$\begin{aligned} a_3^{(1)} &= a_4x_0 + a_3^{(0)} , & a_2^{(1)} &= a_3^{(1)}x_0 + a_2^{(0)} , \\ a_1^{(1)} &= a_2^{(1)}x_0 + a_1^{(0)} , & b_0 &= a_1^{(1)}x_0 + a_0^{(0)} . \end{aligned}$$

Offenbar gilt $b_0 = p(x_0)$.

Diese Prozedur setzen wir fort

$$\begin{aligned} p(x) &= \left\{ \left[(a_4x + a_3^{(2)})x + a_2^{(2)} \right] (x - x_0) + b_1 \right\} (x - x_0) + b_0 \\ &= \left\{ \left[(a_4x + a_3^{(3)}) (x - x_0) + b_2 \right] (x - x_0) + b_1 \right\} (x - x_0) + b_0 \\ &= \{ [(a_4(x - x_0) + b_3) (x - x_0) + b_2] (x - x_0) + b_1 \} (x - x_0) + b_0 \\ &= b_4(x - x_0)^4 + b_3(x - x_0)^3 + b_2(x - x_0)^2 + b_1(x - x_0) + b_0 \end{aligned}$$

Die Koeffizienten schreibt man in ein Schema, das nach WILLIAM G. HORNER (1756–1837) benannt ist.

$a_4^{(0)}$	$a_3^{(0)}$	$a_2^{(0)}$	$a_1^{(0)}$	$a_0^{(0)}$
$a_4^{(1)}$	$a_3^{(1)}$	$a_2^{(1)}$	$a_1^{(1)}$	$a_0^{(1)}$
$a_4^{(2)}$	$a_3^{(2)}$	$a_2^{(2)}$	$a_1^{(2)}$	
$a_4^{(3)}$	$a_3^{(3)}$	$a_2^{(3)}$		
$a_4^{(4)}$	$a_3^{(4)}$			

Am linken Ende sind alle Koeffizienten gleich. Die übrigen Einträge ergeben sich dadurch, daß man zum darüberliegenden den mit x_0 multiplizierten links daneben liegenden dazuaddiert

$$a_k^{(j+1)} = a_k^{(j)} + a_{k+1}^{(j)} x_0 .$$

Das Hornersche Schema beweist den

Satz Sei $p(x)$ ein Polynom vom Grade n mit $p(x_0) = 0$. Dann existiert ein Polynom $q(x)$ vom Grade $n - 1$, so daß

$$p(x) = (x - x_0)q(x) .$$

Satz (Fundamentalsatz der Algebra)

Zu jedem Polynom mit komplexen Koeffizienten vom Grade n gibt es komplexe Zahlen $z_1, \dots, z_n$, so daß

$$a_0 + a_1 z + \dots + a_n z^n = (z - z_1)(z - z_2) \dots (z - z_n) a_n .$$

Die Wurzeln $z_1, \dots, z_n$ sind bis auf ihre Reihenfolge eindeutig bestimmt. (Sie sind nicht notwendig verschieden.)

Beweis Der wesentliche Schritt ist der Nachweis, daß jedes Polynom vom Grade $n \geq 1$ mindestens eine Nullstelle besitzt. Das ist der berühmte Fundamentalsatz der Algebra, den wir hier nicht beweisen können. Vollständige Induktion nach n mit dem obigen Satz ergibt den Beweis.

Korollar Wenn die polynomialen Funktionen vom Grade $\leq n$, $p(z)$ und $q(z)$ in mindestens $n + 1$ Punkten übereinstimmen, dann sind sie identisch.

Gebrochenrationale Funktionen

Definition Eine gebrochenrationale Funktion ist eine Funktion von der Gestalt $\frac{p(z)}{q(z)}$ wo $p(z)$ und $q(z)$ Polynome sind, $q(z) \not\equiv 0$. (Man kann annehmen, daß $p(z)$ und $q(z)$ keine gemeinsamen Wurzeln haben: in den Wurzeln des Nennerpolynoms setzt man den Funktionswert $= \infty$.)

Satz („Teilen mit Rest“)

Sei $q(z)$ ein Polynom vom Grade $m \geq 1$. Zu jedem Polynom $p(z)$ vom Grad n ($n \geq m$) existieren ein eindeutig bestimmtes Polynom $P(z)$ vom Grade $n - m$ und ein $r(z)$ vom Grade $\leq m - 1$, so daß

$$\frac{p(z)}{q(z)} = P(z) + \frac{r(z)}{q(z)} .$$

Beweis (durch Induktion nach n)

$$p(z) - q(z) \frac{a_n}{b_m} z^{n-m}$$

ist vom Grade $\leq n-1$, also von der Form

$$\begin{aligned} &= \tilde{P}(z)q(z) + r(z) \\ p(z) &= \left(\frac{a_n}{b_m} z^{n-m} + \tilde{P}(z) \right) q(z) + r(z) . \end{aligned}$$

Das folgende Beispiel mit $q(z) = z^2 + 1$ zeigt, warum das Verfahren zur Auffindung von $P(\cdot)$ und $r(\cdot)$ „Teilen mit Rest“ genannt wird

$$\begin{array}{r} a_4 \quad a_3 \quad a_2 \quad a_1 \quad a_0 \\ a_4 \quad 0 \quad a_4 \\ \hline 0 \quad a_3 \quad a_2 - a_4 \\ \quad a_3 \quad 0 \quad a_3 \\ \hline 0 \quad a_2 - a_4 \quad a_1 - a_3 \\ \quad a_2 - a_4 \quad 0 \quad a_2 - a_4 \\ \hline 0 \quad a_1 - a_3 \quad a_0 - (a_2 - a_4) \\ \hline a_4 x^4 + a_3 x^3 + a_2 x^2 + a_1 x + a_0 \\ x^2 + 1 \end{array} \quad = \quad \begin{aligned} &(x^2 + 0x + 1)(a_4 x^2 + a_3 x + (a_2 - a_4)) + \\ &+ (a_1 - a_3)x + (a_0 - (a_2 - a_4)) \\ &= a_4 x^2 + a_3 x + (a_2 - a_4) + \\ &+ \frac{(a_1 - a_3)x + (a_0 - a_2 + a_4)}{x^2 + 1} \end{aligned}$$

Wir diskutieren nun die gebrochenrationalen Funktionen $\frac{p(z)}{q(z)}$, in welchem der Grad des Nenners größer ist als der Grad des Zählers.

Die einfachsten solchen Funktionen sind die Potenzen der „Stammbrüche“ $\left(\frac{1}{z-z^*}\right)$ und die Linearkombinationen solcher Potenzen. Wir werden sehen, daß es keine weiteren linear-gebrochenen Funktionen gibt. („Prinzip der Partialbruchzerlegung“)

Satz Sei $q(z) = (z - z_1)(z - z_2) \cdots (z - z_n)$ mit paarweise verschiedenen z_i .

Für jedes $p(z)$ vom Grad $< m$ existieren dann Zahlen $\alpha_1, \dots, \alpha_m$, so daß

$$\frac{p(z)}{q(z)} = \frac{\alpha_1}{z - z_1} + \frac{\alpha_2}{z - z_2} + \dots + \frac{\alpha_m}{z - z_m} .$$

Wir werden gleich den allgemeinen Satz beweisen, der auch mehrfache Wurzeln im Nennerpolynom zuläßt.

Satz („Partialbruchzerlegung“)

Sei $q(z)$ vom Grad m mit den paarweise verschiedenen Wurzeln $z_1, \dots, z_k$

$$q(z) = (z - z_1)^{m_1} \cdot (z - z_2)^{m_2} \cdots (z - z_k)^{m_k}, \quad m_1 + m_2 + \dots + m_k = m .$$

Für jedes Polynom $p(z)$ vom Grad $< m$ ist dann $\frac{p(z)}{q(z)}$ eine Linearkombination der „Stammbrüche“

$$\frac{1}{(z - z_i)^{\ell_i}} \quad \text{mit } i = 1, 2, \dots, k ; \quad 1 \leq \ell_i \leq m_i .$$

Die Koeffizienten sind eindeutig bestimmt.

Beispiele

$$\begin{aligned}\frac{2}{z^2-1} &= \frac{1}{z-1} - \frac{1}{z+1}; \\ \frac{3z-1}{z^2(z^2-1)} &= \frac{1}{z^2} - \frac{3}{z} + \frac{1}{z-1} + \frac{2}{z+1}\end{aligned}$$

Beweis des Satzes

1) Existenz durch Induktion nach m :

Setze $q(z) = (z - z_1)^{m_1} \cdot q_1(z)$. Dann gilt

$$\frac{p(z)}{q(z)} - \frac{p(z_1)}{q_1(z_1) \cdot (z - z_1)^{m_1}} = \frac{1}{q_1(z_1)} \cdot \frac{p(z)q_1(z_1) - p(z_1)q_1(z)}{(z - z_1)^{m_1}q_1(z)}.$$

Der Zähler verschwindet für $z = z_1$; man kann also mindestens den Faktor $z - z_1$ herauskürzen. Das Problem stellt sich nun für $\tilde{q}(z) = (z - z_1)^{m-1}q_1(z)$.

2) Die Eindeutigkeit der Darstellung ist leicht zu sehen.

Der reelle Fall Sei $q(x)$ ein Polynom mit reellen Koeffizienten. Die reellen Wurzeln seien $x_1, \dots, x_k$. Wenn z^* eine nichtreelle Wurzel ist, dann ist auch $\bar{z}^*$ eine Wurzel $\bar{z}^* \neq z^*$

$(x - z^*)(x - \bar{z}^*)$ hat reelle Koeffizienten .

$q(z)$ ist also ein Produkt von linearen und quadratischen Polynomen mit reellen Koeffizienten.

Beispiel $x^4 + 1$ hat keine reellen Wurzeln. Die Wurzeln sind:

$$\frac{1}{\sqrt{2}}(1+i), \quad \frac{1}{\sqrt{2}}(1-i), \quad \frac{1}{\sqrt{2}}(-1+i), \quad \frac{1}{\sqrt{2}}(-1-i).$$

$$\begin{aligned}\left(x - \frac{1}{\sqrt{2}}(1+i)\right)\left(x - \frac{1}{\sqrt{2}}(1-i)\right) &= x - \sqrt{2}x + 1 \\ \left(x - \frac{1}{\sqrt{2}}(-1+i)\right)\left(x - \frac{1}{\sqrt{2}}(-1-i)\right) &= x + \sqrt{2}x + 1 \\ x^4 + 1 &= (x^2 - \sqrt{2}x + 1)(x^2 + \sqrt{2}x + 1).\end{aligned}$$

Man prüft leicht nach

$$\frac{1}{x^4+1} = \frac{1}{2 \cdot \sqrt{2}} \frac{x + \sqrt{2}}{x^2 + \sqrt{2}x + 1} - \frac{1}{2 \cdot \sqrt{2}} \frac{x - \sqrt{2}}{x^2 - \sqrt{2}x + 1}$$

Dies ist die typische Situation, wenn man die Partialbruchzerlegung der reellen rational-gebrochenen Funktion $\frac{p(x)}{q(x)}$ durchführt. Im Zähler treten lineare Funktionen auf.

Bemerke Sei $x^2 + cx + d = (x - z^*)(x - \bar{z}^*)$, $z^* = x^* + iy^*$ mit $y^* \neq 0$

$$q(x) = (x^2 + cx + d)^\ell \tilde{q}(x) \quad \text{mit} \quad \tilde{q}(z^*) \neq 0 .$$

Wähle a, b , so daß

$$ay^* = \operatorname{Im} \left(\frac{p(z^*)}{\tilde{q}(z^*)} \right) , \quad ax^* + b = \operatorname{Re} \left(\frac{p(z^*)}{\tilde{q}(z^*)} \right) .$$

Dann gilt

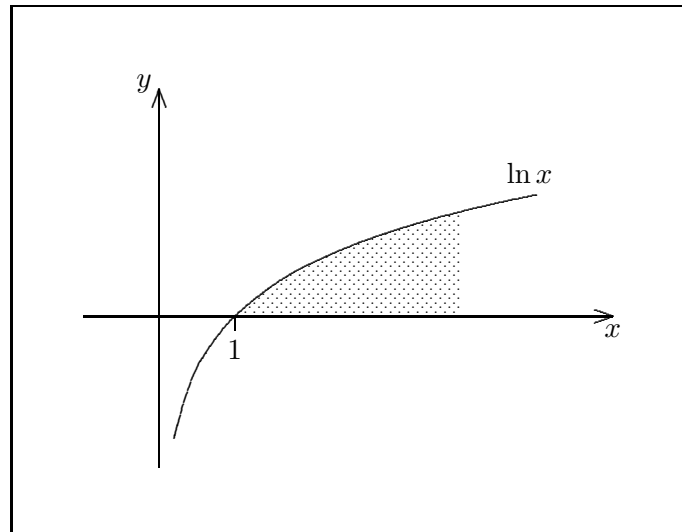
$$\frac{p(x)}{q(x)} - \frac{ax + b}{(x^2 + cx + d)^\ell} = \frac{p(x) - \tilde{q}(x)(ax + b)}{\tilde{q}(x)(x^2 + cx + d)^\ell}$$

und der Zähler verschwindet in z^* und in $\bar{z}^*$.

A.4 Ansätze zur Integrationstheorie, Stammfunktionen

Das Ziel dieses Abschnitts ist es, Formeln der folgenden Art herzuleiten

$$\int_1^x \ln u \, du = x \cdot \ln x - x + 1 \quad \text{für alle } x \geq 0$$



Im Eindimensionalen ist $\int f$ die „Fläche unter der Kurve“: im Zweidimensionalen benutzt man dasselbe Symbol $\int f$ für das „Volumen unter der Fläche“. Dabei denkt man zunächst nur an nichtnegative Funktionen f . $\int f$ ist eine Zahl oder auch $+\infty$.

Definition Eine Funktion f heißt *integrabel*, wenn $\int |f| < \infty$. Für integrable Funktionen, die auch negative Werte annehmen können, definiert man

$$\int f = \int f^+ - \int f^- .$$

Satz (Linearität)

Sind f und g integrabel, so gilt für alle $\alpha, \beta \in \mathbb{R}$

$$\int (\alpha f + \beta g) = \alpha \cdot \int f + \beta \cdot \int g .$$

Satz (Monotone Stetigkeit)

Für jede aufsteigende Folge integrierbarer Funktionen

$$f_1 \leq f_2 \leq \dots$$

gilt

$$\int (\lim \uparrow f_n) = \lim \uparrow \int f_n .$$

Diese Sätze können wir hier natürlich noch nicht beweisen. Der Beweis erfordert einen abstrakten Ansatz, der im zweiten Semester entwickelt wird.

Eine wichtige Konsequenz dieser Hauptsätze der Integrationstheorie ist das

Lemma von Fatou Für jede Folge nichtnegativen Funktionen $f_n \geq 0$ gilt

$$\int (\liminf f_n) \leq \liminf \int f_n .$$

Herleitung aus dem Satz von der monotonen Stetigkeit.

Betrachte für $m = 1, 2, \dots$

$$g_m = \inf \{f_n : n \geq m\} .$$

Offenbar gilt

$$\int g_m \leq \inf_{n \geq m} \int f_n .$$

Weiterhin gilt

$$g_1 \leq g_2 \leq \dots \lim_{\uparrow} g_m = \liminf f_n$$

$$\int (\liminf f_n) = \lim_{\uparrow} \int g_m \leq \lim_{\uparrow} \inf_{n \geq m} \int f_n$$

■

Hinweis Die beiden Hauptsätze (und viele Folgerungen daraus) sind Errungenschaften des 20. Jahrhunderts. Wir nehmen sie hier, wenn immer es nötig ist, als Tatsachen hin, sozusagen als Axiome. Die Autoren des 18. Jahrhunderts waren mit sehr speziellen Funktionen befaßt; sie wollten das Integral konkret berechnen. Fragen nach einem allgemeinen Integralbegriff stellten sich ihnen nicht. Es muß aber dem heutigen Studenten nicht verschwiegen werden, daß die intuitive Vorstellung von der „Fläche unter der Kurve“ im 20. Jahrhundert ein mathematisch strenges Fundament erhalten hat; den Durchbruch brachte 1901 die Dissertation von H. LEBESGUE (1875–1941).

Zur Geschichte

Die Bestimmung von Flächeninhalten und Volumina gehört zu den älteren Aufgaben der Mathematik. Zu den Kenntnissen und Bemühungen der alten Ägypter siehe Struik S. 17. Für die alten Griechen standen praktische Aufgaben nicht im Zentrum der Mathematik. Dennoch kann man die Flächen- und Volumenberechnungen von ARCHIMEDES zu den größten Leistungen der griechischen Mathematik rechnen (vgl. Walter S. 188). Ein großes Thema wurden solche Berechnungen dann in der Renaissance. Bahnbrechendes leisteten KEPLER

(1571–1630), CAVALIERI (1598–1647) und TORRICELLI (1608–1647). Hier, wie schon bei STEVIN (1548–1620), der sowohl über Schwerpunkte als auch über Hydraulik schrieb, standen praktische Aufgaben im Vordergrund. Insbesondere wurden Optimierungsaufgaben angegangen (siehe Walter, S. 222). Kepler sagte, daß die Beweise von Archimedes absolut streng seien, daß er sie aber den Leuten überlasse, die durchaus exakten Beweisen frönen wollten. Jeder folgende Autor nahm sich nun die Freiheit, seine eigene Art der Strenge festzulegen oder auch ganz darauf zu verzichten (Struik S. 109). Cavalieri stützte seine Methode auf den scholastischen Begriff der „Indivisiblen“, wonach durch die Bewegung eines Punktes eine Gerade und durch die Bewegung einer Geraden eine Ebene entsteht (Struik S. 110).

Dies war eine deutliche Abkehr von der Mathematik der Griechen. Bei den Griechen hatte es nur selten eine quantitative Beschreibung eines Bewegungsvorgangs gegeben; das Bewegte, der Veränderung Unterworfenen war das Unvollkommene, in Unordnung Befindliche (siehe Walter S. 221 ff). In der Renaissance bekamen die Maschinen große Bedeutung, und in der Zeit von GALILEI (1564–1642) rückte die Mechanik immer näher an die Mathematik heran. In seinen „Discorsi“ (1638) entwickelte Galilei das mathematische Studium der Bewegung sowie die Beziehung zwischen Weg, Geschwindigkeit und Beschleunigung. (Einzelheiten siehe Edwards)

Der „Kalkül“ von Leibniz und Newton brachte dann schließlich eine wirklich wirksame Verbindung zwischen der Flächenmessung (Integration), der Optimierung und der Bewegungsbeschreibung (Differentiation). Leibniz und Newton dynamisierten sozusagen das Problem der Flächenberechnung. Sie betrachteten die Fläche unter einer Kurve in Abhängigkeit von der oberen Grenze.

Die Lösung des eingangs gestellten Problems ergibt sich im Leibnizschen Kalkül aus der Feststellung, daß der infinitesimale Zuwachs der Funktion

$$x \ln x - x + 1$$

für alle $x > 0$ gleich dem infinitesimalen Zuwachs $\ln x \cdot dx$ der Fläche $\int_1^x \ln u \cdot du$ ist. Aus $d(\ln x) = \frac{1}{x} dx$ ergibt sich mit der Produktregel (siehe unten!)

$$\begin{aligned} d(x \ln x - x + 1) &= d(x \cdot \ln x) - dx \\ &= x \cdot d(\ln x) + \ln x \cdot dx - dx = \ln x \cdot dx \end{aligned}$$

Springen wir ins 20. Jahrhundert! Dort wurde der Integralbegriff nochmals durch eine Dynamisierung gestärkt und erweitert. In der Integrationstheorie von LEBESGUE (1875–1941) bleibt man nicht dabei stehen, einzelnen nichtnegativen Funktionen $f(\cdot)$ die Fläche unter der Kurve $\int f$ zuzuordnen. Man interessiert sich vielmehr für die Eigenschaften des Funktional

$$f \longmapsto \int f .$$

Der von Riemann entwickelte Integralbegriff erscheint demgegenüber als eine Methode für ziemlich allgemeine Funktionen $f(\cdot)$, die durch eine Fourier-Reihe dargestellt werden sollten,

das Integral durch einen Einschließungsprozeß direkt zu bestimmen. Daß dieser Einschließungsprozeß für manche Funktionen nicht zum Ziel führt, bedeutet nicht, daß man diesen Funktionen kein Integral zuordnen könnte. Den „natürlichen“ Definitionsbereich des Integral können wir aber hier nicht diskutieren.

Zur Didaktik

Die heute üblichen Anfängervorlesungen gehen davon aus, daß die Studenten erst dann für das Abstraktionsniveau der reifen, einfachen und geschlossenen Lebesgueschen Integrations-theorie bereit sind, wenn sie es müde sind, konkrete Integrale auszuwerten. Das betrachten wir deshalb als ein didaktisches Unglück, weil die Studenten auf diese Weise mit vielen überflüssigen Klimmzügen belastet und von den wichtigeren Fragestellungen abgelenkt werden. Wir sind der Meinung, daß ein Student, der den allgemeinen Mengenbegriff und den allgemeinen Funktionsbegriff akzeptieren kann mit der gleichen logischen Beherrschung sehr wohl auch den allgemeinen Integralbegriff akzeptieren kann. Das Konstruieren, Beweisen und Hinterfragen kann man u.E. auf später verschieben. Wir glauben, daß man die Intuition der Studenten dadurch am besten in die richtigen Bahnen lenkt, daß man ihnen ohne Beweise den Kern der reifen Integrationstheorie verrät. Damit ist dann auch ein Zielpunkt genannt für den systematischen Aufbau, den wir im „B-Zug“ begonnen haben und im nächsten Semester fortsetzen wollen.

Zurück ins 17. Jahrhundert (nach Struik S. 126)

Leibniz fand seinen „Kalkül“ zwischen 1673 und 1676 in Paris unter dem persönlichen Einfluß von Huygens und durch das Studium von Descartes und Pascal. Er wurde dazu durch die Nachricht angeregt, daß Newton im Besitz einer solchen Methode sein sollte. Während Newtons Vorgehen in erster Linie kinematisch orientiert war, war das von Leibniz geometrischer Natur. Die erste Veröffentlichung von Leibniz' Form der Differential- und Integralrechnung geschah 1684 in einem sechs Seiten langen Artikel in den „*Acta Eruditorum*“, einer mathematischen Zeitschrift, die er 1682 gegründet hatte. Sie enthielt unsere Symbole dx , dy und die Differentiationsregeln einschließlich der „Produktregel“

$$d(uv) = u dv + v du .$$

Dieser Arbeit folgte 1686 eine weitere mit den Regeln der Integralrechnung, die das Symbol $\int$ enthielt. Leibniz war einer der größten Erfinder mathematischer Symbole. Wenige Menschen haben die Einheit von Form und Inhalt so gut wie er verstanden. Seine Erfindung des „Kalküls“ muß auf diesem philosophischen Hintergrund gesehen werden; sie war das Ergebnis seines Forschens nach einer „lingua universalis“ der Veränderung und insbesondere der Bewegung.

Leibniz zwar ebenso wenig wie Newton, der seine Theorie der Fluxionen 1704/36 veröffentlichte, in der Lage, die Grundlagen seines Kalküls so zu erklären, daß die Erklärungen die Zustimmung seiner Zeitgenossen gefunden hätten. Mangels strenger Beweise gab er Analogien an, die das Unendliche und das Unendlichkleine erläuterten. In einem seiner Briefe (an

Foucher 1693) nahm er die Existenz des aktuellen Unendlichen an, um die Schwierigkeiten von Zeno zu überwinden und lobte Grégoire de Saint Vincent, der die Stelle berechnet hatte, wo Achilles die Schildkröte trifft. (Struik S. 128)

Die philosophischen Probleme des Unendlichen brauchen uns heute nicht mehr zu bekümmern. Es hat sich herausgestellt, daß man das Wesentliche des „Kalküls“ auf den Begriff der Stammfunktion und zwei Rechenregeln reduzieren kann. Es wird sich zeigen, daß der Begriff der Stammfunktion stimmig und sehr allgemein verwendbar ist. Die für den „Kalkül“ entscheidenden Regeln, die „Produktregel“ und die „Kettenregel“ werden wir an geeigneter Stelle präzisieren und beweisen. Hier geht es darum, den Sachverhalt (in moderner Notation) darzustellen und dann die Regeln auf einfache Fälle anzuwenden.

Definition Sei $f(\cdot)$ eine Funktion auf (a, b) , welche auf jedem $[c, d] \subseteq (a, b)$ integrierbar ist.

Man sagt von einer Funktion $F(\cdot)$, sie sei *Stammfunktion* für $f(\cdot)$, wenn

$$F(d) - F(c) = \int_c^d f(u) du \quad \text{für alle } c, d.$$

Beachte Wenn $F(\cdot)$ Stammfunktion ist, dann auch $F(\cdot) + \text{const.}$ Die Stammfunktion ist bis auf eine additive Konstante eindeutig bestimmt; denn

$$\begin{aligned} G(d) - G(c) &= \int_c^d f(u) du = F(d) - F(c) \\ \implies G(d) - F(d) &= G(c) - F(c) = \text{const.} \end{aligned}$$

Notation Wenn $F(\cdot)$ Stammfunktion für $f(\cdot)$ auf (a, b) ist, dann notieren wir

$$\begin{aligned} dF &= f(x) dx && \text{für } x \in (a, b) \\ \text{oder} \quad F &= \int f(x) dx && \text{für } x \in (a, b) \end{aligned}$$

Satz (Hauptsatz der Differential- und Integralrechnung)

Wenn auf (a, b)

$$dF = f(x) dx \quad \text{mit einem stetigen } f(\cdot),$$

dann gilt gleichmäßig auf jedem $[c, d] \subseteq (a, b)$

$$\lim_{h \rightarrow 0} \frac{1}{h} [F(x+h) - F(x)] = f(x).$$

Beweis

$$\begin{aligned} \frac{1}{h} [F(x+h) - F(x)] &= \frac{1}{h} \int_x^{x+h} f(y) dy \leq \sup \{f(y) : y \text{ zwischen } x \text{ und } x+h\} \\ \text{ebenso} \quad &\geq \inf \{f(y) : y \text{ zwischen } x \text{ und } x+h\} \end{aligned}$$

Wir werden sehen, daß jede stetige Funktion auf einem kompakten Intervall gleichmäßig stetig ist. Daher

$$\forall [c, d] \subseteq (a, b) \quad \forall \varepsilon > 0 : (x \in [c, d], |h| < \delta) \longrightarrow \left| \frac{1}{h} [F(x+h) - F(x)] - f(x) \right| < \varepsilon$$

Didaktischer Hinweis Sei $F(d) - F(c) = \int_c^d f(y) dy$ für alle $[c, d] \subseteq (a, b)$ mit einem beliebigen integrierbaren $f(\cdot)$.

Nach einem allgemeinen Satz ist $f(\cdot)$ Lebesgue-fast überall eindeutig bestimmt. Man kann sich auch im allgemeinen Fall mit der Identifikation des Integranden befassen. In der Maßtheorie wird das gemacht (Stichwort: Satz von Radon–Nikodym). Man handelt sich aber völlig überflüssige Schwierigkeiten ein, wenn man mit den unzureichenden Mitteln der elementaren Analysis daran geht, das Limesverhalten (für $h \rightarrow 0$) der Funktionenschar $\frac{1}{h} [F(x+h) - F(x)]$ zu studieren. Wer unbedingt wissen will, wie man die Schwierigkeiten mit „relativ elementaren“ Mitteln meistern kann, der sei etwa auf Walter, Analysis 2, S. 342 verwiesen.

Beispiele Die Stammfunktionen für die folgenden elementaren Funktionen waren schon vor der Erfindung des „Kalküls“ bekannt (in anderer Notation natürlich):

- 1) $d(\ln x) = \frac{1}{x} dx$ für $x \in (0, \infty)$
- 2) $d(e^x) = e^x dx$ für $x \in (-\infty, +\infty)$
- 3) $d(\sin x) = \cos x \cdot dx$
 $d(\cos x) = -\sin x \cdot dx$ für $x \in (-\infty, +\infty)$
- 4) $d(\arcsin x) = \frac{1}{\sqrt{1-x^2}} dx$ für $x \in (-1, +1)$
- 5) $d(\arctan x) = \frac{1}{1+x^2} dx$ für $x \in (-\infty, +\infty)$
- 6) $d(x^\alpha) = \alpha \cdot x^{\alpha-1} dx$ für $x \in (0, \infty)$

(Der Exponent α durfte auch damals schon eine rationale Zahl oder eine negative Zahl sein.)

Die Formeln sind den heutigen Studenten aus der Schule bekannt. Ihre Herleitung kann daher hier unterbleiben; sie würde uns von unseren Zielen ablenken. Wir wollen die Formeln hier im Montagszug als Bausteine für die Berechnung weiterer Stammfunktionen benützen. Das liefert uns Anschauungsmaterial für die strenge Begründung der Infinitesimalrechnung im B-Zug.

Beispiele mit unstetigen Integranden

7) Die für $u \neq 0$ definierte Funktion

$$f(u) = \frac{1}{\sqrt{|u|}}$$

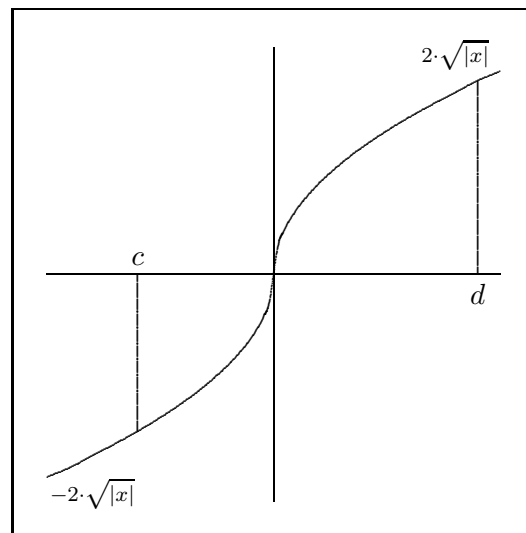
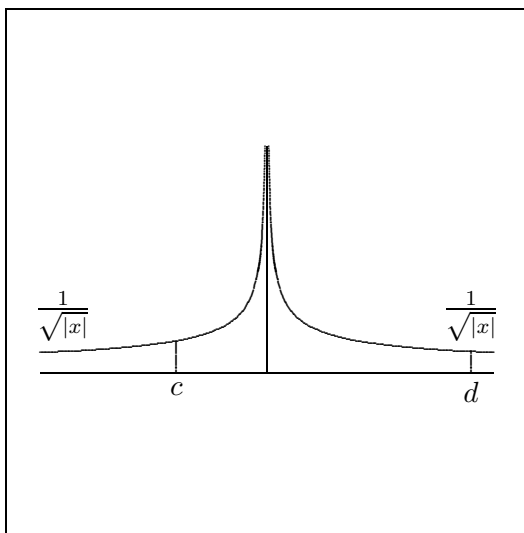
hat eine Stammfunktion auf $\mathbb{R}$, nämlich

$$F(x) = \begin{cases} 2 \cdot \sqrt{x} & \text{für } x \geq 0 \\ -2 \cdot \sqrt{|x|} & \text{für } x \leq 0. \end{cases}$$

Es gilt also

$$dF = \frac{1}{\sqrt{|x|}} dx \quad \text{für } x \in (-\infty, +\infty)$$

$$\int_c^d f(u) du = F(d) - F(c) \quad \text{für alle Paare } c, d.$$



$$8) f(x) = \text{sign}x = \begin{cases} +1 & \text{für } x > 0 \\ 0 & \text{für } x = 0 \\ -1 & \text{für } x < 0. \end{cases}$$

Für $F(x) = |x|$ gilt dann

$$dF = (\text{sign}x)dx \quad \text{für alle } x \in \mathbb{R}.$$

In der Tat gilt für alle $c < d$

$$|d| - |c| = F(d) - F(c) = \int_c^d (\text{sign}x) dx.$$

Produktregel :
$$d(F \cdot G) = F \cdot dG + G \cdot dF$$

Kettenregel :
$$d(H(F)) = h(F) \cdot dF, \text{ wenn } dH = h(y)dy .$$

Wir erläutern zunächst die Produktregel:

Nehmen wir an, F und G seien Stammfunktionen

$$dF = f(x)dx, \quad dG = g(x)dx \quad \text{für } x \in (a, b) .$$

Die Produktregel besagt dann, daß auch $F \cdot G$ eine Stammfunktion ist, und zwar

$$d(F \cdot G) = (F(x) \cdot g(x) + G(x) \cdot f(x))dx \quad \text{für } x \in (a, b) .$$

Die Erklärung kommt ohne den Begriff des Differentialquotienten aus. Die Produktregel gilt ganz allgemein. Für stetig differenzierbare F und G läßt sie sich ganz einfach aus dem obigen „Hauptsatz“ herleiten.

Wenn man die Produktregel ausführlich schreibt, dann heißt sie auch die „Regel der partiellen Integration“:

Für alle c, d mit $a < c < d < b$ gilt

$$F(d) \cdot G(d) - F(c) \cdot G(c) = \int_c^d F(x) \cdot g(x)dx + \int_c^d G(x) \cdot f(x)dx ,$$

was man üblicherweise so schreibt

$$\int_c^d g(x) \cdot F(x)dx = [F \cdot G]_c^d - \int_c^d G(x) \cdot f(x)dx .$$

Die „Methode der partiellen Integration“ geht so vor sich. Das Integral $\int_c^d h(x)dx$ ist zu bestimmen. Nehmen wir an, wir können $h(\cdot)$ als ein Produkt schreiben: $h = g \cdot F$, wo F eine Stammfunktion ist. Dann kann man das Problem zurückführen auf das Problem, die Stammfunktion von $G \cdot f$ zu finden. Diese Reduktion führt manchmal zum Ziel.

Beispiele für partielle Integration

1) Für $x > 0$ gilt

$$\int_1^x \ln u \, du = [u \ln u]_1^x - \int_1^x u \cdot \frac{1}{u} \, du = x \ln x - x + 1 .$$

Symbolisch geschrieben

$$d(x \ln x - x + 1) = \ln x \cdot dx \quad \text{für } x > 0 .$$

2) Für $\alpha > 0$ und alle $x > 0$ gilt

$$\begin{aligned}\int_0^x u^\alpha \cdot \sin u \, du &= -x^\alpha \cdot \cos x + \alpha \cdot \int_0^x u^{\alpha-1} \cdot \cos u \, du \\ \int_0^x u^\alpha \cdot \cos u \, du &= x^\alpha \cdot \sin x - \alpha \cdot \int_0^x u^{\alpha-1} \cdot \sin u \, du\end{aligned}$$

Symbolisch geschrieben

$$\begin{aligned}d(x^\alpha \cdot \cos x) &= [-x^\alpha \cdot \sin x + \alpha \cdot x^{\alpha-1} \cos x] \cdot x \quad \text{für } x > 0 \\ d(x^\alpha \cdot \sin x) &= [x^\alpha \cdot \cos x + \alpha \cdot x^{\alpha-1} \sin x] \cdot x \quad \text{für } x > 0.\end{aligned}$$

Die Funktionen [...] sind in der Tat integrierbar.

3 a) Die Funktion $\cos \frac{1}{x}$ ist zwar im Nullpunkt nicht stetig: sie ist aber beschränkt und daher über jedem endlichen Intervall integrierbar. Es gilt

$$\begin{aligned}\int_c^d \left[2u \cdot \sin \frac{1}{u} - \cos \frac{1}{u} \right] du &= \left[u^2 \cdot \sin \frac{1}{u} \right]_c^d \quad \text{für alle } c, d \\ d \left(x^2 \cdot \sin \frac{1}{x} \right) &= \left[2x \cdot \sin \frac{1}{x} - \cos \frac{1}{x} \right] dx \quad \text{für } x \in (-\infty, +\infty).\end{aligned}$$

3 b) Die Funktion $\frac{1}{x} \cos \frac{1}{x}$ ist in der Nähe des Nullpunkts nicht integrierbar. Es gilt aber immerhin für $x \in (0, \infty)$ und $x \in (-\infty, 0)$.

$$d \left(x^2 \cdot \sin \frac{1}{x} \right) = \left[\sin \frac{1}{x} - \frac{1}{x} \cos \frac{1}{x} \right] dx.$$

4) Für alle $\alpha > 0$ und $x \in (0, \infty)$ gilt

$$\begin{aligned}\Gamma(\alpha, x) &:= \int_0^x u^{\alpha-1} \cdot e^{-u} \, du \\ &= \left[\frac{1}{\alpha} u^\alpha \cdot e^{-u} \right]_0^x + \frac{1}{\alpha} \int_0^x u^\alpha \cdot e^{-u} \, du\end{aligned}$$

4*) Den Limes $\Gamma(\alpha) := \lim_{x \rightarrow \infty} \nearrow \Gamma(\alpha, x)$ nennt man die *Gammafunktion* für $\alpha > 0$. Die Funktion $\Gamma(\alpha, x)$ heißt manchmal die unvollständige Gammafunktion. Die Gammafunktion ist wegen der Funktionalgleichung

$$\Gamma(\alpha) = \frac{1}{\alpha} \cdot \Gamma(\alpha + 1) \quad (\text{hier für } \alpha > 0)$$

berühmt. Es gilt $\Gamma(n+1) = n!$ für $n = 0, 1, 2, \dots$

Eine sehr lesenswerte Beschreibung der Geschichte der Gammafunktion gibt Davis, The Chauvenet Papers, Herausgeber ABBOTT.

Wir werden später sehen

$$\Gamma\left(\frac{1}{2}\right) = \sqrt{\pi} = \int_0^{\infty} \frac{1}{\sqrt{x}} e^{-x} dx ,$$

eine weitere bemerkenswerte Beziehung zwischen der Exponentialfunktion und der Zahl π .

5) Für alle α und alle $\delta > 0$

$$d((\sin x)^{\alpha-1} \cdot \cos x) = [-\alpha \sin^{\alpha} x + (\alpha-1) \sin^{\alpha-2} x] dx \text{ für } x \in (\delta, \pi - \delta) .$$

Besonders interessant ist die Rekursion ($n = 2, 3, \dots$)

$$\int_0^x \sin^n u du = -\frac{1}{n} \sin^{n-1} x \cdot \cos x + \frac{n-1}{n} \int_0^x \sin^{n-2} u du .$$

5*) Setzen wir $x = \frac{\pi}{2}$, so erhalten wir

$$\begin{aligned} \int_0^{\pi/2} \sin^n x dx &= \frac{n-1}{n} \int_0^{\pi/2} \sin^{n-2} x dx \quad \text{für } n = 2, 3, \dots \\ \int_0^{\pi/2} \sin^{2m} x dx &= \frac{2m-1}{2m} \cdot \frac{2m-3}{2m-2} \cdot \dots \cdot \frac{1}{2} \cdot \frac{\pi}{2} \\ \int_0^{\pi/2} \sin^{2m+1} x dx &= \frac{2m}{2m+1} \cdot \frac{2m-2}{2m-1} \cdot \dots \cdot \frac{2}{3} \cdot 1 \\ \frac{\int_0^{\pi/2} \sin^{2m} x dx}{\int_0^{\pi/2} \sin^{2m+1} x dx} &= \frac{1}{2} \cdot \frac{2}{2} \cdot \frac{3}{4} \cdot \frac{5}{4} \cdot \frac{5}{6} \cdot \frac{7}{6} \cdot \dots \cdot \frac{2m-1}{2m} \cdot \frac{2m+1}{2m} \cdot \frac{\pi}{2} . \end{aligned}$$

Mit Hilfe von

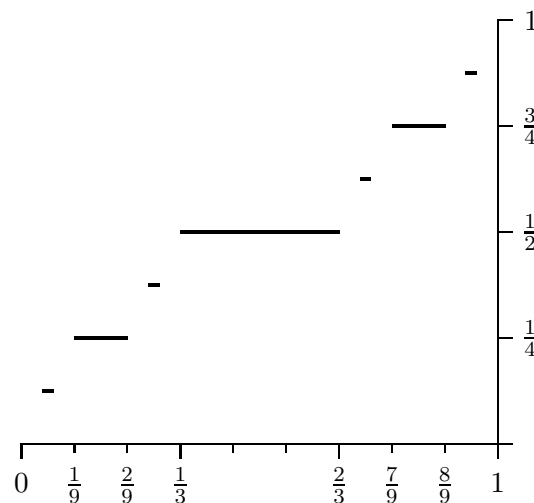
$$1 \leq \frac{\int_0^{\pi/2} \sin^{2m} x dx}{\int_0^{\pi/2} \sin^{2m+1} x dx} \leq 1 + \frac{1}{2m}$$

findet man die Wallissche Formel

$$\lim_{m \rightarrow \infty} \left(\frac{(2m)!}{(m!)^2} \cdot \frac{1}{2^{2m}} \cdot \sqrt{m} \right) = \sqrt{\pi} .$$

Hinweis für den Kenner

- 1) Die Funktionen $F(\cdot)$, die als Stammfunktionen auf einem Intervall auftreten können, heißen absolutstetige Funktionen. Jede absolutstetige Funktion läßt sich als Differenz von isotonen Funktionen darstellen. Aber nicht jede isotone stetige Funktion ist absolut stetig.
- 2) Ein berühmtes Beispiel ist Cantors Funktion



$$\begin{aligned}
 C(x) &= \frac{1}{2} && \text{für } x \in \left[\frac{1}{3}, \frac{2}{3}\right] \\
 C(x) &= \frac{1}{4} && \text{für } x \in \left[\frac{1}{9}, \frac{2}{9}\right] \\
 C(x) &= \frac{3}{4} && \text{für } x \in \left[\frac{7}{9}, \frac{8}{9}\right] \\
 C(x) &= \frac{1}{8} && \text{für } x \in \left[\frac{1}{27}, \frac{2}{27}\right] \\
 C(x) &= \frac{3}{8} && \text{für } x \in \left[\frac{4}{27}, \frac{5}{27}\right], \dots
 \end{aligned}$$

$C(\cdot)$ ist differenzierbar im Innern der Intervalle mit den Längen $\frac{1}{3}, \frac{1}{9}, \frac{1}{9}, \frac{1}{27}, \frac{1}{27}, \frac{1}{27}, \frac{1}{27}, \dots$

$$\frac{1}{3} + \frac{2}{9} + \frac{4}{27} + \frac{8}{81} + \dots = \frac{1}{3} \left(1 + \frac{2}{3} + \left(\frac{2}{3}\right)^2 + \dots \right) = \frac{1}{3} \cdot \frac{1}{1 - \frac{2}{3}} = 1.$$

Die Menge der x , in welchen $C(\cdot)$ nicht differenzierbar ist, ist eine Lebesguesche Nullmenge; sie heißt das Cantorsche Diskontinuum. Ein Punkt x gehört dann zum

Cantorsche Diskontinuum, wenn in der ternären Entwicklung

$$x = \sum_{i=1}^{\infty} \omega_i \cdot 3^{-i} \quad \omega_i = 0 \text{ oder } \omega_i = 2 \text{ für alle } i.$$

- 3) Die Funktionen beschränkter Schwankung (und zwar nicht nur die stetigen) erfreuen sich in der Maßtheorie seit langem großer Beliebtheit. Das Stichwort ist die Lebesgue-Stieltjes-Integration. Nehmen wir der Einfachheit halber an, daß $F(\cdot)$ isoton und rechtsstetig ist. Für alle borelmeßbaren nichtnegativen $h(\cdot)$ wird das „Integral“

$$\int h(x) dF(x)$$

definiert; und dieses Integral bzgl. dF hat die grundlegenden Eigenschaften der Linearität und monotonen Stetigkeit.

- 4) Die Funktionen unbeschränkter Schwankung treten erst seit neuerem als Elemente einer Integrationstheorie auf. Die stochastische Integration ist ein höchst lebendiger Zweig der Wahrscheinlichkeitstheorie. Besonders wichtig und beliebt ist die Integration bzgl. des Wiener-Maßes oder bzgl. des Prozesses der Brownschen Bewegung. Die zufälligen und deterministischen Objekte, die hier integriert werden

$$\int_0^t h(s) dW_s$$

sind aber nicht so einfach zu beschreiben. Die Produktregel und die Kettenregel wird in der stochastischen Integration durch eine etwas komplizierte Regel abgelöst („Itô-Formel“).

- 5) Von Differentialquotienten ist in keiner dieser weitergehenden Theorien die Rede. Die Fortführung des Begriffs des Differentialquotienten in der abstrakten Maßtheorie ist der Begriff der Radon-Nikodym-Ableitung.

A.5 Variablensubstitution. Spezielle Kurven.

Zu Newtons Ansatz.

Die Aussage „ $J(\cdot)$ ist Stammfunktion zu $j(\cdot)$ “ schreiben wir

$$dJ = j(x)dx \quad \text{oder auch} \quad J = \int j(x)dx .$$

Sie ist genau dann wahr auf dem Intervall (a, b) , wenn

$$J(d) - J(c) = \int_c^d j(x)dx \quad \text{für alle } [c, d] \subseteq (a, b) .$$

Nicht jede elementare Funktion $j(\cdot)$ besitzt eine elementare Stammfunktion $J(\cdot)$. Auch dann, wenn es zu $j(\cdot)$ eine elementare Stammfunktion gibt, ist diese nicht immer leicht zu finden. Geschick und Erfahrung ist erforderlich; die Produktregel (Regel von der partiellen Integration) und die Kettenregel (Regel von der Variablensubstitution) sind die wichtigsten Hilfsmittel. Die Kettenregel hilft weiter, wenn es gelingt, den Integranden $j(x)$ so zu faktorisieren, daß

$$j(x)dx = h(F(x)) \cdot dF(x) .$$

Wenn $H(\cdot)$ die Stammfunktion zu $h(\cdot)$ ist, dann hat man

$$\int_c^d j(x)dx = \int_c^d h(F(x)) \cdot f(x)dx = H(F(d)) - H(F(c)) .$$

Den genauen Geltungsbereich dieser Formeln wollen wir hier nicht zu bestimmen versuchen. Im Montagszug geht es um den Kalkül.

Beispiele

$$1) \int \frac{1}{1 + \sin^2 x} \cdot \cos x \, dx = \arctan(\sin x).$$

In der Tat gilt mit

$$y = F(x) = \sin x , \quad dy = f(x)dx = \cos x \, dx$$

$$\begin{aligned} \frac{1}{1 + \sin^2 x} \cdot \cos x \cdot dx &= \frac{1}{1 + F^2(x)} \cdot f(x) \, dx = \frac{1}{1 + y^2} \, dy \\ &= d(\arctan y) = d(\arctan F(x)) . \end{aligned}$$

$$2) \int e^{\sqrt{x}} \, dx = 2 \cdot (\sqrt{x} - 1) \cdot e^{\sqrt{x}} \quad \text{für } x \in (0, \infty)$$

Das ergibt sich mit der Substitutionsmethode wie folgt

$$y = \sqrt{x}, \quad dy = \frac{1}{2} \frac{1}{\sqrt{x}} dx \quad \text{oder} \quad x = y^2, \quad dx = 2y dy$$

$$\begin{aligned} \int e^{\sqrt{x}} dx &= \int e^y \cdot 2y dy = e^y \cdot 2y - \int e^y \cdot 2 \cdot dy \quad (\text{partielle Integration}) \\ &= e^y \cdot 2y - 2e^y = 2(y-1) \cdot e^y = 2 \cdot (\sqrt{x}-1)e^{\sqrt{x}}. \end{aligned}$$

$$3) \int \frac{1}{\sin x} dx = \ln \left(\tan \frac{x}{2} \right) \quad \text{für } x \in (0, \pi).$$

Daß die Funktion $j(x) = \frac{1}{\sin x}$ in $(0, \pi)$ eine Stammfunktion besitzt, ist offensichtlich; $j(\cdot)$ hat über jedem Intervall $[c, d] \subseteq (0, \pi)$ ein endliches Integral. Wer von irgendwo die Vermutung gewonnen hat, daß $J(x) = \ln \left(\tan \frac{x}{2} \right)$ diese (bis auf eine additive Konstante eindeutig bestimmte) Stammfunktion ist, prüft sie leicht nach

$$\begin{aligned} dJ &= \frac{1}{\tan \frac{x}{2}} \cdot d \left(\tan \frac{x}{2} \right) = \frac{1}{\tan \frac{x}{2}} \cdot \frac{1}{\cos^2 \left(\frac{x}{2} \right)} \cdot \frac{1}{2} dx \\ &= \frac{1}{2 \sin \frac{x}{2} \cdot \cos \frac{x}{2}} dx \end{aligned}$$

denn

$$\begin{aligned} d \left(\frac{\sin y}{\cos y} \right) &= \frac{1}{\cos y} \cdot \cos y dy + \sin y \cdot \frac{-1}{\cos^2 y} \cdot (-\sin y) dy \\ &= \frac{\cos^2 y + \sin^2 y}{\cos^2 y} dy = \frac{1}{\cos^2 y} dy \\ 2 \cdot \sin \frac{x}{2} \cdot \cos \frac{x}{2} &= \sin x \end{aligned}$$

Um zur Vermutung zu gelangen, braucht es die schlaue Idee, $y = \tan \frac{x}{2}$ als neue Integrationsvariable einzuführen. Diese Idee liefert

$$\begin{aligned} dy &= \frac{1}{\cos^2 \frac{x}{2}} \cdot \frac{1}{2} dx \\ \int \frac{1}{\sin x} dx &= \int \frac{2 \cdot \cos^2 \frac{x}{2}}{2 \cdot \sin \frac{x}{2} \cdot \cos \frac{x}{2}} \cdot \frac{1}{\cos^2 \frac{x}{2}} \cdot \frac{1}{2} dx \\ &= \int \frac{1}{y} dy = \ln y = \ln \left(\tan \frac{x}{2} \right) \end{aligned}$$

$$4) \int \frac{1}{\sqrt{1+x^2}} dx = \ln \left(x + \sqrt{1+x^2} \right) \quad \text{für } x \in (-\infty, +\infty)$$

Eine gute Chance, die Lösung zu vermuten, hat derjenige, der die sog. Hyperbelfunktionen kennt

$$\begin{aligned}\sinh t &:= \frac{1}{2}(e^t - e^{-t}) \\ \cosh t &:= \frac{1}{2}(e^t + e^{-t}) .\end{aligned}$$

Man prüft leicht nach

$$\begin{aligned}d(\sinh t) &= \cosh t \cdot dt \\ d(\cosh t) &= \sinh t \cdot dt \\ (\cosh t)^2 - (\sinh t)^2 &= 1 \quad \text{für } t \in (-\infty, +\infty) .\end{aligned}$$

Der hyperbolische Sinus besitzt eine Umkehrfunktion auf $\mathbb{R}$; denn

$$\begin{aligned}x = \frac{1}{2}(e^t - e^{-t}) &\iff 2x \cdot e^t = e^{2t} - 1 \iff e^{2t} - 2xe^t = 1 \\ &\iff (e^t - x)^2 = 1 + x^2 \iff e^t = x + \sqrt{1 + x^2} \\ &\iff t = \ln(x + \sqrt{1 + x^2}) .\end{aligned}$$

Für die neue Integrationsvariable t haben wir

$$\begin{aligned}x = \sinh t \quad dx &= (\cosh t)dt \\ \text{andererseits } \frac{1}{\sqrt{1+x^2}} &= \frac{1}{\cosh t} \\ \int \frac{1}{\sqrt{1+x^2}} dx &= \int \frac{1}{\cosh t} \cdot (\cosh t)dt = t = \ln(x + \sqrt{1+x^2}) .\end{aligned}$$

Die Integration der gebrochenrationalen Funktionen

- A) Wenn das quadratische Polynom $x^2 + cx + d$ keine reelle Nullstelle besitzt, d.h. wenn $d - \frac{c^2}{4} > 0$, dann gilt

$$\begin{aligned}\int \frac{1}{x^2 + cx + d} dx &= \frac{1}{\sqrt{d - \frac{c^2}{4}}} \cdot \arctan \left(\frac{x + \frac{c}{2}}{\sqrt{d - \frac{c^2}{4}}} \right) \\ \int \frac{2x + c}{x^2 + cx + d} dx &= \ln(x^2 + cx + d) .\end{aligned}$$

Durch eine affine Transformation können wir das Problem, $\frac{ax+b}{(x^2+cx+d)^n}$ zu integrieren auf die Probleme,

$$\frac{1}{(1+x^2)^n} \quad \text{und} \quad \frac{2x}{(1+x^2)^n}$$

zu integrieren, zurückführen. $n = 2, 3, \dots$

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

Betrachten wir für $n = 1, 2, \dots$

$$F_n(x) := \int \frac{1}{(1+x^2)^n} \cdot dx.$$

Es gilt die Rekursionsformel (für $n = 2, 3, \dots$)

$$F_n(x) = \frac{x}{2(n-1)(1+x^2)^{n-1}} + \frac{2n-3}{2(n-1)} F_{n-1}(x)$$

B) Für jedes $x^* \in \mathbb{R}$ gilt in $(-\infty, x^*)$ und in (x^*, ∞)

$$\int \frac{1}{(x-x^*)^n} dx = \begin{cases} \ln|x-x^*| & \text{falls } n=1 \\ \frac{1}{1-n} \frac{1}{(x-x^*)^{n-1}} & \text{falls } n=2, 3, \dots \end{cases}$$

C) Sei

$$f(x) = \frac{p(x)}{q(x)}$$

eine gebrochenrationale Funktion. Mit Hilfe der Partialbruchzerlegung finden wir in den Intervallen zwischen den Nullstellen des Nennerpolynoms die elementare Stammfunktion. Diese Stammfunktion ist im allgemeinen keine gebrochenrationale Funktion.

Beispiele

$$1) \int \frac{1}{x^2-1} dx = \frac{1}{2} \ln \left| \frac{x-1}{x+1} \right| \quad \text{für } x \in (-\infty, -1) \cup (-1, +1) \cup (+1, +\infty)$$

$$2) \quad 2 \cdot \sqrt{2} \int \frac{1}{1+x^4} dx = \frac{1}{2} \ln(x^2 + \sqrt{2} \cdot x + 1) - \frac{1}{2} \ln(x^2 - \sqrt{2} \cdot x + 1) \\ + \arctan(\sqrt{2} \cdot x + 1) + \arctan(\sqrt{2} \cdot x - 1).$$

Dem Buch von Courant (S. 204) entnehmen wir, daß Leibniz mit dem Beispiel Schwierigkeiten hatte. Offenbar hat er sich verrechnet; und das Prinzip der Partialbruchzerlegung stand ihm nicht so klar vor Augen, daß er aufgab, bevor er den Rechenfehler entdeckt hatte. Courants Buch behandelte einige weitere Klassen von elementar integrierbaren Funktionen, indem er sie mit Hilfe der Substitutionsregel auf rationalgebrochene Funktionen zurückführt.

Hinweis zum Begriff der elementaren Funktion Es hat sich im 19. Jahrhundert in der Lehrbuchliteratur eingebürgert, diejenigen Funktionen als elementare Funktionen zu bezeichnen, die sich aus den Konstanten und den Funktionen

$$x, \ln x, e^x, \sin x, \cos x, \arcsin x, \arctan x$$

mittels der Grundrechenoperationen, dem Hintereinanderschalten und der Bildung der Umkehrfunktion gewinnen lassen. Diese Namensgebung hat keinen tieferen Grund; Funktionen wie die Gammafunktion oder das „Fehlerintegral“

$$\Phi(x) := \int_{-\infty}^x \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}y^2} dy$$

verdienen das Etikett „elementar“ mindestens ebenso sehr wie irgendwelche komplizierte zusammengesetzte Funktion wie z.B. $\ln(\sin(e^x))$. Wie immer man den Begriff der elementaren Funktion fassen mag, wenn man verlangt, daß nicht nur die Ableitungen sondern auch die Stammfunktion von elementaren Funktionen elementar heißen sollen, dann wird man keine wirklich überschaubare Klasse von Funktionen erhalten. Die Stammfunktion von so einfachen Funktionen wie

$$\frac{1}{\sqrt{1+x^4}}, \quad \frac{1}{x} e^x, \quad \frac{1}{x} \sin x, \quad \frac{1}{\ln x}, \quad \frac{1}{\sqrt{\cos 2x}}$$

sind im traditionellen Sinne nicht elementar — was allerdings nicht ganz leicht zu beweisen ist. Manche dieser Stammfunktionen haben Anlaß zu interessanten mathematischen Überlegungen gegeben.

a) $\int \frac{dy}{\sqrt{\cos 2y}}$ wird durch die Variablensubstitution $x = \sin y$ zu

$$\int \frac{dx}{\sqrt{(1-x^2)(1-2x^2)}}$$

b) Die Stammfunktion

$$\int \frac{dy}{\sqrt{\cos \alpha - \cos y}}$$

wird durch die Variablensubstitution $x = \cos \frac{y}{2}$ zu

$$-k \cdot \sqrt{2} \int \frac{dx}{\sqrt{1-x^2} \sqrt{1-k^2 x^2}} \quad \text{mit} \quad k = \frac{1}{\cos \frac{\alpha}{2}}.$$

c) $\int \frac{dy}{\sqrt{1-k^2 \sin^2 y}}$ wird zu $\int \frac{dx}{\sqrt{(1-x^2)(1-k^2 x^2)}}$.

Bekanntlich gilt

$$\arcsin(s) = \int_0^s \frac{dx}{\sqrt{1-x^2}}.$$

Man kann den Sinus als die Umkehrfunktion dieses unbestimmten Integrals einführen (mit 2π -periodischer Fortsetzung).

- d) Man erhält höchst interessante Funktionen $s(u)$, die sogenannten elliptischen Funktionen, wenn man die Umkehrfunktion von

$$u(s) = \int_0^s \frac{dx}{\sqrt{(1-x^2)(1-k^2x^2)}}$$

studiert. Diese Funktionen $s(\cdot)$ sind im 19. Jahrhundert eingehend untersucht worden. Die Anwendungen gehören nicht zum traditionellen Stoff der Anfängervorlesung.

Die exzessive Beschäftigung mit den sogenannten elementaren Funktionen, die manche Lehrbücher für lehrreich halten, wollen wir hier nicht mitmachen.

Geschichtliches : Newtons Ansatz

Wie schon in A.4 erwähnt, geben die Notationen, die wir heute verwenden, im großen Ganzen auf Leibniz zurück. Leibniz konnte die Grundlagen des neuen Kalküls ebensowenig schlüssig erklären wie Newton. Seine Symbole dx, dy wurden manchmal wie endliche Größen behandelt und manchmal als Größen aufgefaßt, die kleiner als jede angebbare Zahl und doch nicht Null waren. Ähnlich war es mit den „Fluxionen“ bei Newton. Beide Väter des Kalküls wurden wegen der mangelnden Grundlagen heftig kritisiert, Newton insbesondere von BERKELEY (1685–1753), der die Fluxionen als die „Geister von abgeschiedenen Größen“ verspottete. Berkeley glaubte, daß die offenbar richtigen Resultate des Kalküls durch eine Kompensation von Fehlern erzielt würden. (siehe Struik, insbesondere S. 141 ff). Die Differentiale $dx, dy, \dots$ kamen im 19. Jahrhundert außer Mode; man arbeitete nur noch mit Differentialquotienten. (Genaueres siehe Struik S. 151). Erst im 20. Jahrhundert kamen sie wieder in der nun streng begründeten Theorie der Differentialformen.

Newton dachte nicht im gleichen Sinne wie Leibniz an Funktionen. Er dachte nicht an eine „unabhängige Variable“ x und eine „abhängige Variable“ y . Für ihn waren x und y gleichberechtigte „Fluents“. Das folgende Zitat (nach Struik S. 123) kann so verstanden werden, daß Newton an Bahnen dachte, auf denen man sich infinitesimal bewegen kann. In Newtons „Methode der Fluxionen“ (1736) heißt es: *Die Variablen für Fluents werden mit $v, x, y, z, \dots$ bezeichnet „und die Geschwindigkeiten, mit denen jede Fluente durch ihre Bewegung vergrößert wird (die ich Fluxionen oder einfach Geschwindigkeiten nennen will) werde ich durch dieselben (punktierten) Buchstaben bezeichnen, also $\dot{v}, \dot{x}, \dot{y}, \dot{z}$.“*

Die infinitesimalen Größen werden bei Newton „Momente der Fluxionen“ genannt und durch $\dot{x}o, \dot{y}o, \dot{z}o$ bezeichnet, wo o eine „unendlich kleine Größe“ ist. Newton fährt dann fort:

„So sei eine beliebige Gleichung

$$x^3 - ax^2 + axy - y^3 = 0$$

gegeben. Man setze darin $x + \dot{x}o$ für $x, y + \dot{y}o$ für y und erhält dann

$$\begin{aligned} x^3 + 3x^2\dot{x}o + 3x\dot{x}o\dot{x}o + \dot{x}^3o^3 - ax^2 - 2ax\dot{x}o - a\dot{x}o\dot{x}o + axy \\ + ay\dot{x}o + a\dot{x}o\dot{y}o + ax\dot{y}o - y^3 - 3y^2\dot{y}o - 3y\dot{y}o\dot{y}o - \dot{y}^3o^3 = 0. \end{aligned}$$

Nun gilt nach Voraussetzung $x^3 - ax^2 + axy - y^3 = 0$, so daß sich nach Streichung dieser Größen und Division des verbleibenden Ausdrucks durch o ergibt

$$\begin{aligned} 3x^2\dot{x} - 2ax\dot{x} + ay\dot{x} + ax\dot{y} + 3x\dot{x}\dot{x}o - a\dot{x}\dot{x}o \\ + a\dot{x}\dot{y}o - 3y\dot{y}\dot{y}o + \dot{x}^3oo - \dot{y}^3oo = 0 \end{aligned}$$

Da aber o als unendlich klein vorausgesetzt wird, so daß es Momente von Größen darstellen kann, ergeben die damit multiplizierten Ausdrücke im Vergleich zu den übrigen nichts. Ich lasse sie deshalb weg, und es bleibt übrig

$$3x^2\dot{x} - 2ax\dot{x} + ay\dot{x} + ax\dot{y} - 3y^2\dot{y} = 0. "$$

Die Division des verbleibenden Ausdrucks durch o und das Weglassen der Glieder mit einem mehrfachen o hat allerseits großes Unbehagen ausgelöst. (vgl. Struik S. 128 und S. 140). Ähnlichkeiten mit Newtons Notation haben sich aber dennoch in manchen Zusammenhängen bis heute gehalten; der Punkt (wie in $\dot{v}, \dot{x}, \dot{y}, \dot{z}$) wird heute häufig verwandt, wenn man an Ableitungen nach einem Zeitparameter denkt. Wenn man Newtons Kurve

$$x^3 - ax^2 + axy - y^3 = 0$$

durchläuft, dann stehen die infinitesimalen Änderungen von x und y in der Beziehung

$$(3x^2 - 2ax + ay)\dot{x} + (ax - 3y^2)\dot{y} = 0.$$

Newton redet aber nicht von einem Zeitparameter; wenn man die Zeit anders parametrisiert, dann multiplizieren sich die „Geschwindigkeiten“ $\dot{x}$ und $\dot{y}$ mit derselben Funktion von t .

Hinweis In der modernen Auffassung des 20. Jahrhunderts würde man das Resultat von Newtons Überlegung etwa folgendermaßen ausdrücken:

Auf der Mannigfaltigkeit

$$\{(x, y) : x^3 - ax^2 + axy - y^3 = 0\}$$

verschwindet die Differentialform

$$(3x^2 - 2ax + ay)dx + (ay - 3y^2)dy .$$

Für eine differenzierbare Kurve $(x(t), y(t))$, die auf der Mannigfaltigkeit verläuft, gilt für den Geschwindigkeitsvektor $(\dot{x}(t), \dot{y}(t))$ die Beziehung

$$(3x^2(t) - 2ax(t) + ay(t))\dot{x}(t) + (ay(t) - 3y^2(t))\dot{y}(t) = 0 .$$

Im Gegensatz zu Newtons Bezeichnung mit dem Punkt über der Variablen wird die Ableitung in Bezug auf den Ort seit der Zeit von Euler mit einem Strich bezeichnet. Das überraschende Ansehen der Lehrbücher von Euler (insbesondere die „Introductio ad analysin infinitorum“ von 1748) bereinigte überhaupt viele strittige Bezeichnungsfragen (vgl. Struik S. 136). Zum Gebrauch der Bezeichnungen

$$y', f'(x), \frac{dy}{dx} \text{ etc.}$$

siehe auch Walter, S. 238.

Kurven

Wir zitieren aus W. Walter, Analysis 2 (S. 151): „*Es gibt zwei Arten, den Begriff der Kurve zu bestimmen. In der geometrischen Auffassung ist eine Kurve der Ort von Punkten in der Ebene oder im Raum, die durch gewisse Eigenschaften charakterisiert sind. So wird etwa in der Ebene ein Kreis durch den konstanten Abstand zu einem Punkt und eine Ellipse durch die konstante Abstandssumme zu zwei Punkten beschrieben. Im mechanischen Bild erscheint die Kurve als Bahnkurve eines bewegten Punktes. Beide Auffassungen finden sich bereits bei den Griechen. Die Kegelschnitte, ein Hauptgegenstand der griechischen Mathematik, sind durch geometrische Eigenschaften definiert. Die erste mechanisch erklärte Kurve ist die Archimedische Spirale.*“

Beispiel für Kurven

- 1) **Ellipsen** x, y seien cartesische Koordinaten in der Ebene. Betrachten wir zwei Punkte im Abstand $2e$, etwa

$$(x, y)(P_-) = (-e, 0) ; (x, y)(P_+) = (+e, 0) .$$

Der geometrische Ort aller Punkte P , für welche die Summe der Abstände zu diesen Punkten gleich $2a$ ist, ist folgendermaßen zu beschreiben

$$2a = \sqrt{(x - e)^2 + y^2} + \sqrt{(x + e)^2 + y^2}$$

Die Gleichung kann man umformen, etwa durch Erweitern mit der Differenz der Wurzeln

$$2a \cdot \left[\sqrt{(x + e)^2 + y^2} - \sqrt{(x - e)^2 + y^2} \right] = 4ex .$$

Aus den beiden Gleichungen erhält man durch Linearkombination

$$4a \cdot \sqrt{(x+e)^2 + y^2} = 4a^2 + 4ex$$

und schließlich

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1 \quad \text{mit } b^2 = a^2 - e^2.$$

Als Parametrisierung der Ellipse bietet sich an

$$x = a \cdot \cos t, \quad y = b \cdot \sin t.$$

t heißt die exzentrische Anomalie; wenn t das Intervall $[0, 2\pi]$ durchläuft, dann durchläuft P_t mit den Koordinaten $(a \cdot \cos t, b \cdot \sin t)$ die Ellipse.

2) Die Funktionen

$$x = \sinh t, \quad y = \cosh t$$

heißen deswegen Hyperbelfunktionen, weil (x, y) den Hyperbelast

$\{(x, y) : y^2 - x^2 = 1, y > 0\}$ durchläuft, wenn t das Intervall $(-\infty, +\infty)$ durchläuft.

Hinweis Die Kurve $y = \cosh x$ heißt Kettenlinie. Zur geometrischen Begründung dieses Namens siehe Heuser, S. 296.

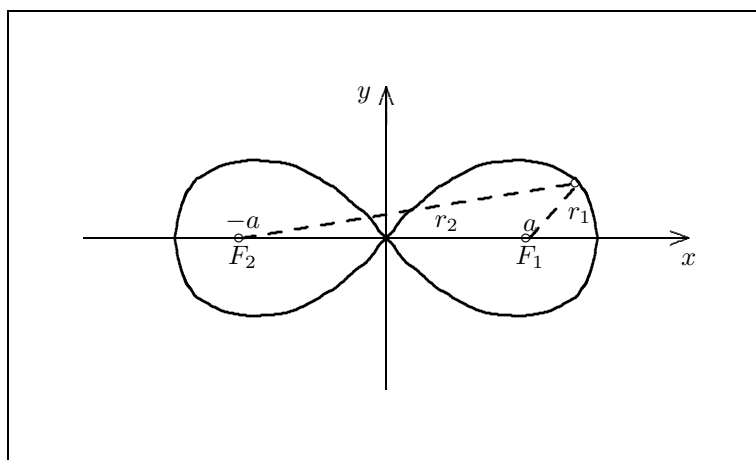
3) Die Lemniskate ist geometrisch definiert als der Ort aller Punkte P , für welche das Produkt der Abstände $r(P_-)$ und $r(P_+)$ von zwei festgehaltenen Punkten P_- bzw. P_+ im Abstand $2a$ den Wert a^2 hat. Für

$$(x, y)(P_-) = (-a, 0), \quad (x, y)(P_+) = (+a, 0)$$

$$r_-^2 = (x+a)^2 + y^2, \quad r_+^2 = (x-a)^2 + y^2$$

ergibt sich aus $r_-^2 \cdot r_+^2 = a^4$ durch eine einfache Umrechnung die Lemniskatengleichung

$$(x^2 + y^2)^2 - 2a^2(x^2 - y^2) = 0.$$



Führen wir Polarkoordinaten ein

$$x = r \cdot \cos \vartheta, \quad y = r \cdot \sin \vartheta,$$

so erhalten wir

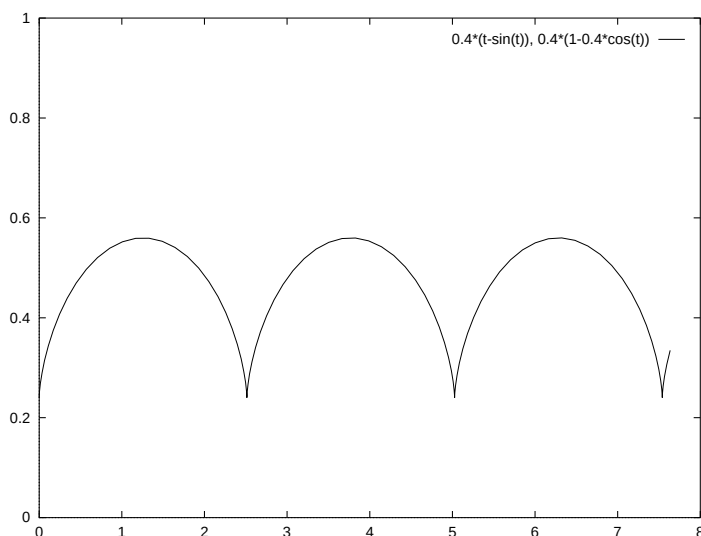
$$r^4 - 2a^2 \cdot r^2 (\cos^2 \vartheta - \sin^2 \vartheta) = 0$$

$$r^2 = 2a^2 \cdot \cos 2\vartheta.$$

- 4) **Zykloiden** entstehen, wenn ein Kreis auf einer Geraden oder auf einem anderen Kreis abrollt. Jeder Punkt der rollenden Kreisscheibe beschreibt dabei eine Zykloide. Betrachten wir den einfachsten Fall, daß ein Kreis vom Radius a auf der x -Achse rollt. Wählen wir den Anfangspunkt des Koordinatensystems und den der Zeitrechnung so, daß der Kurvenpunkt für $t = 0$ in den Ursprung fällt, dann ergibt sich für diese „gewöhnliche“ Zykloide die Parameterdarstellung

$$x = a(t - \sin t), \quad y = a(1 - \cos t);$$

dabei bedeutet t den Winkel, um den sich der rollende Kreis aus seiner Anfangslage herausgedreht hat und der bei gleichförmiger Rollbewegung der Zeit proportional ist.



Zykloide mit $a = 0,4$

Man kann nach Elimination des Parameters t die Gleichung der Zykloide auch in rechtwinkligen Koordinaten schreiben

$$\cos t = \frac{a - y}{a}, \quad t = \arccos \frac{a - y}{a}, \quad \sin t = \pm \sqrt{1 - \frac{(a - y)^2}{a^2}}$$

$$x = a \cdot \arccos \frac{a - y}{a} \pm \sqrt{(2a - y)y}.$$

Wir haben schon erwähnt, daß Leibniz 1686 das Symbol $\int$ eingeführt hat. Leibniz schrieb die Gleichung der Zykloide folgendermaßen

$$y = \sqrt{2x - x^2} + \int \frac{dx}{\sqrt{2x - x^2}}.$$

5) Die van der Waals Gleichung

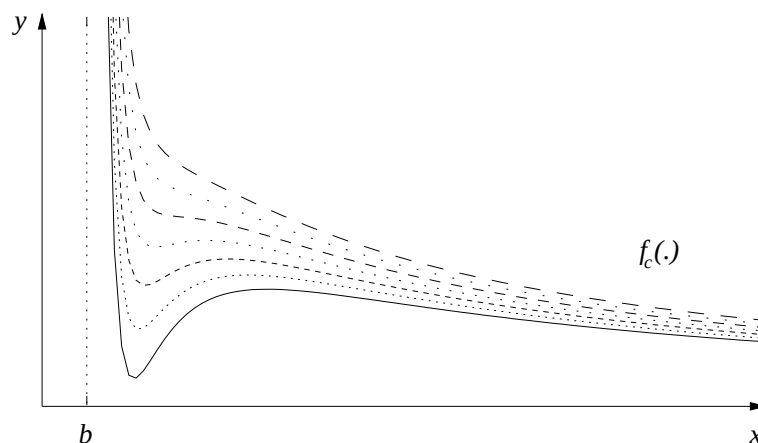
Es seien a, b, c positive Zahlen. Wir studieren im Bereich $x > b$ die Funktionen $f_c(x)$, die durch die Gleichung

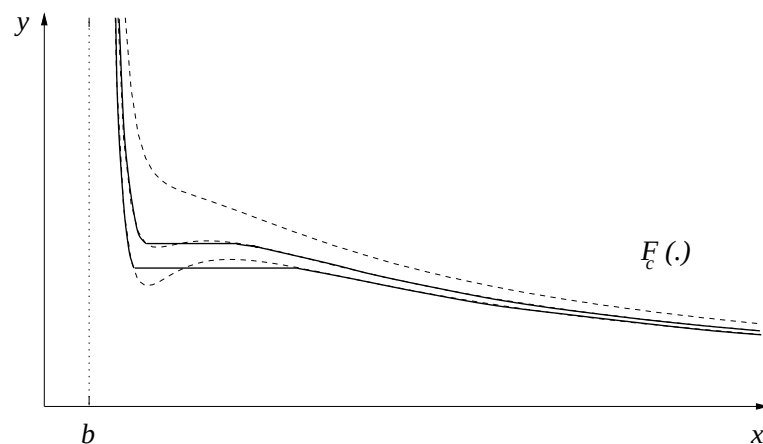
$$\left(y + \frac{a}{x^2}\right)(x - b) = c$$

gegeben sind. (y steht für $f_c(x)$).

Es handelt sich um eine auf den Nobelpreisträger van der Waals (1837–1923) zurückgehende approximative Beschreibung der Zustandsgleichung realer Gase. x steht für das Volumen, y für den Druck, c ist proportional zur absoluten Temperatur. Die Konstanten a und b hängen davon ab, mit welchem Gas man es zu tun hat. Der Grenzfall $a = 0 = b$ liefert die Zustandsgleichung des idealen Gases.

Das Bild zeigt die Funktion $f_c(\cdot)$ für einige c . Für die physikalische Interpretationen ist das zweite Bild maßgeblich. Für nicht allzu hohe Temperaturen ändert sich der Druck durch die Veränderung des Volumens in einem gewissen Bereich nicht; Volumenänderung führt zum Verdampfen einer Flüssigkeitsmenge bzw. zum Kondensieren einer Gasmenge.





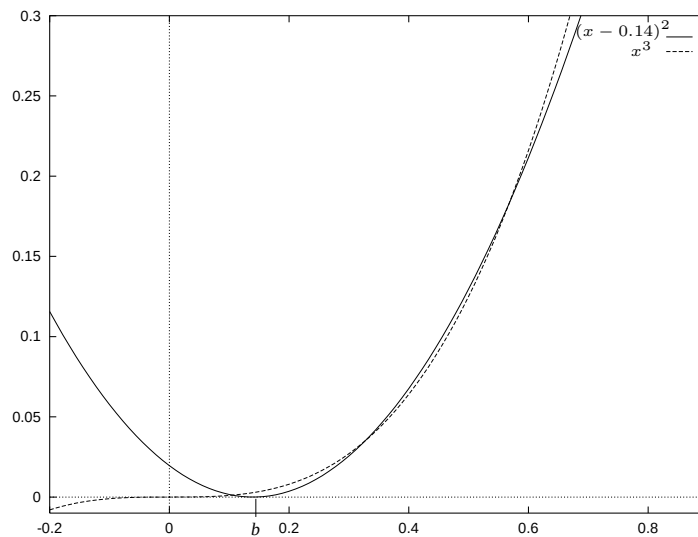
Eine qualitative Diskussion der Gleichung für festes (a, b) in einem gewissen Bereich ergibt, daß $f_c(\cdot)$ für $c \geq c^*$ (die kritische Temperatur) antiton ist.

Wir bestimmen c^* :

$$f_c(x) = \frac{a}{x^2} + \frac{c}{x-b}$$

$$f'_c(x) = \frac{2a}{x^3} - \frac{c}{(x-b)^2} .$$

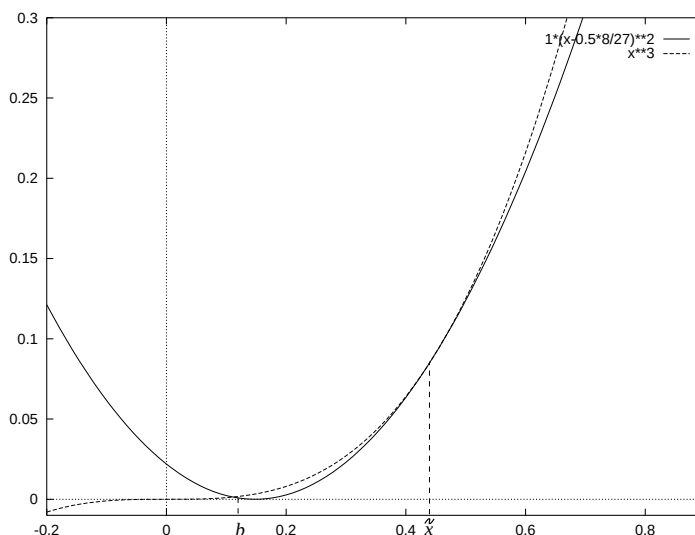
Wir interessieren uns, ob $f'_c(x)$ für gewisse $x > b$ den Wert 0 annehmen kann, ob also die Gleichung $\frac{2a}{c} (x-b)^2 = x^3$ Lösungen $x > b$ besitzt.



Das Bild zeigt die Funktionen

$$x \mapsto \frac{2a}{c} (x-b)^2 \quad \text{und} \quad x \mapsto x^3 .$$

Für $x \rightarrow \infty$ steigt die zweite schneller an als die erste. Für kleine Parameter $\frac{2a}{c}$ schneiden sich die Funktionsgraphen im Bereich $x > b$ nicht. Für große $\frac{2a}{c}$ gibt es im Bereich $x > b$ genau zwei Lösungen; besonders interessant ist der Fall, wo diese beiden Lösungen zusammenfallen, wo also die Ableitungen im Schnittpunkt gleich sind.



In diesem Fall gibt es ein $\tilde{x} > b$ mit

$$\frac{2a}{c} (\tilde{x} - b)^2 = \tilde{x}^3 \quad \text{und} \quad \frac{4a}{c} (\tilde{x} - b) = 3\tilde{x}^2 .$$

Aus den beiden Gleichungen ergibt sich durch Division

$$\frac{2a}{c} \cdot \frac{c}{4a} (\tilde{x} - b) = \frac{1}{3} \tilde{x} , \quad \text{also} \quad \tilde{x} - b = \frac{2}{3} \tilde{x} , \quad \tilde{x} = 3b .$$

Dies, eingesetzt in die zweite Gleichung, ergibt $\frac{a}{c} = \frac{27}{8} b$. Die kritische Temperatur ist also $c^* = \frac{8}{27} \cdot \frac{a}{b}$.

Die Funktion $F_c(\cdot)$, die die Isotherme physikalisch beschreibt, ergibt sich aus $f_c(\cdot)$ dadurch, daß man die Nichtisotonie so bereinigt, daß

$$\int_{x_n}^{x_0} [F_c(x) - f_c(x)] dx = 0 .$$

Für $x \leq x_u(c)$ ist aller Dampf verflüssigt; für $x \geq x_0(c)$ ist alle Flüssigkeit verdampft.

Die mathematische Funktion $F_c(\cdot)$ beschreibt die realen Verhältnisse natürlich nur approximativ. Die Modellvorstellungen, die zu dieser einfachen Gleichung führen, sind recht grob. Eine Diskussion findet man z.B. im berühmten Lehrbuch: R.W. POHL: Einführung in die Physik, Mechanik, Akustik und Wärmelehre. Springer, 17. Aufl. 1969.

Anhang : Zum Lehrprogramm

Struik schrieb 1963 (er war damals Professor emeritus der Mathematik am MIT): „*Eine ungewöhnlich fruchtbare Periode mathematischer Schaffenskraft wurde durch die Veröffentlichung der Arbeiten von 1684 und 1686 eingeleitet. Nach 1687 war Leibniz eng mit den Brüdern Bernoulli verbunden, die seine Methode begierig aufnahmen. Noch vor 1700 hatten diese Männer das meiste von dem gefunden, was heute den Studenten der Differential- und Integralrechnung geboten wird, dazu aber wichtige Teile höherer Gebiete einschließlich einiger Probleme aus der Variationsrechnung. . . .*“

Heute haben die Studenten der Mathematik keine Muße, sich wie die Bernoullis über all die interessanten Kurven zu freuen: Kettenlinie, Lemniskate, Isochrone, Brachistochrone etc. Man befaßt die Studenten hauptsächlich mit der Begründung des Kalküls. Es wird allenfalls oberflächlich auf die Spezialitäten der alten Meister verwiesen. Courant bemerkt dazu in seinem seit dem ersten Erscheinen 1927 bis heute höchst populärem Lehrbuch (Vorbemerkungen zur 3. Auflage von 1955): „*Die Grundtendenz aller neueren Mathematik ist die Ersetzung von Einzelbetrachtungen durch immer allgemeinere systematische Methoden, welche vielleicht nicht immer den individuellen Zügen des einzelnen Falles in vollem Maße gerecht werden, aber doch durch die Kraft ihrer Allgemeinheit eine Fülle neuer Resultate versprechen. Auf der anderen Seite gewinnt die Zahl, die analytische Behandlung der Dinge, immer mehr ihr selbständiges Recht, um schließlich gänzlich zur Herrschaft über die Geometrie zu gelangen. . . . Dabei sei auf eine Gefahr, die aus der eingangs erwähnten Diskontinuität (zur geometrischen Betrachtung der einzelnen mathematischen Objekte in der Elementarmathematik) entspringt, besonders hingewiesen. Der elementare Standpunkt der Schulmathematik verleitet zwar dazu, an Einzelheiten haften zu bleiben und den Blick für die allgemeinen Zusammenhänge und systematischen Methoden zu verlieren. Der „höhere Standpunkt“ der allgemeinen Methode birgt aber die umgekehrte Gefahr in sich, daß der Zusammenhang mit dem konkreten Einzelnen verlorengeht und daß man den einfachsten individuellen Schwierigkeiten ratlos gegenübersteht, weil man in der Welt allgemeiner Begriffe verlernt hat, das Konkrete zu sehen und zu fassen. Der Leser muß mit eigenen Kräften dafür sorgen, durch dieses Dilemma hindurchzukommen. Nur das immer wiederholte selbständige Durchdenken der Einzelheiten und das vollständige Erfassen der allgemeinen Gedanken im speziellen Beispiel kann zu diesem Ziele führen, und hierin liegt die Hauptaufgabe für jeden, der sich um das Studium der Wissenschaft bemüht.*“

A.6 Elementare Lösungen von Differentialgleichungen

Beispiel zum Einstieg

Bei gewissen einfachen Bimolekülreaktionen wächst die Menge x des Reaktionsprodukts mit einer Geschwindigkeit, die proportional zum Produkt der gerade vorhandenen Mengen der reagierenden Ausgangsstoffe ist („Massenwirkungsgesetz“). Am Anfang sei a die Menge des einen und b die Menge des anderen Ausgangsstoffes (o.B.d.A. $a < b$). Dies führt auf die Differentialgleichung

$$dx = k \cdot (a - x) \cdot (b - x) dt$$

Wegen

$$\frac{b - a}{(a - x)(b - x)} = \frac{1}{a - x} - \frac{1}{b - x} \quad (\text{Partialbruchzerlegung})$$

kann man das auch so schreiben

$$\left(\frac{1}{a - x} - \frac{1}{b - x} \right) dx = k \cdot (b - a) dt .$$

Der Zuwachs auf der linken Seite ist gleich dem Zuwachs auf der rechten Seite. Wir gehen zu den Stammfunktionen über

$$\ln \frac{a - x}{b - x} + \text{const} = (a - b) \cdot k \cdot t .$$

Die Konstante ergibt sich aus $x(0) = 0$

$$\begin{aligned} \ln \frac{1 - \frac{x}{a}}{1 - \frac{x}{b}} &= (a - b)kt \\ 1 - \frac{x}{a} &= \left(1 - \frac{x}{b}\right) e^{(a-b)kt} \\ x(t) &= \frac{a \cdot b \cdot (1 - e^{(a-b)kt})}{b - a \cdot e^{(a-b)kt}} \quad \text{für } t \geq 0 . \end{aligned}$$

Damit ist die Menge des Reaktionsprodukts zur Zeit t bestimmt.

Die Suche nach einer (elementaren) Stammfunktion $y = Y(x)$ zu einer gegebenen (elementaren) Funktion $f(x)$ (in einer Umgebung von x_0 mit $Y(x_0) = y_0$) kann man als die Aufgabe interpretieren, eine Lösung der Differentialgleichung

$$dy = f(x) dx \quad \text{durch den Punkt } (x_0, y_0)$$

zu finden. Integration liefert

$$Y(x) - y_0 = \int_{x_0}^x f(\xi) d\xi .$$

Wir betrachten allgemeinere Differentialgleichungen

$$dy = f(x, y) dx .$$

Definition Die Funktion $Y(x)$ ist eine Lösung mit $Y(x_0) = y_0$, wenn

$$Y(x) - y_0 = \int_{x_0}^x f(\xi, Y(\xi)) d\xi .$$

Hier wird x als unabhängige und y als abhängige Variable betrachtet. Man kann (x, y) auch symmetrisch behandeln, indem man einen Kurvenparameter einführt und nach $(x(t), y(t))$ für t in einer Umgebung von t_0 sucht, so daß

$$\dot{y}(t) dt = f(x(t), y(t)) \cdot \dot{x}(t) dt .$$

In integrierter Form ergibt das für die Lösung mit $(x(t_0), y(t_0)) = (x_0, y_0)$

$$(*) \quad y(t) - y_0 = \int_{t_0}^t f(x(t), y(t)) \cdot \dot{x}(t) dt .$$

Die Frage ist nun, ob es eine (lokale) Lösung gibt: Gibt es $(x(t), y(t))$ mit dieser Eigenschaft (*) für alle t in einer Umgebung von t_0 ?

Existenz- und Eindeutigkeitsfragen für lokale Lösungen von Differentialgleichungen werden uns im C-Zug beschäftigen. Hier geht es darum, heuristisch mit Hilfe der Methode der Stammfunktionen in einfachen Fällen Lösungen zu gewinnen. Ob das, was beim heuristischen Vorgehen herauskommt, das Problem löst, muß dann durch eine Probe bestätigt werden. Die Frage nach der Eindeutigkeit muß zurückgestellt werden.

Beispiel Mit der Methode von oben gelangt man zu der Vermutung, daß

$$Y(x) = \frac{a}{b} \cdot \frac{1}{1 - c \cdot e^{ax}} \quad \text{mit } c = \frac{y_0 - \frac{a}{b}}{y_0}$$

eine Lösung der Differentialgleichung

$$dy = (a - by) \cdot y \cdot dx$$

ist. Die Probe bestätigt das.

Betrachten wir allgemeiner eine Differentialgleichung

$$\boxed{M(x, y) dx + N(x, y) dy = 0 .}$$

In zwei Fällen können wir mit der Methode der Stammfunktionen Lösungen angeben

A) („Getrennte Variable“)

Nehmen wir an, daß

$$\frac{M(x, y)}{N(x, y)} = -f(x) \cdot g(y) .$$

Die Differentialgleichung hat dann die Gestalt

$$dy = f(x) \cdot g(y) dx \quad \text{oder} \quad \frac{1}{g(y)} dy = f(x) dx .$$

Diese Beziehung zwischen den infinitesimalen Zuwächsen liefert eine Beziehung zwischen den Stammfunktionen

$$\int_{y_0}^y \frac{1}{g(y)} dy = \int_{x_0}^x f(\xi) d\xi .$$

Beispiel

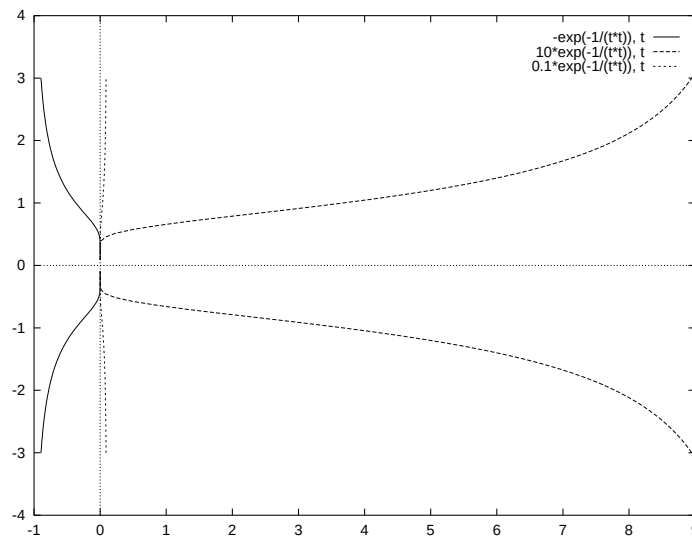
$$-\frac{1}{2} y^3 dx + x dy = 0 .$$

Die Variablen lassen sich trennen. Man sieht aber auch, daß mit den Lösungen etwas Irreguläres passieren kann, wenn sich x oder y dem Wert 0 nähert. Wir studieren die Lösungskurve durch (x_0, y_0) mit $x_0 \neq 0$, $y_0 \neq 0$.

$$\begin{aligned} \frac{2}{y^3} dy &= \frac{1}{x} dx \\ \int_{y_0}^y \frac{2}{y^3} dy &= \int_{x_0}^x \frac{1}{\xi} d\xi \\ \left[-\frac{1}{y^2} \right]_{y_0}^y &= [\ln |x|]_{x_0}^x \\ -\frac{1}{y^2} + \frac{1}{y_0^2} &= \ln \frac{x}{x_0} . \end{aligned}$$

Diese Gleichung kann man sowohl nach y als auch nach x auflösen.

$$\begin{aligned} Y(x) &= \pm \left(\frac{1}{y_0^2} - \ln \frac{x}{x_0} \right)^{-1/2} \\ X(y) &= \text{const} \cdot \exp \left(-\frac{1}{y^2} \right) \end{aligned}$$

**B) („Exakte Form“)**

$$M(x, y) dx + N(x, y) dy$$

heißt eine *exakte* Differentialform, wenn es eine Funktion $F(x, y)$ gibt, so daß

$$M(\cdot, \cdot) = F_1(\cdot, \cdot), \quad N(\cdot, \cdot) = F_2(\cdot, \cdot).$$

Hierbei bezeichnet $F_1(\cdot, \cdot)$ die partielle Ableitung nach der „ersten“ Variablen x und $F_2(\cdot, \cdot)$ die partielle Ableitung nach der „zweiten“ Variablen y .

Mit partiellen Ableitungen werden wir uns später ausführlich beschäftigen.

Nehmen wir das Beispiel von Newton (siehe oben)

$$F(x, y) = x^3 - ax^2 + axy - y^3$$

$$F_1(x, y) = 3x^2 - 2ax + ay$$

$$F_2(x, y) = ax - 3y^2$$

Dieses $F(\cdot, \cdot)$ führt auf die exakte Differentialform

$$(3x^2 - 2ax + ay) dx + (ax - 3y^2) dy.$$

Betrachten wir nun die Differentialgleichung

$$F_1(x, y) dx + F_2(x, y) dy = 0.$$

Wir erhalten die **Lösung in impliziter Gestalt**

$$F(x, y) = \text{const}.$$

Die Lösung durch (x_0, y_0) erhalten wir, wenn wir die Konstante gleich $F(x_0, y_0)$ setzen.

Beispiel Zu lösen ist die Differentialgleichung

$$(3y + e^x) dx + (3x + \cos y) dy = 0$$

sagen wir durch den Punkt $(x_0, y_0) = (0, \pi)$. Die Differentialform ist exakt; denn für

$$F(x, y) = 3xy + e^x + \sin y$$

erhalten wir

$$F_1(x, y) = 3y + e^x ; \quad F_2(x, y) = 3xy + \cos y .$$

Die gesuchte Lösung der Differentialgleichung (in implizierter Gestalt) ist

$$3xy + e^x + \sin y = 1 .$$

Es war A.C. CLAIRAUT (1713–1765), der die Bedingung der Exaktheit von Differentialformen

$$M(x, y) dx + N(x, y) dy$$

herausgearbeitet hat. Clairaut hat sie in seiner Theorie der Gleichgewichte von Flüssigkeiten und der Anziehung von Rotationsellipsoiden entwickelt, um die Erdgestalt zu verstehen („Théorie de la figure de terre“, 1743). Clairaut hat seine Methode weitergeführt in seiner „Théorie de la lune“, 1752. Die Begründung stand zwar auf schwachen Füßen; aber die Methode funktionierte. In moderner Notation lautet die Bedingung von Clairaut

Theorem (Clairaut) Die Differentialform

$$M(x, y) dx + N(x, y) dy$$

ist genau dann exakt, wenn

$$M_2(x, y) = N_1(x, y) .$$

(Vorausgesetzt $M(\cdot, \cdot)$ und $N(\cdot, \cdot)$ sind entsprechend glatt).

Den Beweis verschieben wir auf später.

Es kann als glücklicher Zufall gelten, wenn eine Differentialgleichung

$$M(x, y) dx + N(x, y) dy = 0$$

in der Form gegeben ist, daß $M(x, y) dx + N(x, y) dy$ exakt ist. Beispielsweise ist

$$(a - by)y \cdot dx - dy = 0$$

eine unpassende, jedoch

$$dx - \frac{1}{(a - by)y} dy = 0$$

ein passende Form.

Definition Die Funktion $m(x, y)$ ($\neq 0$ in der Nähe von (x_0, y_0)) heißt ein *integrierender Faktor* für die Differentialgleichung $M(x, y) dx + N(x, y) dy = 0$, wenn

$$m(x, y) \cdot M(x, y) dx + m(x, y) \cdot N(x, y) dy$$

eine exakte Differentialform ist.

Leute mit Erfahrung und Fingerspitzengefühl erraten manchmal zu einer gegebenen Differentialgleichung

$$M(x, y) dx + N(x, y) dy = 0$$

einen integrierenden Faktor $m(x, y)$.

Mit der Methode der Stammfunktionen erhalten sie dann eine Funktion $F(x, y)$ mit

$$F_1(x, y) = m(x, y)M(x, y)$$

$$F_2(x, y) = m(x, y)N(x, y) ;$$

und das liefert ihnen dann eine Lösung der Differentialgleichung (in impliziter Gestalt)

$$F(x, y) = F(x_0, y_0) .$$

Beispiel

$$(xy^2 + y) dx - x dy = 0$$

Die Differentialform ist nicht exakt

$$m(x, y) = \frac{1}{y^2}$$

ist ein integrierender Faktor; denn

$$m(x, y) \cdot (x \cdot y^2 + y) dx + m(x, y)(-x) dy = \left(x + \frac{1}{y}\right) dx - \frac{x}{y^2} dy$$

ist exakt. Für

$$F(x, y) = \frac{x}{y} + \frac{x^2}{2}$$

erhalten wir

$$F_1(x, y) = x + \frac{1}{y}$$

$$F_2(x, y) = -\frac{x}{y^2} .$$

Die Lösungen haben also die Gestalt

$$\frac{x}{y} + \frac{x^2}{2} = c .$$

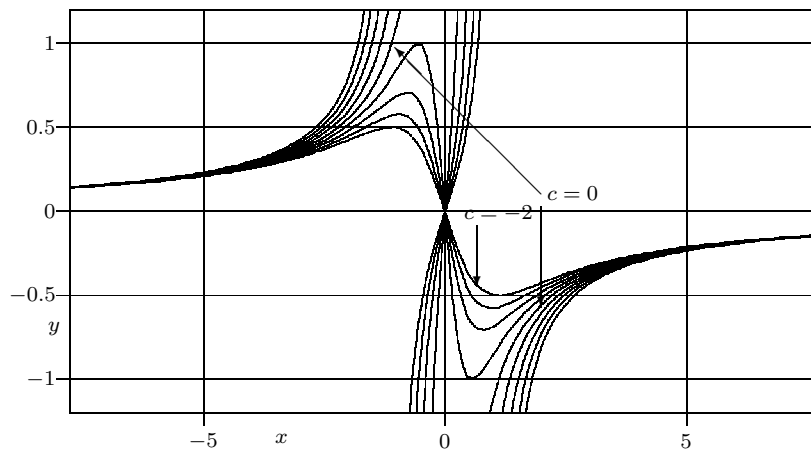
Man kann sowohl nach y als auch nach x auflösen

$$Y(x) = \frac{2x}{2c - x^2} .$$

(Die Abszissenachse $y \equiv 0$ ist ebenfalls eine Lösungskurve.)

Im Falle $c < 0$ erstrecken sich die Lösungen über die ganze Achse. Für $C > 0$ erhalten wir Lösungskurven in $(-\infty, -\sqrt{2c}) \cup (-\sqrt{2c}, +\sqrt{2c}) \cup (\sqrt{2c}, +\infty)$

$$X(y) = \frac{1}{y} \pm \sqrt{2c + \frac{1}{y^2}}$$



Einige elementarlösbare Differentialgleichungen

Wir betrachten Differentialgleichungen der Form

$$dy = f(t, y) dt$$

für einige spezielle $f(\cdot, \cdot)$. Die „unabhängige Variable“ t wird als die Zeit interpretiert. Zur Zeit t_0 setzen wir $y(t_0) = y_0$ fest und studieren den zeitlichen Verlauf $y(\cdot)$ in einer Umgebung von t_0 . Man schreibt auch

$$\dot{y} = f(y, t) .$$

Lineare Differentialgleichungen

Man spricht von einer **linearen Differentialgleichung**, wenn $f(\cdot, \cdot)$ als Funktion von y affin ist

$$\dot{y} + a(t)y = b(t) .$$

Dabei sind $a(\cdot)$ und $b(\cdot)$ als stetig angenommen. Im Falle $b(\cdot) \equiv 0$ spricht man von einer *homogenen* linearen Differentialgleichung

$$\dot{y} + a(t)y = 0 .$$

Offenbar sind Linearkombinationen von Lösungen ebenfalls Lösungen. Betrachte

$$\begin{aligned} A(t) &= \int_{t_0}^t a(s) ds \\ z(t) &= \exp(-A(t))y_0 . \end{aligned}$$

Es gilt

$$\begin{aligned} dz &= e^{-A(t)}(-a(t))y_0 dt = -a(t) \cdot z(t) dt \\ \dot{z} + a(t) \cdot z(t) &= 0 \end{aligned}$$

Die Funktionen $z(t)$ lösen also die homogene Gleichung. Sie sind die einzigen Lösungen der homogenen linearen Gleichung, denn wenn $\bar{y}(t)$ irgendeine Lösung ist, dann ist $e^{A(t)}\bar{y}(t)$ eine Konstante wegen

$$d(e^{A(t)}\bar{y}(t)) = e^{A(t)}(d\bar{y} + a(t)\bar{y} \cdot dt) \equiv 0 .$$

Die Lösungen der *inhomogenen* Gleichung

$$\dot{y} + a(t)y = b(t)$$

ergeben sich als die Summe einer speziellen Lösung und einer beliebigen Lösung der homogenen Gleichung

$$y(t) = e^{-A(t)}y_0 + \exp(-A(t)) \int_{t_0}^t \exp(A(s))b(s) ds .$$

Man kann das einfach nachrechnen. Es gibt aber auch einen Trick, wie man die Lösung finden kann, den sog. Ansatz der „**Variation der Konstanten**“:

Gesucht ist $C(t)$ so, daß

$$z(t) := \exp(-A(t))C(t)$$

die inhomogene Gleichung löst. Nach der Produktregel gilt

$$\begin{aligned}\dot{z}(t) &= \exp(-A(t))\dot{C}(t) - \exp(-A(t))a(t)C(t) \\ &= \exp(-A(t))\dot{C}(t) - a(t)z(t) .\end{aligned}$$

Zu lösen ist also die Gleichung

$$\exp(-A(t))\dot{C}(t) = b(t) .$$

Das ist eine einfache Integrationsaufgabe. Sie führt auf

$$C(t) = C(t_0) + \int_{t_0}^t \exp(A(s))b(s) ds$$

und damit auf die obige Formel.

Hinweis Wenn $a(t)$ und $b(t)$ Konstante sind, dann spricht man von einer **linearen Differentialgleichung mit konstanten Koeffizienten**. In diesem Falle haben wir

$$A(t) = tA \quad \text{also} \quad y(t) = \exp(-tA)y(t_0)$$

als allgemeine Lösung der homogenen Gleichung.

Verallgemeinerung In den Anwendungen interessiert man sich für Systeme linearer Differentialgleichungen

$$\dot{y} + a(t) \cdot y = b(t)$$

$$\begin{pmatrix} \dot{y}_1 \\ \dot{y}_2 \\ \vdots \\ \dot{y}_n \end{pmatrix} + \begin{pmatrix} a_{11}(t) & a_{12}(t) & \cdots & a_{1d}(t) \\ a_{21}(t) & a_{22}(t) & \cdots & a_{2d}(t) \\ \vdots & \vdots & \ddots & \vdots \\ a_{d1}(t) & a_{d2}(t) & \cdots & a_{dd}(t) \end{pmatrix} \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_d \end{pmatrix} = \begin{pmatrix} b_1(t) \\ b_2(t) \\ \vdots \\ b_d(t) \end{pmatrix}$$

Wer etwas Matrizenkalkül beherrscht, sieht sofort, daß im d -dimensionalen Fall alles fast genau so ist wie im eindimensionalen Fall.

$$A(t) = \int_{t_0}^t a(s) ds$$

und

$$\exp(-A(t)) = I - A(t) + \frac{1}{2!}A^2(t) - \frac{1}{3!}A^3(t) + \dots$$

sind $d \times d$ -Matrizen. Zu jeder Spalte y_0 ist

$$y(t) = \exp(-A(t))y_0$$

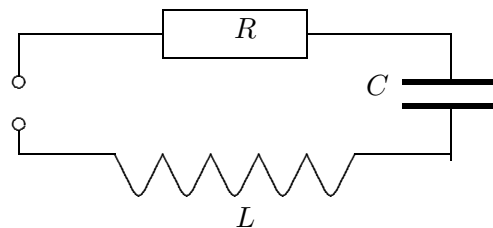
die eindeutig bestimmte Lösung der homogenen Gleichung mit der Anfangsbedingung $y(t_0) = y_0$

$$y(t) = \exp(-A(t)) \int_{t_0}^t \exp(A(s))b(s) ds$$

ist eine spezielle Lösung der inhomogenen Gleichung.

Alle weiteren Lösungen der inhomogenen Gleichung ergeben sich durch Addition einer Lösung der homogenen Gleichung.

Anwendung Ein aufgeladener Kondensator (mit Kapazität C) wird über eine Spule (mit der „Impedanz“ oder „Selbstinduktivität“ L) und einem Ohmschen Widerstand R entladen



Die Spannung U ist durch die Ladung des Kondensators gegeben. Die zeitliche Änderung $\dot{U}$ durch den Strom I :

$$C \cdot \dot{U} = -I .$$

Am Ohmschen Widerstand ist die Spannung (wegen der Selbstinduktivität) gleich

$$U - L \cdot \dot{I} = R \cdot I .$$

Das ergibt ein System linearer Differentialgleichungen mit konstanten Koeffizienten

$$\begin{pmatrix} \dot{U} \\ \dot{I} \end{pmatrix} = \begin{pmatrix} 0 & -\frac{1}{C} \\ \frac{1}{L} & -\frac{R}{L} \end{pmatrix} \begin{pmatrix} U \\ I \end{pmatrix}$$

Wenn man nicht einfach den Schalter schließt um die Entladung herbeizuführen, sondern stattdessen dort eine von einer äußern Spannungsquelle gelieferte Spannung anlegt, dann führt das auf ein inhomogenes Gleichungssystem („erzwungene Schwingungen“).

Man kann das System in eine Differentialgleichung zweiter Ordnung umschreiben

$$L \cdot \ddot{I} + R \cdot \dot{I} + \frac{1}{C} I = 0 .$$

Die Lösungen findet man durch den Ansatz der komplexen Amplituden

$$I(t) = \exp(i\omega t)I(0) .$$

Von ω ist zu fordern, daß die folgende quadratische Gleichung erfüllt ist.

$$(-\omega^2)L + (i\omega)R + \frac{1}{C} = 0 .$$

In der Einführung in die Elektrizitätslehre diskutiert man die verschiedenen Fälle, die da mit reellen oder komplexen Wurzeln auftreten können. Wir bemerken hier nur ganz kursorisch:

Im Falle $R = 0$ schwingt der Schwingkreis periodisch. Im Falle $R \neq 0$ haben wir eine sog. gedämpfte Schwingung mit der Frequenz

$$\sqrt{\frac{1}{LC} - \frac{R^2}{4L^2}}$$

und der „Dämpfungskonstanten“ $\frac{R}{2L}$.

Auf eine formal identische Gleichung führt die Konstellation des „harmonischen Oszillators“. Man stelle sich einen Massenpunkt vor, der von einer Feder gehalten wird, bei welcher die rücktreibende Kraft proportional zur Auslenkung x ist.

$$\text{Kraft} = \text{Masse} \times \text{Beschleunigung} = m \cdot \ddot{x} \stackrel{!}{=} -kx .$$

Die Reibung setzen wir proportional zu $\dot{x}$ an. (Reibungskraft $= -r \cdot \dot{x}$)

Daraus ergibt sich die Schwingungsgleichung

$$m\ddot{x} + r\dot{x} + kx = 0 .$$

Weitere klassische Typen von Differentialgleichungen

Zum Kreis um die Bernoullis gehörte auch RICATTI (1676–1754). Nach ihm ist die Gleichung $\dot{y} = f(t, y)$ benannt, wo $f(t, \cdot)$ quadratisch in y ist

$$\dot{y} + a(t)y + b(t)y^2 = c(t) .$$

Den Fall $c(t) \equiv 0$ kann man in einen allgemeineren Rahmen stellen. Wir behandeln diesen Fall zuerst und kommen dann zur Riccati-Gleichung zurück.

Sei α beliebig reell ($\neq 0, \neq 1$) und $g(t), h(t)$ stetig;

$$\dot{y} + g(t)y + h(t)y^\alpha = 0$$

heißt dann eine *Bernoullische Differentialgleichung*. Man führt sie durch eine Transformation der abhängigen Variablen y auf eine lineare Differentialgleichung zurück. Setze

$$x(t) = (y(t))^{1-\alpha}$$

$$\dot{x}(t) = (1-\alpha)(y(t))^{-\alpha} \cdot \dot{y}(t)$$

$$\begin{aligned} (1-\alpha)y^{-\alpha}(\dot{y} + g(t)y + h(t)y^\alpha) &= \dot{x} + (1-\alpha)g(t)y^{1-\alpha} + (1-\alpha)h(t) \\ &= \dot{x} + (1-\alpha)g(t)x + (1-\alpha)h(t) . \end{aligned}$$

Fazit $y(\cdot)$ löst die Bernoullische Differentialgleichung genau dann, wenn $x(\cdot) = (y(\cdot))^{1-\alpha}$ die angegebene lineare Differentialgleichung löst.

Beispiel $\dot{y} + \frac{1}{3} y + \frac{t}{3} y^4 = 0$.

Hier ist $\alpha = 4$, $g(t) = \frac{1}{3}$, $h(t) = \frac{t}{3}$.

$$x(\cdot) = (y(\cdot))^{-3}$$

löst die lineare Differentialgleichung

$$\begin{aligned}\dot{x} - x &= t \\ x(t) &= c \cdot e^t - (t+1).\end{aligned}$$

Die allgemeine Lösung der gegebenen Bernoullischen Differentialgleichung hat also die Gestalt

$$y(t) = [\text{const} \cdot e^t - (t+1)]^{-1/3}.$$

(Der Punkt t^* , in welchem $y(\cdot)$ verschwindet, ist ein Ausnahmepunkt).

Zurück zur Riccati-Gleichung

Es gibt kein allgemeingültiges Lösungsverfahren. Man kann aber alle Lösungen beschreiben, wenn man eine spezielle Lösung $y_1(\cdot)$ erraten kann.

$$\dot{y}_1 + a(t)y_1 + b(t)y_1^2 = c(t).$$

Wenn $y(\cdot)$ irgendeine weitere Lösung ist, dann haben wir für die Differenz

$$\begin{aligned}z(t) &= y(t) - y_1(t) \\ \dot{z} = \dot{y} - \dot{y}_1 &= -a(t)(y - y_1) - b(t)(y^2 - y_1^2) \\ &= -a(t)z - b(t)z(z + 2y_1) \\ \dot{z} + (a(t) + 2b(t)y_1(t))z + b(t)z^2 &= 0.\end{aligned}$$

Dies ist eine Bernoulli-Gleichung, die man bekanntlich durch Quadratur (d.h. Stammfunktionsbildung) lösen kann..

Beispiel $\dot{y} + (2t-1)y - y^2 = (1-t+t^2)$.

Eine spezielle Lösung ist $y_1(t) = t$. Für

$$z(t) = y(t) - t$$

erhalten wir die Bernoulli-Gleichung

$$\dot{z} - z - z^2 = 0.$$

Für $u = z^{1-\alpha} = \frac{1}{z}$ ergibt sich

$$\dot{u} + u = -1$$

$$u(t) = c e^{-t} - 1$$

$$z(t) = \frac{1}{c e^{-t} - 1}$$

$$y(t) = t + \frac{1}{c e^{-t} - 1} \quad \text{für } t \neq \ln c .$$

Schlußbemerkung Es gibt eine Reihe von speziellen Differentialgleichungen, die in der Geometrie, der Mechanik oder auch in anderen Disziplinen immer wieder auftreten. Einige dieser Gleichungen haben zu ausgedehnten mathematischen Theorien Anlaß gegeben. Die Modellbildungen der Anwender führen nicht immer direkt auf die bekannten wohlstudierten Gleichungen; oft ist viel Kunstfertigkeit nötig um durch geeignete Variablentransformationen die Beziehungen aufzudecken. Diese Kunstfertigkeiten haben aber wenig zu tun mit dem, was heute die “Theorie der gewöhnlichen Differentialgleichungen” heißt. Dort geht es weniger darum, spezielle Differentialgleichungen „explizit“ zu lösen, als vielmehr darum, das qualitative Verhalten der Lösungskurven von recht allgemeinen Typen von Differentialgleichungen zu verstehen.

A.7 Die Gammafunktion

Das „erste Eulersche Integral“ ist definiert für $\alpha > 0$, $\beta > 0$:

$$B(\alpha, \beta) := \int_0^1 x^{\alpha-1} (1-x)^{\beta-1} dx .$$

Wir wollen uns zunächst mit dem „zweiten Eulerschen Integral“ beschäftigen

$$\Gamma(\lambda) := \int_0^\infty x^{\lambda-1} e^{-x} dx \quad \text{für } \lambda > 0 .$$

Für $\lambda \in (0, 1)$ ist der Integrand unbeschränkt. Wegen

$$\int_0^\varepsilon x^{\lambda-1} e^{-x} dx \leq \left[\frac{1}{\lambda} x^\lambda \right]_0^\varepsilon \longrightarrow 0 \quad \text{für } \varepsilon \rightarrow 0$$

und dem starken Abfallen von e^{-x} für $x \rightarrow \infty$ ist das Integral aber endlich. Für große λ kommt der (prozentual) überwiegende Anteil zum Integral von dem Integrationsgebiet $\lambda \pm \text{const} \cdot \sqrt{\lambda}$.

Das Integral wächst sehr schnell an für $\lambda \rightarrow \infty$.

$$\Gamma(n) = (n-1)! \quad \text{für } n \in \mathbb{N} .$$

Dies ergibt sich aus $\Gamma(1) = \int_0^\infty e^{-x} dx = 1$ und der

Funktionalgleichung Für alle $\lambda > 0$ gilt

$$\lambda \cdot \Gamma(\lambda) = \Gamma(\lambda + 1) .$$

Beweis durch partielle Integration.

Die Gamma-Funktion $\Gamma(\lambda)$ interpoliert also die Fakultätsfunktion. Im 18. Jahrhundert wurden auch andere Interpolationen der Fakultätsfunktion diskutiert. Einen sehr lesenswerten Bericht gibt PHILIP J. DAVIS: *Leonhard Euler's Integral. A historical profile of the Gamma function* in der Sammlung „The Chauvenet Papers“, ed. Abbott.

Mit Hilfe des Eulerschen Integrals kann man $\Gamma(z)$ auch für komplexe z mit $\lambda = \Re z > 0$ erklären

$$\Gamma(z) = \int_0^\infty x^{z-1} e^{-x} dx \quad \text{für } \Re z > 0 .$$

Es gilt nämlich für $x > 0$, $\lambda = \Re z > 0$

$$\begin{aligned} |x^{z-1} e^{-x}| &= |e^{(z-1)\ln x} e^{-x}| \\ &= e^{(\lambda-1)\ln x} e^{-x} = x^{\lambda-1} e^{-x} . \end{aligned}$$

Für die Fortsetzung auf den natürlichen Definitionsbereich

$$\mathbb{C} \setminus \{0, -1, -2, \dots\}$$

ist aber das Eulersche Integral nicht geeignet. Eine für alle z passende Darstellung hat Gauß angegeben.

Satz Für alle $z \in \mathbb{C} \setminus \{0, -1, -2, \dots\}$ existiert der Limes

$$\Gamma_G(z) = \lim_{n \rightarrow \infty} \frac{n! n^z}{z(z+1) \dots (z+n)}$$

und für $\Re z > 0$ gilt

$$\Gamma_G(z) = \int_0^\infty x^{z-1} e^{-x} dx .$$

Bemerkung

$$z \cdot \frac{n! n^z}{z(z+1) \dots (z+n)} = \frac{n! n^{z+1}}{(z+1)(z+2) \dots (z+1+n)} \cdot \frac{z+1+n}{n} .$$

Folgerung Wenn der Limes existiert, genügt er der Funktionalgleichung

$$z \cdot \Gamma_G(z) = \Gamma_G(z+1) \quad \text{für alle } z \in \mathbb{C} \setminus \{0, 1, 2, \dots\} .$$

Es genügt daher, die Konvergenz für z mit $\Re z \in (0, 1]$ zu beweisen. Es genügt auch, die Konvergenz für reelle Argumente $\lambda \in (0, 1]$ zu beweisen. Betrachten wir die Ableitungen des Logarithmus

$$\begin{aligned} \frac{d}{d\lambda} \ln \frac{n! n^\lambda}{\lambda(\lambda+1) \dots (\lambda+n)} &= - \left(\frac{1}{\lambda} + \frac{1}{\lambda+1} + \dots + \frac{1}{\lambda+n} \right) + \ln n \\ \frac{d^2}{d\lambda^2} \ln \frac{n! n^\lambda}{\lambda(\lambda+1) \dots (\lambda+n)} &= \frac{1}{\lambda^2} + \frac{1}{(\lambda+1)^2} + \dots + \frac{1}{(\lambda+n)^2} . \end{aligned}$$

Diese Funktionenfolgen haben einen Limes für $n \rightarrow \infty$. Die zweite Ableitung strebt gegen eine positive Funktion.

Nach dem Gesagten ist es nach den Maßstäben des A-Zuges evident, daß der Limes existiert und logarithmisch konvex ist. (Eine gewissenhafte Herleitung findet sich z.B. in Barner-Flohr *Analysis I*.) Wir werden sehen, daß diese Feststellungen genügen, um nachzuweisen, daß

$$\Gamma_G(\lambda) = \int_0^\infty x^{\lambda-1} e^{-x} dx \quad \text{für alle } \lambda > 0 .$$

Wir studieren nun das Eulersche Integral $\Gamma(\lambda)$ für $\lambda \in (0, \infty)$. Es handelt sich um eine positive Funktion. $k(\lambda) = \ln \Gamma(\lambda)$ ist also wohldefiniert. Wir zeigen, daß $k(\cdot)$ konvex ist.

Satz $\Gamma(\lambda)$ ist logarithmisch konvex in $(0, \infty)$.

Es gibt viele lehrreiche Beweise dieses fundamentalen Satzes. Wir führen hier drei Beweise, die aber zum Teil Fakten benützen, die wir noch nicht vollständig bewiesen haben. Einen kompletten Beweis ohne Anleihen bei Unbewiesenem zeichnet die Übungsaufgabe vor.

1. Beweis Wir wollen zeigen, daß für alle $\lambda, \mu > 0$ und alle $p, q > 1$ mit $\frac{1}{p} + \frac{1}{q} = 1$ gilt

$$k\left(\frac{\lambda}{p} + \frac{\mu}{q}\right) \leq \frac{1}{p} k(\lambda) + \frac{1}{q} k(\mu)$$

oder, äquivalent damit

$$\Gamma\left(\frac{\lambda}{p} + \frac{\mu}{q}\right) \leq (\Gamma(\lambda))^{1/p} \cdot (\Gamma(\mu))^{1/q}.$$

Dies bewerkstelligen wir mit Hilfe einer Verallgemeinerung der Hölderschen Ungleichung auf Integrale. (Den Beweis führen wir in C8.)

Betrachte

$$f(x) = x^{\frac{1}{p}(\lambda-1)} e^{-x/p}$$

mit der p -Norm $\|f\|_p = (\int (f(x))^p dx)^{1/p} = (\Gamma(\lambda))^{1/p},$

$$g(x) = x^{\frac{1}{q}(\mu-1)} e^{-x/q}$$

mit der q -Norm $\|g\|_q = (\int (g(x))^q dx)^{1/q} = (\Gamma(\mu))^{1/q}.$

Wir haben

$$f(x) \cdot g(x) = x^{(\frac{1}{p}\lambda + \frac{1}{q}\mu - 1)} e^{-x} \quad \text{mit} \quad \int f(x)g(x) dx = \Gamma\left(\frac{\lambda}{p} + \frac{\mu}{q}\right).$$

Die Höldersche Ungleichung

$$\int f(x)g(x) dx \leq \|f\|_p \cdot \|g\|_q$$

liefert die Behauptung.

2. Beweis $k(\lambda)$ ist zweimal differenzierbar. Wir zeigen $k''(\lambda) \geq 0$.

$$\begin{aligned} k'(\lambda) &= (\ln \Gamma(\lambda))' = \frac{1}{\Gamma(\lambda)} \cdot \Gamma(\lambda)' \\ k''(\lambda) &= \frac{1}{\Gamma(\lambda)} \cdot \Gamma''(\lambda) - \left(\frac{\Gamma'(\lambda)}{\Gamma(\lambda)} \right)^2. \end{aligned}$$

Wir zeigen

$$\Gamma(\lambda) \cdot \Gamma''(\lambda) - (\Gamma'(\lambda))^2 \geq 0$$

indem wir zeigen

$$\Gamma''(\lambda) \cdot \xi^2 - 2\Gamma'(\lambda) \cdot \xi + \Gamma(\lambda) \geq 0 \quad \text{für alle } \xi \in \mathbb{R}.$$

In der Integrationstheorie zeigt man, daß auch die Ableitungen durch Integrale dargestellt werden und zwar

$$\begin{aligned} \Gamma(\lambda) &= \int x^\lambda e^{-x} \frac{1}{x} dx \\ &= \int e^{\lambda \ln x} e^{-x} \frac{1}{x} dx \\ \Gamma'(\lambda) &= \int \ln x e^{\lambda \ln x} e^{-x} \frac{1}{x} dx \\ &= \int \ln x x^\lambda e^{-x} \frac{1}{x} dx \\ \Gamma''(\lambda) &= \int (\ln x)^2 x^\lambda e^{-x} \frac{1}{x} dx \\ \Gamma''(\lambda)\xi^2 - 2\Gamma'(\lambda)\xi + \Gamma(\lambda) &= \int (\xi^2(\ln x)^2 - 2\xi \ln x + 1) x^\lambda e^{-x} \frac{1}{x} dx \\ &= \int (\xi \ln x - 1)^2 e^{-x} \frac{1}{x} dx \geq 0. \end{aligned}$$

■

Ein wirkliches Verständnis für die logarithmische Konvexität der Gamma-Funktion auf $(0, \infty)$ kann man wohl nur dadurch gewinnen, daß man die Aussage weitreichend verallgemeinert. Man muß von Wahrscheinlichkeitsdichten (oder von Wahrscheinlichkeitsmaßen) im $\mathbb{R}^d$ und deren kumulantenenerzeugenden Funktionen sprechen. Wahrscheinlichkeitsdichten auf dem $\mathbb{R}^3$ kann man auch im Rahmen der Mechanik interpretieren.

Der Einfachheit halber betrachten wir nur den eindimensionalen Fall. Für eine nichtnegative Funktion auf $\mathbb{R}$

$$p(x) \geq 0 \quad \text{mit} \quad \int p(x) dx = 1,$$

welche für $|x| \rightarrow \infty$ genügend schnell abfällt, definiert man den Schwerpunkt und das Trägheitsmoment der Massenverteilung $p(\cdot)$:

$$\text{Schwerpunkt : } x^* = \int x p(x) dx$$

$$\text{Trägheitsmoment : } \sigma^2 = \int (x - x^*)^2 p(x) dx.$$

Außerdem definiert man die kumulantenenerzeugende Funktion

$$\psi(\lambda) := \ln \int e^{\lambda x} p(x) dx .$$

Lemma Kumulantenenerzeugende Funktionen sind stets konvex

Beweis (im eindimensionalen Fall)

Für $M(\lambda) = \int e^{\lambda x} p(x) dx$ gilt im Innern des Endlichkeitsintervalls

$$\begin{aligned} M'(\lambda) &= \int x e^{\lambda x} p(x) dx \\ M''(\lambda) &= \int x^2 e^{\lambda x} p(x) dx . \end{aligned}$$

Setze für λ im Innern des Endlichkeitsintervalls von $\psi(\lambda) = \ln M(\lambda)$

$$p_\lambda(x) = e^{\lambda x} p(x) e^{-\psi(\lambda)} .$$

Es gilt $\int p_\lambda(x) dx = 1$

$$\begin{aligned} \psi'(\lambda) &= \frac{M'(\lambda)}{M(\lambda)} = \int x p_\lambda(x) dx \\ \psi''(\lambda) &= \frac{M''(\lambda)}{M(\lambda)} - \left(\frac{M'(\lambda)}{M(\lambda)} \right)^2 = \int x^2 p_\lambda(x) dx - (\psi'(\lambda))^2 \\ &= \int (x - \psi'(\lambda))^2 p_\lambda(x) dx \geq 0 . \end{aligned}$$

Durch Spezialisierung dieses Lemmas erhalten wir den

3. Beweis für die logarithmische Konvexität der Gammfunktion

$$\Gamma(\lambda) = \int_0^\infty x^\lambda e^{-x} \frac{1}{x} dx = \int_0^\infty e^{\lambda \ln x} e^{-x} \frac{1}{x} dx .$$

Mit $y = \ln x$, $dy = \frac{1}{x} dx$

$$\Gamma(\lambda) = \int_{-\infty}^{+\infty} e^{\lambda y} \exp(-e^y) dy \quad \text{für } \lambda > 0 .$$

$\ln \Gamma(\lambda)$ ist eine kumulantenenerzeugende Funktion.

Theorem (Charakterisierung der Gamma-Funktion)

Es gibt genau eine Funktion $F(\lambda)$ auf $(0, \infty)$ mit den Eigenschaften

- (i) $\lambda F(\lambda) = F(\lambda + 1)$
- (ii) $F(\cdot)$ ist logarithmisch konvex
- (iii) $F(1) = 1$

Beweis Angenommen wir hätten ein $F(\cdot)$ mit den Eigenschaften (i), (ii), (iii).

- 1) Es genügt $F(\lambda)$ für $\lambda \in (0, 1)$ zu bestimmen. Die Beweisidee orientiert sich an der Vermutung, daß für große n und $\lambda \in (0, 1)$

$$F(n + \lambda) \approx F(n) \cdot n^\lambda, \quad F(n + \lambda)(n + \lambda)^{1-\lambda} \approx F(n + 1)$$

in dem Sinne, daß der Quotient nahe bei 1 liegt. In der Tat gilt

$$\begin{aligned} F(n + \lambda) &= F((1 - \lambda)n + \lambda(n + 1)) \\ &\leq (F(n))^{1-\lambda} (F(n + 1))^\lambda = (F(n))^{1-\lambda} (nF(n))^\lambda = F(n)n^\lambda \\ F(n + 1) &= F(\lambda(n + \lambda) + (1 - \lambda)(n + \lambda + 1)) \\ &\leq (F(n + \lambda))^\lambda (F(n + \lambda + 1))^{1-\lambda} = F(n + \lambda)(n + \lambda)^{1-\lambda} \\ n^\lambda &\geq \frac{F(n + \lambda)}{F(n)} = \frac{nF(n + \lambda)}{F(n + 1)} \\ &\geq \frac{n}{(n + \lambda)^{1-\lambda}} = n^\lambda \left(1 + \frac{\lambda}{n}\right)^{\lambda-1}. \end{aligned}$$

Der Quotient strebt nach 1 für $n \rightarrow \infty$.

- 2) Wegen $F(1) = 1$ gilt $F(n) = (n - 1)!$ für $n \in \mathbb{N}$. Für $\lambda \in (0, 1)$ haben wir

$$\begin{aligned} F(\lambda) &= \frac{F(\lambda + n)}{\lambda(\lambda + 1) \dots (\lambda + n - 1)} \leq \frac{(n - 1)! n^\lambda}{\lambda(\lambda + 1) \dots (\lambda + n - 1)} \\ F(\lambda) &= \frac{F(\lambda + n)}{\lambda(\lambda + 1) \dots (\lambda + n - 1)} \geq \frac{n! n^\lambda}{\lambda(\lambda + 1) \dots (\lambda + n)} \left(1 + \frac{\lambda}{n}\right)^\lambda \end{aligned}$$

Da der Quotient der oberen und der unteren Abschätzung nach 1 konvergiert, ergibt sich

$$F(\lambda) = \lim_{n \rightarrow \infty} \frac{n! n^\lambda}{\lambda(\lambda + 1) \dots (\lambda + n)}.$$

Es gibt nur ein $F(\cdot)$ mit den geforderten Eigenschaften.

Korollar Für alle $\lambda \in (0, \infty)$ gilt

$$\lim_{n \rightarrow \infty} \frac{n! n^\lambda}{\lambda(\lambda + 1) \dots (\lambda + n)} = \int_0^\infty x^{\lambda-1} e^{-x} dx = \Gamma(\lambda).$$

Wir benützen nun die Charakterisierung der Gammafunktion zum Studium von Eulers erstem Integral, der Betafunktion $B(\alpha, \beta)$ für $\alpha > 0$, $\beta > 0$.

Satz (Funktionalgleichung für die Betafunktion)

$$B(\alpha + 1, \beta) = \frac{\alpha}{\alpha + \beta} B(\alpha, \beta) \quad \text{für alle } \alpha, \beta > 0.$$

Beweis

$$\begin{aligned} B(\alpha + 1, \beta) &= \int_0^1 x^\alpha (1-x)^{\beta-1} dx \\ &= \int_0^1 \left(\frac{x}{1-x} \right)^\alpha (1-x)^{\alpha+\beta-1} dx \end{aligned}$$

Partielle Integration liefert wegen

$$\begin{aligned} \left(\frac{x}{1-x} \right)' &= \left(-1 + \frac{1}{1-x} \right)' = \frac{1}{(1-x)^2} \\ B(\alpha + 1, \beta) &= \int_0^1 \alpha \left(\frac{x}{1-x} \right)^{\alpha-1} \frac{1}{(1-x)^2} \frac{1}{\alpha + \beta} (1-x)^{\alpha+\beta} dx \\ &= \frac{\alpha}{\alpha + \beta} \int_0^1 x^{\alpha-1} (1-x)^{\beta-1} dx = \frac{\alpha}{\alpha + \beta} B(\alpha, \beta). \end{aligned}$$

Satz Für alle $\alpha, \beta > 0$ gilt

$$B(\alpha, \beta) = \frac{\Gamma(\alpha) \cdot \Gamma(\beta)}{\Gamma(\alpha + \beta)}.$$

Beweis

- 1) Mit der Methode mit welcher wir den 3. Beweis der logarithmischen Konvexität des zweiten Eulerschen Integrals geführt haben, ergibt sich auch die Konvexität von $\ln B(\alpha, \beta)$. Das Produkt zweier logarithmisch konvexer Funktionen ist logarithmisch konvex.

Für jedes feste $\beta > 0$ ist also $B(\alpha, \beta)$ und auch

$$F(\alpha) := \frac{1}{\Gamma(\beta)} \Gamma(\alpha + \beta) B(\alpha, \beta)$$

logarithmisch konvex.

- 2) Wir beweisen auf der Grundlage unserer Charakterisierung der Gammafunktion, daß $F(\cdot)$ die Gammafunktion ist.

Aus der Funktionsgleichung für die Betafunktion ergibt sich

$$\begin{aligned}\alpha F(\alpha) &= \frac{1}{\Gamma(\beta)} \cdot \Gamma(\alpha + \beta) \cdot (\alpha + \beta) B(\alpha + 1, \beta) \\ &= \frac{1}{\Gamma(\beta)} \cdot \Gamma(\alpha + 1 + \beta) B(\alpha + 1, \beta) = F(\alpha + 1) \\ B(1, \beta) &= \int_0^1 (1-x)^{\beta-1} dx = \left[-\frac{1}{\beta} (1-x)^\beta \right]_0^1 = \frac{1}{\beta} \\ F(1) &= \frac{1}{\Gamma(\beta)} \cdot \Gamma(1 + \beta) \frac{1}{\beta} = 1.\end{aligned}$$

Also ist $F(\cdot)$ die Gammafunktion.

Hinweis Die Zurückführung der Betafunktion auf die Gammafunktion ist ein eindrucksvolles Beispiel für die Nützlichkeit unserer Kennzeichnung der Gammafunktion, die auf keine spezielle Darstellung (wie Eulers Integral, die Formel von Gauß oder die von Weierstraß in den Übungen) Bezug nimmt.

Bemerke, daß die Funktionalgleichung

$$\lambda \cdot G(\lambda) = G(\lambda + 1), \quad G(1) = 1$$

viele Möglichkeiten zuläßt. Für jede in $(0, 1)$ beliebig vorgegebene Funktion $p(\lambda)$ genügt nämlich $G(\lambda) = p(\lambda)\Gamma(\lambda)$ der Funktionalgleichung.

Die Kennzeichnung durch die logarithmische Konvexität findet sich zum ersten Mal im Lehrbuch von H. BOHR und J. MØLLERUP (1922, dänisch). EMIL ARTIN hat den Gedanken populär gemacht in seinem Büchlein „Einführung in die Theorie der Gammafunktion“. Das Anliegen von Artin war es, zu zeigen, „daß die Gammafunktion in jeder Hinsicht zu den Elementarfunktionen gerechnet werden kann und daß für alle Eigenschaften dieser Funktion elementare, durchsichtige und dem Standpunkt eines Kollegs über Integralrechnung angepaßte Beweise gegeben werden können.“

Welche Eigenschaften der Gammafunktion so bedeutsam sind, daß sie eine Behandlung in der Anfängervorlesung verdienen, ist strittig. Das Studium der Gammafunktion eröffnet jedenfalls ein weites Feld für allerlei bezaubernden Umgang mit den Methoden der klassischen Analysis.

Wir erwähnen einige, die in Betracht kommen als

Merkwürdigkeiten

$$1) \Gamma\left(\frac{1}{2}\right) = \sqrt{\pi}$$

Der Beweis ist sehr einfach: Wegen $\Gamma(1) = 1$ gilt

$$\begin{aligned}\Gamma\left(\frac{1}{2}\right) \cdot \Gamma\left(\frac{1}{2}\right) &= B\left(\frac{1}{2}, \frac{1}{2}\right) \\ &= \int_0^1 \frac{1}{\sqrt{x(1-x)}} dx = 2 \int_0^1 \frac{1}{\sqrt{1-y^2}} dy = \pi.\end{aligned}$$

Als eine Konsequenz erhalten wir

$$\int_{-\infty}^{+\infty} \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{1}{2} x^2\right) dx = 1.$$

In der Tat haben wir

$$2 \int_0^{\infty} \exp\left(-\frac{1}{2} x^2\right) dx = \sqrt{2} \int_0^{\infty} \frac{1}{\sqrt{y}} e^{-y} dy = \sqrt{2} \Gamma\left(\frac{1}{2}\right) = \sqrt{2\pi}.$$

So erklärt sich der merkwürdige Normierungsfaktor für die „Gauß'sche Glockenkurve“.

- 2) Die Binomialkoeffizienten sind eng verwandt mit gewissen Werten der Betafunktion. Es gilt

$$\begin{aligned}[k B(k, n-k+1)]^{-1} &= \frac{1}{k} \frac{\Gamma(n+1)}{\Gamma(k)\Gamma(n-k+1)} \\ &= \frac{n!}{k!(n-k)!} = \binom{n}{k}.\end{aligned}$$

Für festes n steigen die Binomialkoeffizienten $\binom{n}{k}$ bis $k \leq \frac{n}{2}$ an; dann fallen sie wieder ab

$$\binom{n}{n-k} = \binom{n}{k}.$$

Eine ungefähre Auskunft über den größten Binomialkoeffizienten gibt die folgende Formel:

$$\binom{2n}{n} \frac{1}{2^{2n}} \sqrt{n} \longrightarrow \sqrt{\pi} \quad \text{für } n \rightarrow \infty.$$

- 3) Eine Approximation von $n!$, die in der elementaren Wahrscheinlichkeitstheorie hervorragende Dienste leistet und dort auch ausführlich diskutiert wird, ist die folgende Formel, die ursprünglich auf den Newtonschüler J. STIRLING (1692–1770) zurückgeht

$$n! = \sqrt{2\pi n} \left(\frac{n}{e}\right)^n \cdot \exp\left(S\left(\frac{1}{n}\right)\right)$$

mit einem Fehler $S(\cdot)$, welcher wie folgt abgeschätzt werden kann:

$$\frac{1}{12n} - \frac{1}{360} \cdot \frac{1}{n^3} \leq S\left(\frac{1}{n}\right) \leq \frac{1}{12n}.$$

Es ist umstritten, wie lehrreich die elementaren Beweise sind. Die professionellen Beweise benützen die sogenannten „Methoden von Laplace“, die wir hier aber nicht behandeln wollen. Alternative Formulierungen der Stirlingschen Formel sind z.B.:

$$\begin{aligned} \ln \Gamma(\alpha + 1) &= \alpha \ln \alpha - \alpha + \frac{1}{2} \ln \alpha + \ln \sqrt{2\pi} + \frac{1}{12\alpha} - \frac{1}{360\alpha^3} + \dots \\ &= \left(\alpha + \frac{1}{2}\right) \ln \left(\alpha + \frac{1}{2}\right) - \left(\alpha + \frac{1}{2}\right) + \ln \sqrt{2\pi} - \frac{1}{24\left(\alpha + \frac{1}{2}\right)} + \dots \end{aligned}$$

Die Approximation ist schon für recht kleine α sehr gut.

- 4) Mit mäßiger Mühe findet man die Abschätzungen

$$\sqrt{\frac{\lambda}{2} - \frac{1}{4}} \cdot \Gamma\left(\frac{\lambda}{2}\right) < \Gamma\left(\frac{\lambda+1}{2}\right) < \Gamma\left(\frac{\lambda}{2}\right) \cdot \sqrt{\left(\frac{\lambda}{2} - \frac{1}{4}\right) + \frac{1}{4(4\lambda-2)}}$$

Hinweis zum Beweis der unteren Abschätzung: Betrachte

$$c(\lambda) = \frac{1}{\sqrt{\frac{\lambda}{2} - \frac{1}{4}}} \cdot \frac{\Gamma\left(\frac{\lambda+1}{2}\right)}{\Gamma\left(\frac{\lambda}{2}\right)}$$

Bemerke $c(\lambda) \cdot c(\lambda+1) > 1$ und beweise

$$\begin{aligned} \frac{c(\lambda+1)}{c(\lambda)} &< 1 \text{ (wegen der logarithmischen Konvexität)} \\ \lim_{n \rightarrow \infty} c(\lambda+n) &= 1. \end{aligned}$$

- 5) Eine Identität, deren Bedeutung für die Anwendungen weniger offensichtlich ist, ist die „Verdoppelungsrelation“ von LEGENDRE (1752–1833). Sie besagt

$$\Gamma\left(\frac{\lambda}{2}\right) \cdot \Gamma\left(\frac{\lambda+1}{2}\right) = \frac{\sqrt{\pi}}{2^{\lambda-1}} \cdot \Gamma(\lambda).$$

Den Beweis kann man z.B. als eine Übungsaufgabe zur Gaußschen Darstellung oder auch auf der Grundlage der Charakterisierung der Gammafunktion führen. Mit derselben Methode findet man für $p = 3, 4, 5, \dots$

$$\Gamma\left(\frac{\lambda}{p}\right) \cdot \Gamma\left(\frac{\lambda+1}{p}\right) \cdot \dots \cdot \Gamma\left(\frac{\lambda+p-1}{p}\right) = \frac{(2\pi)^{\frac{p-1}{2}}}{p^{\lambda-\frac{1}{2}}} \cdot \Gamma(\lambda).$$

- 6) Für die in $\mathbb{C} \setminus \{0, -1, -2, \dots\}$ definierte Gammafunktion findet man eine merkwürdige Beziehung zur komplexen Sinusfunktion

$$\Gamma(z) \cdot \Gamma(1-z) = \frac{\pi}{\sin \pi z} = \prod_{k=1}^{\infty} \left(1 - \frac{z^2}{k^2}\right)^{-1}.$$

Artin gibt einen elementaren Beweis. Wirklich transparent wird die Situation aber wohl nur in der Weierstraßschen Theorie der Produktdarstellungen holomorpher Funktionen.

A.8 Die Fehlerfunktion, Kettenbrüche, eine Riccati-Gleichung, Eulers Behandlung des Wallisschen Produkts

Die Gaußsche Glockenkurve ist so berühmt, daß man sie für wert befunden hat, auf jedem 10 DM-Schein in Formel und Bild zu erscheinen. Die womöglich noch wichtigere Stammfunktion, **Fehlerfunktion** genannt, ist keine elementare Funktion im üblichen Sinn. Die Statistiker entnehmen ihre Werte bis heute aus Tabellen, was als ein Anachronismus erscheint, wenn man die Rechenkraft der Taschenrechner bedenkt und das (seit den Tagen von Lagrange und Euler verfügbare) theoretische Wissen über diese spezielle Funktion. Die üblichen Bezeichnungen sind

$$\begin{aligned}\varphi(x) &= \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}x^2\right) && \text{ („Glockenkurve“)} \\ \Phi(x) &= \int_{-\infty}^x \varphi(y) dy && \text{ („Fehlerfunktion“)}\end{aligned}$$

Offenbar gilt $\Phi(-x) = 1 - \Phi(x) = \int_x^{\infty} \varphi(y) dy$. $2\Phi(-x)$ ist die Wahrscheinlichkeit, daß eine normalverteilte Zufallsgröße um mehr als x ($x > 0$) Standardabweichungen von ihrem Mittelwert abweicht. Beispielsweise muß man mit einer Abweichung um mehr als 2 Standardabweichungen mit der Wahrscheinlichkeit 5% rechnen. (Den genaueren Wert $\Phi(-1.960) = 0.025$ entnimmt man Tabellen (oder dem Taschenrechner), wie man sie in den meisten Büchern zur elementaren Statistik findet.)

Das Studium von $\Phi(\cdot)$ wird uns Gelegenheit geben, einige Techniken zu skizzieren, die die Mathematiker des 18. Jahrhunderts entwickelt haben, um konkrete Funktionen darzustellen und zu behandeln. Insbesondere wollen wir am Beispiel einige grundlegende Tatsachen über **Kettenbrüche** vorstellen. Diejenigen, die die Theorie von Grund auf erlernen wollen, verweisen wir auf die beiden Standardwerke:

O. PERRON: „Lehre von den Kettenbrüchen“, Teubner 1913

H.S. WALL: „Analytical Theory of Continued Fractions“, Chelsea 1948

Als Ausgangspunkt unserer theoretischen Überlegungen zur Fehlerfunktion $\Phi(\cdot)$ soll uns ein Kettenbruch dienen, der bereits Lagrange bekannt war (bei Perron S. 472; bei Wall S. 358, Formel 92.15)

$$\frac{\varphi(x)}{\Phi(-x)} = x + \frac{1}{|x|} + \frac{2}{|x|} + \frac{3}{|x|} + \dots \quad \text{für } x > 0$$

Der unendliche Kettenbruch konvergiert nur für große x so gut, daß man ihn für effiziente **Approximationen** von $\Phi(-x)$ benutzen kann.

Zur numerischen Berechnung siehe

ABRAMOWITZ und STEGUN: "Handbook of Mathematical Functions"
(7. Error Function & Fresnel Integral, 26. Probability Functions)

Ad-hoc-Approximationen für $\Phi(x)$ oder $\frac{\varphi(x)}{\Phi(-x)}$ („Mills' ratio“) findet man in

KENDALL-STUART: The Advanced Theory of Statistics, Vol. 1, p. 137 ff

Für kleine positive x benutzt man lieber die endlichen Kettenbrüche

$$\begin{aligned}\frac{\varphi(x)}{\Phi(-x)} &= x + r_1(x) = x + \frac{1}{x + r_2(x)} \\ &= x + \frac{1}{x + \frac{2}{x + r_3(x)}} = \dots\end{aligned}$$

zusammen mit einer Approximation des Restterms $r_\lambda(x)$.

Mit diesen Resttermen und ihrer zu einer Kettenbruchentwicklung führenden Beziehung

$$r_\lambda(x) = \frac{\lambda}{x + r_{\lambda+1}(x)}$$

werden wir uns ausführlich beschäftigen.

Ein wesentliches Hilfsmittel wird uns dabei die Theorie der **Riccati-Gleichungen** sein. Dabei befinden wir uns auf den Spuren von Euler. Euler hat (wie Perron in § 80 seiner Monographie ausführt) in den meisten seiner Arbeiten, die von Kettenbrüchen handeln, einen Zusammenhang mit der Riccatischen Differentialgleichung hergestellt. Allerdings hat Euler das uns hier interessierende Thema nicht mit der Strenge behandelt, zu der wir uns heute verpflichtet fühlen müssen.

Interessant ist auch der Grenzfall $x = 0$. Wir haben

$$\begin{aligned}\sqrt{\frac{2}{\pi}} &= \frac{\varphi(0)}{\Phi(0)} = r_1(0) = \frac{1}{r_2(0)} = \frac{1 \cdot r_3(0)}{2} = \frac{1 \cdot 3}{2 \cdot r_4(0)} = \dots \\ &= \frac{1 \cdot 3 \cdot 5 \cdot \dots \cdot (2m-1)}{2 \cdot 4 \cdot 6 \cdot \dots \cdot (2m-2) \cdot r_{2m}(0)} \\ &= \frac{1 \cdot 3 \cdot \dots \cdot (2m-1) \cdot r_{2m+1}(0)}{2 \cdot 4 \cdot \dots \cdot (2m-2) \cdot 2m}\end{aligned}$$

$$r_{2m}(0) \cdot r_{2m+1}(0) = 2m, \quad r_{2m+1}(0) \cdot r_{2m+2}(0) = 2m+1$$

Die Konvergenz des Wallisschen Produkts ist äquivalent mit der Tatsache $r_\lambda(0) \approx \sqrt{\lambda}$ für große natürliche λ .

Mit Hilfe des „Eulerschen Integrals“ (Gamma-Funktion) können wir $r_\lambda(0)$ auch explizit angeben

$$r_\lambda(0) = \sqrt{2} \cdot \frac{\Gamma(\frac{\lambda+1}{2})}{\Gamma(\frac{\lambda}{2})}; \quad 2r_\lambda^2(0) = 4 \left(\frac{\Gamma(\frac{\lambda+1}{2})}{\Gamma(\frac{\lambda}{2})} \right)^2$$

(Zunächst sind natürlich nur die $r_\lambda(0)$ mit natürlichem λ vorgegeben. Die Konstruktion von $r_\lambda(\cdot)$ für beliebige reelle $\lambda > 0$ ist noch zu leisten.)

Wir werden unten (auf den Spuren von Euler) die Kettenbruchdarstellung herleiten

$$2r_\lambda^2(0) = (2\lambda - 1) + \frac{1}{|4\lambda - 2|} + \frac{9}{|4\lambda - 2|} + \frac{25}{|4\lambda - 2|} + \dots \quad \text{für } \lambda > \frac{1}{2}.$$

Auch hier gilt, daß man für kleine λ nicht die „Näherungsbrüche“ des unendlichen Kettenbruchs zur effizienten Approximation der Zahlenwerte $2r_\lambda^2(0)$ (wie z.B. $2r_1^2(0) = \frac{4}{\pi} = 1 + \frac{1}{2} + \frac{9}{2} + \frac{25}{2} + \frac{49}{2} + \dots$) benützen wird, sondern lieber die endlichen Kettenbrüche

$$4 \left(\frac{\Gamma\left(\frac{\lambda+1}{2}\right)}{\Gamma\left(\frac{\lambda}{2}\right)} \right)^2 = 2\lambda - 1 + \frac{4\left(\frac{1}{2}\right)^2}{|4\lambda - 2|} + \frac{4\left(\frac{3}{2}\right)^2}{|4\lambda - 2|} + \dots + \frac{4\left(\frac{1}{2} + m - 1\right)^2}{|4\lambda - 2 - (2\lambda - 1) + W\left(\frac{1}{2} + m, \lambda\right)|}$$

mit einer Abschätzung der Reste $W(\eta, \lambda)$, für welche gilt

$$[W(\eta, \lambda) - (2\lambda - 1)][W(\eta + 1, \lambda) + (2\lambda - 1)] = 4\eta^2 \quad \text{für alle } \lambda.$$

(Für große η, λ kommt die folgende Approximation in Betracht

$$W(\eta, \lambda) \approx \sqrt{(2\eta - 1)^2 + (2\lambda - 1)^2}.)$$

Dazu müssen wir aber weiter ausholen. Es handelt sich um Themen, die Euler so fasziniert haben, daß er mehrere Arbeiten darüber verfaßt hat. Der Herausgeber von Eulers Gesammelten Werken, Georg Faber, meinte 1934, daß zur vollständigen Aufklärung der Sachverhalte eine gründliche Kenntnis der Theorie der Kettenbrüche nötig sei (siehe Vorwort zur Serie prima, Band 16/2). Wir werden aber sehen, daß man alles mit elementaren Mitteln erledigen kann.

I) Kettenbrüche

Im 17. Jahrhundert, als die Dezimalbruchentwicklungen noch nicht die Alleinherrschaft bei der Zahlendarstellung angetreten hatten, wurden die (regelmäßigen) Kettenbrüche als ein Verfahren konzipiert, um Irrationalzahlen durch Rationalzahlen mit möglichst kleinen Nennern zu approximieren.

Wichtigen Anteil an der Entwicklung hatte CHRISTIAAN HUYGENS (1629–95), der nicht nur eine bedeutende Abhandlung über Wahrscheinlichkeitstheorie schrieb (1657), den Saturnring erklärte (1655), eine Theorie der Bewegung der Körper durch den Stoß und eine Wellentheorie des Lichts entwickelte (1690), sondern auch die Pendeluhr und die Federuhr mit Unruh erfand. Das Werk „*Horologium ocillatorium*“ (1673) über die Pendeluhr mit einer mathematischen Theorie des Pendels gilt sogar als sein Hauptwerk. In den Uhren mußten vorgegebene Verhältnisse von Drehgeschwindigkeiten durch Zahnräder mit nicht allzuvielen Zähnen realisiert werden. Die Berechnungen wurden auf Kettenbrüche gestützt.

Das Verfahren, eine Zahl $x > 0$ in einen **regelmäßigen** Kettenbruch zu entwickeln, ist das folgende: Finde zu $x_0 = x$ die größte ganze Zahl darunter, $b_0 := [x_0]$ und setze x_1 gleich dem Kehrwert des Rests

$$x_1 = \frac{1}{x_0 - b_0} (> 1); \quad x_0 = b_0 + \frac{1}{x_1}.$$

Konstruiere nun zu der Zahl $x_1 > 1$ in derselben Weise den „Teilnenner“ b_1 und die Zahl $x_2 > 1$

$$b_1 = [x_1] \in \mathbb{N}, \quad x_2 = \frac{1}{x_1 - b_1} > 1, \quad x_0 = b_0 + \frac{1}{b_1 + \frac{1}{x_2}}.$$

Genau dann, wenn $x = x_0$ rational ist, bricht der Kettenbruch ab. (Man beachte die Analogie zum Euklidischen Algorithmus.) Es sind zweierlei Schreibweisen in Gebrauch

$$\begin{aligned} x_0 &= b_0 + \frac{1|}{|b_1|} + \frac{1|}{|b_2|} + \dots + \frac{1|}{|b_{n-1} + \frac{1}{x_n}|} \\ &= (b_0; 1, b_1; 1, \dots; 1, b_{n-1} + \frac{1}{x_n}; 0, \dots) \end{aligned}$$

Beispiel Der reguläre Kettenbruch zur Zahl π lautet

$$\begin{aligned} \pi &= 3 + \frac{1|}{|7|} + \frac{1|}{|15|} + \frac{1|}{|1|} + \frac{1|}{|292|} + \frac{1|}{|1|} + \dots \\ &= (3; 1, 7; 1, 15; 1, 1; 1, 292; 1, 1; \dots) \end{aligned}$$

Die Approximation durch den ersten Näherungsbruch $\pi \approx 3 + \frac{1}{7} = \frac{22}{7}$ ist bereits von Archimedes benutzt worden. Der dritte Näherungsbruch

$$\pi \approx 3 + \frac{1}{7 + \frac{1}{15 + \frac{1}{1}}} = 3 + \frac{1}{7 + \frac{1}{16}} = \frac{355}{113}$$

geht auf METIUS (1527–1607) zurück. π ist in Wirklichkeit etwas kleiner. Daß der Näherungsbruch äußerst genau ist, zeigt sich in der Größe des vierten Teilnenners $b_4 = 292$. Der Fehler ist kleiner als $\frac{1}{113 \cdot 33102}$, wie man mit den unten bereitgestellten Mitteln leicht nachrechnen kann.

Als **Verallgemeinerung** der regelmäßigen Kettenbrüche betrachtet man Zahldarstellungen der Art

$$\begin{aligned} x &= b_0 + \frac{a_1|}{|b_1|} + \frac{a_2|}{|b_2|} + \dots + \frac{a_n|}{|b_n + r_{n+1}|} \\ &= (b_0; a_1, b_1; \dots; a_n, b_n + r_{n+1}; 0, \dots) \end{aligned}$$

Man spricht von einem endlichen Kettenbruch der Tiefe n .

Ein **unendlicher Kettenbruch** ist ein formaler Ausdruck der Art

$$b_0 + \frac{a_1|}{|b_1|} + \frac{a_2|}{|b_2|} + \frac{a_3|}{|b_3|} + \dots,$$

auch geschrieben $(b_0; a_1, b_1; a_2, b_2; \dots)$.

Zu einem solchen formalen Ausdruck definiert man rekursiv die Folgen $(A_n)_n, (B_n)_n$ wie folgt

$$\begin{aligned} A_0 &= b_0; & B_0 &= 1; \\ A_1 &= b_1 A_0 + a_1 \cdot 1; & B_1 &= b_1 B_0 + a_1 \cdot 0; \\ A_2 &= b_2 A_1 + a_2 \cdot A_0; & B_2 &= b_2 B_1 + a_2 \cdot B_0; \\ &\dots & & \\ A_{n+1} &= b_{n+1} A_n + a_{n+1} A_{n-1}; & B_{n+1} &= b_{n+1} B_n + a_{n+1} B_{n-1}. \end{aligned}$$

Die Quotienten $\frac{A_k}{B_k}$ heißen die Näherungsbrüche des unendlichen Kettenbruchs. Offenbar gilt

$$\frac{A_0}{B_0} = b_0, \quad \frac{A_1}{B_1} = b_0 + \frac{a_1}{b_1}, \quad \frac{A_2}{B_2} = b_0 + \frac{a_1}{b_1 + \frac{a_2}{b_2}}.$$

Wir werden sehen, daß für alle n

$$\begin{aligned} (*) \quad \frac{A_n}{B_n} &= b_0 + \frac{a_1}{|b_1|} + \frac{a_2}{|b_2|} + \dots + \frac{a_n}{|b_n| + 0} \\ &= (b_0; a_1, b_1; \dots; a_n, b_n; 0, \dots) \end{aligned}$$

Uns interessiert hier nur der Fall $a_\nu \geq 0, b_\nu > 0$ für alle ν . Beachte: Wenn $a_{n+1} = 0$, dann spielen die weiteren a_ν, b_ν für die Berechnung der weiteren $\frac{A_k}{B_k}$ keine Rolle; wir wollen einen solchen Kettenbruch mit einem endlichen Kettenbruch der Tiefe n identifizieren.

Lemma Die folgenden Kettenbrüche haben dieselben Näherungsbrüche

$$\begin{aligned} &(b_0; a_1, b_1; \dots; a_n, b_n; a_{n+1}, b_{n+1}; 0, \dots) \\ &\approx (b_0; a_1, b_1; \dots; a_n, b_n + \frac{a_{n+1}}{b_{n+1}}; 0, \dots) \end{aligned}$$

Beweis Für $k < n$ ist die Aussage trivial. Für den zweiten Kettenbruch ergibt sich der n -te Näherungsbruch aus dem A_k, B_k des ersten Kettenbruchs wie folgt

$$\begin{aligned} A_n^* &= \left(b_n + \frac{a_{n+1}}{b_{n+1}} \right) A_{n-1} + a_n A_{n-2} = A_n + \frac{a_{n+1}}{b_{n+1}} A_{n-1} = \frac{1}{b_{n+1}} A_{n+1}, \\ B_n^* &= \left(b_n + \frac{a_{n+1}}{b_{n+1}} \right) B_{n-1} + a_n B_{n-2} = B_n + \frac{a_{n+1}}{b_{n+1}} B_{n-1} = \frac{1}{b_{n+1}} B_{n+1}. \end{aligned}$$

Also $\frac{A_n^*}{B_n^*} = \frac{A_{n+1}}{B_{n+1}}$; und alle weiteren Näherungsbrüche sind gleich diesem ultimativen Wert. Und das beweist (*).

Lemma Wenn die a_ν, b_ν allesamt positiv sind, dann gilt

$$\frac{A_0}{B_0} < \frac{A_2}{B_2} < \frac{A_4}{B_4} < \dots < \frac{A_5}{B_5} < \frac{A_3}{B_3} < \frac{A_1}{B_1}$$

mit

$$\left[\frac{A_{n+1}}{B_{n+1}} - \frac{A_n}{B_n} \right] = (-1) \left[1 - \frac{b_{n+1}B_n}{B_{n+1}} \right] \left[\frac{A_n}{B_n} - \frac{A_{n-1}}{B_{n-1}} \right].$$

Beweis $\frac{A_0}{B_0} < \frac{A_1}{B_1}$ und $\frac{A_1}{B_1} > \frac{A_2}{B_2}$ sind offensichtlich. Für alle n gilt

$$\begin{aligned} \frac{A_{n+1}}{B_{n+1}} - \frac{A_n}{B_n} &= \frac{1}{B_{n+1}B_n} [B_n(b_{n+1}A_n + a_{n+1}A_{n-1}) - (b_{n+1}B_n + a_{n+1}B_{n-1})A_n] \\ &= \frac{(-a_{n+1})B_{n-1}}{B_{n+1}} \left[\frac{A_n}{B_n} - \frac{A_{n-1}}{B_{n-1}} \right] = (-1) \left[1 - \frac{b_{n+1}B_n}{B_{n+1}} \right] \left[\frac{A_n}{B_n} - \frac{A_{n-1}}{B_{n-1}} \right] \end{aligned}$$

Sprechweise Man sagt von einem unendlichen Kettenbruch mit positiven Einträgen a_ν, b_ν , daß er **konvergiert**, wenn

$$\frac{A_{2n+1}}{B_{2n+1}} - \frac{A_{2n}}{B_{2n}} \rightarrow 0.$$

Der Limes der Näherungsbrüche heißt der **Wert** des unendlichen Kettenbruchs.

Bemerke Für einen konvergenten Kettenbruch (mit dem Wert w^*) gibt es eindeutig bestimmte $r_2 > 0, r_3 > 0, \dots$, so daß

$$w^* = b_0 + \frac{a_1}{b_1 + r_2} = b_0 + \frac{a_1}{|b_1|} + \frac{a_2}{|b_2 + r_3|} = \dots$$

Für diese Reste gilt

$$r_n(b_n + r_{n+1}) = a_n \quad \text{für alle } n.$$

Den n -ten Näherungsbruch erhält man, wenn man r_{n+1} durch 0 ersetzt. Es ist klar, daß eine bessere Approximation von r_{n+1} eine bessere Approximation von w^* liefert

$$w^* \approx b_0 + \frac{a_1}{|b_1|} + \dots + \frac{a_n}{|b_n + \tilde{r}_{n+1}|}.$$

Dabei ist zu beachten, daß die Abbildung

$$w \mapsto b_0 + \frac{a_1}{|b_1|} + \dots + \frac{a_n}{|b_n + w|}$$

auf der Halbachse $(-b_n, \infty)$ antiton ist für ungerades n und isoton für gerades n .

II) Verallgemeinerung des Kettenbruchs von Lagrange

Für alle $\lambda > 0$ studieren wir die Funktionen

$$M_\lambda(x) := \int_0^\infty y^{\lambda-1} \exp\left(-\frac{1}{2}(y+x)^2\right) dy$$

$$r_\lambda(x) := \frac{M_{\lambda+1}(x)}{M_\lambda(x)}$$

Satz A.8.1 (Funktionalgleichung)

$$\lambda \cdot M_\lambda(x) = x \cdot M_{\lambda+1}(x) + M_{\lambda+2}(x)$$

Beweis

$$\begin{aligned} \lambda \cdot M_\lambda(x) &= \int_0^\infty \lambda y^{\lambda-1} \exp\left(-\frac{1}{2}(y+x)^2\right) dy \\ &= \left[y^\lambda \exp\left(-\frac{1}{2}(y+x)^2\right) \right]_0^\infty + \int_0^\infty y^\lambda (y+x) \exp\left(-\frac{1}{2}(y+x)^2\right) dy \\ &= x M_{\lambda+1}(x) + M_{\lambda+2}(x) . \end{aligned}$$

Korollar

$$\frac{\lambda}{r_\lambda(x)} = x + r_{\lambda+1}(x) .$$

Beweis Dividiere die Funktionalgleichung durch $M_{\lambda+1}(x)$.

Bemerkung 1 Man kann $M_{\lambda+1}(\cdot)$ (bis auf eine multiplikative Konstante) aus $r_\lambda(x)$ zurückgewinnen. Es gilt nämlich

$$\begin{aligned} -\frac{d}{dx} M_{\lambda+1}(x) &= \int_0^\infty y^\lambda \exp\left(-\frac{1}{2}(y+x)^2\right) (y+x) dy \\ &= M_{\lambda+2}(x) + x M_{\lambda+1}(x) = \lambda M_\lambda(x) \\ -\frac{d}{dx} \ln M_{\lambda+1}(x) &= \lambda r_\lambda(x)^{-1} \\ M_{\lambda+1}(x) &= M_{\lambda+1}(x_0) \exp\left(-\lambda \int_{x_0}^x r_\lambda(y)^{-1} dy\right) \end{aligned}$$

Satz A.8.2 (Riccati–Gleichung)

Für alle $\lambda \geq 1$ löst $r_\lambda(\cdot)$ die Differentialgleichung

$$(R_\lambda): \quad r'_\lambda(x) = r_\lambda^2(x) + x r_\lambda(x) - \lambda .$$

Beweis

$$\begin{aligned}
r'_\lambda(x) &= \left(\frac{M_{\lambda+1}(x)}{M_\lambda(x)} \right)' = \frac{-M'_\lambda}{M_\lambda^2} M_{\lambda+1} + \frac{1}{M_\lambda} M'_{\lambda+1} \\
&= \frac{xM_\lambda + M_{\lambda+1}}{M_\lambda^2} M_{\lambda+1} - \lambda \\
&= xr_\lambda(x) + r_\lambda^2(x) - \lambda.
\end{aligned}$$

Bemerkung 2 Die Funktionen $M_\lambda(x)$ liefern ein Bindeglied zwischen der Gamma-Funktion ($x = 0$) und der Gaußschen Fehlerfunktion ($\lambda = 1$). Es gelten

$$\begin{aligned}
\text{(i)} \quad \frac{1}{\sqrt{2\pi}} M_1(x) &= \int_0^\infty \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}(y+x)^2\right) dy \\
&= \int_x^\infty \varphi(v) dv = \Phi(-x) \\
\frac{1}{\sqrt{2\pi}} M_2(x) &= \int_x^\infty \frac{1}{\sqrt{2\pi}} (v-x) \exp\left(-\frac{1}{2}v^2\right) dv \\
&= \varphi(x) - x\Phi(-x) \\
r_1(x) &= \frac{M_2(x)}{M_1(x)} = \frac{\varphi(x)}{\Phi(-x)} - x \\
\text{(ii)} \quad M_\lambda(0) &= \int_0^\infty y^{\lambda-1} \exp\left(-\frac{1}{2}y^2\right) dy \\
&= \int_0^\infty u^{\lambda/2-1} \exp(-u) du \cdot 2^{\lambda/2-1} = \Gamma\left(\frac{\lambda}{2}\right) \cdot 2^{\lambda/2-1} \\
r_\lambda(0) &= \sqrt{2} \cdot \frac{\Gamma\left(\frac{\lambda+1}{2}\right)}{\Gamma\left(\frac{\lambda}{2}\right)} \\
2r'_\lambda(0) + 1 &= 2r_\lambda^2(0) - (2\lambda - 1) = 4 \left(\frac{\Gamma\left(\frac{\lambda+1}{2}\right)}{\Gamma\left(\frac{\lambda}{2}\right)} \right)^2 - (2\lambda - 1)
\end{aligned}$$

Korollar Die Funktionen $r_m(x)$ sind die Reste in den Kettenbruchdarstellungen

$$\frac{\varphi(x)}{\Phi(-x)} = x + r_1(x) = x + \frac{1}{x + r_2(x)} = x + \frac{1}{x + \frac{2}{x + r_3(x)}} = \dots$$

Als Approximation für kleine $x > 0$ empfehlen sich aber nicht die Näherungsbrüche des unendlichen Kettenbruchs von Lagrange, sondern vielmehr die endlichen Kettenbrüche

$$\frac{\varphi(x)}{\Phi(-x)} \approx x + \frac{1}{|x|} + \frac{2}{|x|} + \dots + \frac{(m-1)}{|x + \tilde{r}_m(x)|}$$

mit einer passenden Approximation $\tilde{r}_m(x) \approx r_m(x)$. In Betracht kommt z.B.

$$r_\lambda(x) = -\frac{x}{2} + \sqrt{\frac{x^2}{4} + 2 \frac{\Gamma^2(\frac{\lambda+1}{2})}{\Gamma^2(\frac{\lambda}{2})}} \approx -\frac{x}{2} + \sqrt{\frac{x^2}{4} + \left(\lambda - \frac{1}{2}\right)} \quad \text{für große } \lambda.$$

Entscheidend für eine gute Approximation von

$$\frac{\varphi(x)}{\Phi(-x)} = x + \frac{1}{|x|} + \frac{2}{|x|} + \dots + \frac{m-1}{|x + r_m(x)|}$$

für x nahe bei 0 ist offenbar eine gute Approximation von

$$h_\lambda(0) = 4 \left(\frac{\Gamma(\frac{\lambda+1}{2})}{\Gamma(\frac{\lambda}{2})} \right)^2 - (2\lambda - 1).$$

Diese Approximationsaufgabe führt auf den Brounckerschen Kettenbruch.

Die Art, wie die Analytiker des 18. Jahrhunderts operierten, ist im 19. Jahrhundert in Frage gestellt worden. Es wird aber heute weithin als Defizit in der Ausbildung angesehen, wenn die Studenten der Mathematik mit der naßforschenden Analysis der älteren Zeit gar nicht mehr in Berührung gebracht werden. Laplace pflegte zu jüngeren Mathematikern zu sagen: „*Lest Euler, er ist unser aller Meister.*“ Gauß: „*Das Studium der Werke Eulers bleibt die beste Schule in den verschiedenen Gebieten der Mathematik und kann durch nichts anderes ersetzt werden.*“ (zitiert nach Struik S. 135) Der Anhang stellt einen Versuch dar, den Geist der Analysis des 18. Jahrhunderts an einem Beispiel sichtbar werden zu lassen, ohne auf die heute zu fordernde Strenge zu verzichten.

G. FABER schreibt im Vorwort zu Eulers Gesammelten Werken, Series prima, 16/2, Seite XII:

„Eulers unermüdliche Freude am Zahlenrechnen war stets beherrscht von dem Gedanken, die Tragweite solcher Verfahren darzutun, die geeignet waren, die Arbeit des Rechnens durch die Überlegung des Mathematikers abzukürzen. Rechenfehler größeren Ausmaßes sind ihm selten unterlaufen; doch befolgt er nicht die strenge, erst später zur allgemeinen Geltung gelangte Regel, daß der Fehler nicht mehr als eine halbe Einheit der letzten noch beibehaltenen Stelle betragen darf. Bei den halbkonvergenten Reihen, die er als erster ausgiebig benutzte, und die er auch in seiner Wesensart richtig erkannte, findet er mit sicherem Takt die Stelle, an der am günstigsten mit der Summation aufgehört wird. Dabei wird er sich wenigstens an Beispielen klar, daß es sich um divergente Reihen handelt. Aber auch bei solchen divergenten Reihen, die keineswegs zur Klasse der semikonvergenten gehören, kam nach seiner Auffassung, die sich freilich schon bei einem Teil seiner Zeitgenossen und erst recht später nicht durchsetzen konnte, ein bestimmter Summenwert zu, den er durch folgende Festsetzung sicher stellen zu können glaubte: Die Summe jeder unendlichen Reihe ist der Wert des endlichen Ausdrucks,

durch dessen Entwicklung die Reihe entsteht. Diese Auffassung hat mit gehöriger Einschränkung hinterher durch den Weierstraßschen Begriff der analytischen Fortsetzung einen bestimmten mathematischen Sinn und damit eine gewisse Berechtigung erlangt. Nachdem man noch vor wenigen Jahrzehnten die Eulerschen mit divergenten Reihen arbeitenden Ansätze allgemein als rettungslos verfehlt angesehen hatte, braucht man heute nur manche seiner Entwicklungen etwas anders zu begründen, um ihnen einen guten Sinn zu geben. ...“

Allgemeine Bemerkungen zum Begriff „Approximation“

- 1) Aus der Sicht eines reinen Mathematikers unserer Tage kann jede Folge, die gegen die Zahl a^* konvergiert, als Approximation dieser Zahl gelten. Beispiele sind etwa

$$e = \lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n = 1 + \frac{1}{1!} + \frac{1}{2!} + \frac{1}{3!} + \dots$$

$$\frac{\pi}{4} = \text{Limes der Partialsummen von } \left(1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} \pm \dots\right)$$

Eine Funktionenfolge, die gleichmäßig oder in irgendeiner passenden Norm gegen die Funktion $f^*(\cdot)$ konvergiert, kann (in abstrakter Sicht) zur Approximation von $f^*(\cdot)$ herangezogen werden. So können z.B. Potenzreihen, die bekanntlich in jedem Kompaktum innerhalb des Konvergenzkreises gleichmäßig konvergieren, zur Approximation der Grenzfunktion herangezogen werden.

Die Approximation von Funktionen durch die Näherungsbrüche von Kettenbrüchen, die Euler in seinem Lehrbuch „Introductio ad analysim infiniterum“ den Partialsummen von Reihen mehr oder weniger ebenbürtig zur Seite gestellt hat, führt heute eher ein Schattendasein.

- 2) In den Anfängervorlesungen verschweigt man üblicherweise, daß es auch andere (nicht auf den Konvergenzbegriff gestützte) Vorstellungen von (guter und effizienter) Approximation gibt.

Beispiel Die Stirlingsche Formel schreibt man manchmal

$$n! = \sqrt{2\pi n} \cdot \left(\frac{n}{e}\right)^n \cdot \exp\left(\frac{1}{12n} - \frac{1}{360n^3} \pm \dots\right)$$

und man gibt vielleicht sogar eine Formel an für die weiteren Koeffizienten der Reihe

$$S\left(\frac{1}{n}\right) = \ln n! - \left[n \ln n - n + \frac{1}{2} \ln(2\pi n)\right]$$

$$= \frac{1}{12} \left(\frac{1}{n}\right) - \frac{1}{360} \left(\frac{1}{n}\right)^3 \pm \dots$$

Hier handelt es sich aber nicht um eine (in irgendeinem Bereich der Argumente) konvergente Reihe. Es handelt sich vielmehr um eine sog. asymptotische Entwicklung des Restterms $S\left(\frac{1}{n}\right)$.

Die Approximation durch die Partialsummen der Reihe ist nützlich, obwohl man für festes n die Approximation von $S\left(\frac{1}{n}\right)$ durch Hinzunahme weiterer Summanden nicht beliebig genau machen kann. Man beweist: Der Unterschied zwischen $S\left(\frac{1}{n}\right)$ und einer Partialsumme mit fester Länge strebt schneller nach 0 für $n \rightarrow \infty$ als alle Terme in der Partialsumme.

- 3) Der Anfänger muß vor dem Denkfehler gewarnt werden, die Taylor-Approximation einer n -mal stetig differenzierbaren Funktion generell für eine abgebrochene Potenzreihen-Entwicklung zu halten.

$$f(x) = f(x_0) + f'(x_0)(x - x_0) + \dots + \frac{1}{n!} f^{(n)}(x_0)(x - x_0)^n + o(|x - x_0|^n)$$

Der Approximationsfehler kann im allgemeinen für festes $x \neq x_0$ nicht beliebig klein gemacht werden. Man beweist nur: Der Unterschied zwischen $f(x)$ und dem Taylor-Polynom strebt für $x \rightarrow x_0$ schneller nach 0 als alle Terme des Polynoms. Der allgemeine Fall ähnelt also mehr der in 2) skizzierten Situation als der in 1) beschriebenen.

- 4) Praktiker sind interessiert an „guten“ Approximationen. Sie wollen mit **geringem Rechenaufwand** die gemeinte Zahl a^* mit **hinreichender Genauigkeit** approximieren. Limesprozesse sind als solche ohne praktisches Interesse; asymptotische Aussagen haben nur dann praktisches Interesse, wenn sie „gute“ Approximationen liefern. „Hinreichende Genauigkeit“ ist kein mathematischer Begriff. Eine Garantie, daß eine gewisse (einigermaßen realistisch geschätzte) Genauigkeit eingehalten wird, ist bei vielen den Praktikern geläufigen Approximationen schwer theoretisch zu untermauern. Praktiker verlassen sich lieber auf die Erfahrung. Konvergente Folgen und asymptotische Entwicklungen werden genau dann hochgeschätzt, wenn sie Approximationen suggerieren, die sich „bewähren“.
- 5) Nicht alle von Praktikern verwendeten Approximationen gründen auf einer nur den Einzelfall betreffenden Theorie. Viele sind Anwendungsfälle einer allgemeinen Approximationstheorie; die Koeffizienten, die die Anwendung auf den konkreten Fall regulieren, müssen dann aus Tabellen entnommen werden. — Wir denken etwa an die Theorie der Approximation durch Tschebyschev-Polynome. Bei Abramowitz & Stegun finden sich mehrere Vorschläge dieser Art für die Approximation der Fehlerfunktion $\Phi(\cdot)$.
- 6) Wir wollen rekapitulieren, was wir über die Approximation der Zahl π wissen.

- a) In vielen Formelsammlungen findet man

$$\pi = 3,141592653\dots$$

Ähnlichen Charakter hat die Information

$$\pi = 3 + \frac{1}{7} + \frac{1}{15} + \frac{1}{1} + \frac{1}{292} + \frac{1}{1} + \frac{1}{1} + \frac{1}{2} + \dots$$

Für die zweite Approximation muß man sich etwas weniger Bits Informationen für die gleiche Genauigkeit merken. Leider gibt es anscheinend keinen einfachen Algorithmus, um sich die Dezimalstellen oder die Teilnenner des regelmäßigen Kettenbruchs zu besorgen.

- b) Wir hatten $\frac{\pi}{2}$ als die kleinste positive Nullstelle der durch die Potenzreihe gegebenen Cosinusfunktion „definiert“. Die Suche nach der kleinsten positiven Nullstelle der Gleichung

$$1 - \frac{x^2}{2!} + \frac{x^4}{4!} - \frac{x^6}{6!} + \dots = 0$$

kann als ein Anwendungsfall von allgemeineren Verfahren zur numerischen Lösung von Gleichungen behandelt werden. In der „Numerischen Mathematik“ lernt man die Techniken und die Abschätzung der Fehler. Das für den Anwendungsfall $\frac{\pi}{2}$ Spezifische sind nur die (nun wirklich leicht zu merkenden) Koeffizienten der Cosinusreihe.

- c) In einer Übungsaufgabe haben wir gelernt, wie man mit Hilfe des Additionstheorems für den Tangens das Problem der Berechnung von $\frac{\pi}{4} = \arctan 1$ so umformen kann, daß Potenzreihenentwicklung (oder auch die Kettenbruchdarstellung) des Arcustangens

$$\begin{aligned} \arctan x &= \frac{x}{|1|} + \frac{x^2}{|3|} + \frac{4x^2}{|5|} + \frac{9x^2}{|7|} + \dots \\ &= x - \frac{1}{3}x^3 + \frac{1}{5}x^5 - \frac{1}{7}x^7 + \dots \end{aligned}$$

tatsächlich effizient benützt werden kann.

- d) Euler war fasziniert von dem Brounckerschen Kettenbruch

$$\frac{4}{\pi} = 1 + \frac{1}{|2|} + \frac{9}{|2|} + \frac{25}{|2|} + \dots$$

Dieser Kettenbruch und die Verallgemeinerungen sind aber wohl weniger als Vorschläge zur Approximation gedacht gewesen, als vielmehr als Ausdruck von theoretisch interessanten Zusammenhängen.

Zum Abschluß unseres Einblicks in die konkrete Analysis des 18. Jahrhunderts zitieren wir aus dem bereits erwähnten Artikel

PHILIP DAVIS : *Leonard Euler's Integral*,
abgedruckt in den Chauvenet Papers, ed. Abbot (S. 345).

„Functions are the building blocks of mathematical analysis. In the 18th and 19th century mathematicians devoted much time and loving care to developing the properties and interrelationships between special functions. Powers, roots, algebraic functions, the gamma function, the beta function, the hypergeometric

function, the elliptic functions, the theta function, the Bessel function, the Mathieu function, the Weber function, Struve function, the Airy function, Lamé function. Literally hundreds of special functions were singled out for scrutiny and their main features were drawn. This is an art which is not much cultivated in these days. Times have changed and emphasis has shifted. Mathematicians on the whole prefer more abstract fare. Large classes of functions are studied instead of individual ones. Sociology has replaced biography. The field of special functions, as it is now known, is left largely to a small but ardent group of enthusiasts plus those whose work in physics or engineering confronts them directly with the necessity of dealing with such matters.“

A.9 Frühe Ansätze zur Theorie der Fourier-Reihen

1.) Anlässe für die Theorie der trigonometrischen Reihen

Das Problem der schwingenden Saite spielte eine wesentliche Rolle bei der Herausbildung des Begriffs der „willkürlichen Funktion“ im 18. und 19. Jahrhundert. Eine knappe Einführung in die physikalische Modellierung und eine Heranführung an moderne Fragestellungen findet man in ROLF LEIS „*Zur Entwicklung der angewandten Analysis und mathematischen Physik in den letzten hundert Jahren*“ in: Ein Jahrhundert Mathematik, Festschrift zum Jubiläum der DMV, Vieweg 1990 (S. 491–535).

Sehr lesenswert sind auch die Ausführungen von HEUSER im „Lehrbuch der Analysis, Teil 2“, S. 685–700. Wir stützen uns auf Heuser.

DANIEL BERNOULLI ging 1753 von der physikalischen Vorstellung aus, daß jeder Ton durch eine Überlagerung von Grund- und Obertönen entsteht und schloß daraus, daß die Bewegung einer beidseitig eingespannten Saite der Länge L stets in der Form

$$u(x, t) = \sum_{n=1}^{\infty} a_n \sin \frac{n\pi x}{L} \cdot \cos \frac{n\pi \alpha t}{L}$$

dargestellt werden könne. Für $t = 0$ erhält man die Entwicklung

$$g(x) = \sum_{n=1}^{\infty} a_n \sin \frac{n\pi x}{L}, \quad 0 \leq x \leq L$$

für die Anfangsauslenkung g . Ohne weitere mathematische Argumente kam D. Bernoulli zu dem Schluß, daß sich jede („willkürlich wählbare“) Anfangsauslenkung in der angegebenen Weise schreiben läßt. Euler, d'Alembert und später auch Lagrange verwarfen diese Argumentation von Bernoulli. Sie glaubten auch das Resultat nicht. Die trigonometrischen Reihen schienen ihnen zu starke analytische Eigenschaften zu besitzen, als daß sie „willkürliche Funktionen“ g darstellen könnten. Bernoulli ließ sich nicht beirren und meinte, daß die unendlich vielen Koeffizienten sehr wohl ausreichen könnten, um die Reihe „jedem“ g anpassen zu können.

Für die These von D. Bernoulli sprach, daß man auch sehr bizarre $g(\cdot)$ als trigonometrische Reihen darstellen können. Dafür sprach z.B. auch die Entdeckung von Euler, daß die Reihe

$$\sin x + \frac{1}{2} \sin 2x + \frac{1}{3} \sin 3x + \dots$$

die unstetige 2π -periodische Funktion darstellte, welche im Intervall $(0, 2\pi)$ mit der Funktion $\frac{1}{2}(\pi - x)$ übereinstimmt.

Euler hat an verschiedenen Stellen seines umfangreichen Werks verschiedene Definitionen angeboten, was man sich unter einer Funktion vorstellen sollte. Das Übergewicht hatte zunächst der auf analytische Ausdrücke eingeschränkte Funktionsbegriff: „Eine Funktion einer veränderlichen Größe ist ein analytischer Ausdruck, der in beliebiger Weise aus dieser

veränderlichen Größe und aus Zahlen oder konstanten Größen zusammengesetzt ist.“ Dabei bleibt aber unklar, was unter einem „analytischen Ausdruck“ zu verstehen ist. Euler dachte nicht nur an die gängigen finiten Operationen sondern sehr wohl auch an Grenzprozesse, was er aber nicht genauer faßte. Eine wirkliche Theorie der Konvergenz von Funktionenfolgen entstand erst im späten 19. Jahrhundert.

Die Theorie der trigonometrischen Reihen gab immer wieder bedeutsame Anstöße, über den allgemeinen Funktionsbegriff nachzudenken. FOURIER (1768–1830) befaßte sich in seiner vielbeachteten Theorie der Wärmeleitung sehr gründlich mit den trigonometrischen Reihen. Er stellte die Tatsache klar, daß eine Funktion, die sich durch ein stetiges Kurvenstück oder durch die Aneinanderreihung von solchen darstellen läßt, durch eine trigonometrische Reihe dargestellt werden kann. Es wird berichtet, daß er 1807, als er seine Idee zum erstenmal bekanntgab, auf den scharfen Widerspruch von Lagrange stieß (vgl. Struik S. 168).

2.) Formale Fourier–Reihen

Der moderne Begriff der formalen Fourier–Reihe nimmt keinen Bezug auf den Funktionsbegriff; formale Fourier–Reihen sind abstrakte Rechenobjekte. Erst durch Konvergenzbetrachtungen für die Folge der Partialsummen wird der Bezug zu 2π –periodischen Funktionen (oder auch Maßen) auf dem Intervall $(0, 2\pi]$ hergestellt.

Definition

- a) Eine **formale Fourier–Reihe** ist ein formaler Ausdruck der Gestalt

$$\sum_{-\infty}^{+\infty} c_n \cdot e^{inx}$$

mit komplexen Koeffizienten c_n .

- b) Wenn nur endlich viele der c_n ungleich 0 sind, spricht man von einem **trigonometrischen Polynom**.

- c) Die N –te **Partialsumme** der formalen Fourier–Reihe ist das trigonometrische Polynom

$$S_N(x) = \sum_{-N}^{+N} c_n \cdot e^{inx} .$$

Für viele Anwendungen ist es bequem, die formalen Ausdrücke folgendermaßen umzuschreiben

$$\begin{aligned} c_0 &= \frac{1}{2} a_0 & \text{und für } k = 1, 2, \dots \\ a_k &= c_k + c_{-k} , & b_k &= i(c_k - c_{-k}) \end{aligned}$$

$$\text{d.h. } c_k = \frac{1}{2}(a_k - ib_k), \quad c_{-k} = \frac{1}{2}(a_k + ib_k).$$

Man hat dann

$$\begin{aligned} c_k \cdot e^{ikx} + c_{-k} \cdot e^{-ikx} &= a_k \cdot \cos kx + b_k \cdot \sin kx \\ S_N(x) &= \frac{a_0}{2} + \sum_{k=1}^N (a_k \cdot \cos kx + b_k \cdot \sin kx). \end{aligned}$$

Dieses Umschreiben ist besonders dann bequem, wenn

$$c_{-n} = \bar{c}_n \quad \text{für alle } n \in \mathbb{Z};$$

in diesem Falle sind nämlich die a_k, b_k reelle Zahlen.

Beispiel $\sum_{n \neq 0} \frac{1}{2ni} e^{inx}$ liefert $\sum_{k=1}^{\infty} \frac{1}{k} \cdot \sin kx$.

Rechnen mit formalen Fourier-Reihen

Man kann formale Fourier-Reihen addieren; man kann sie mit Skalaren multiplizieren; und man kann sie komplex konjugieren.

$$\begin{aligned} \text{(i)} \quad & \sum c_n e^{inx} + \sum d_n e^{inx} = \sum (c_n + d_n) e^{inx} \\ \text{(ii)} \quad & \lambda \cdot \sum c_n e^{inx} = \sum (\lambda c_n) e^{inx} \\ \text{(iii)} \quad & \left(\sum c_n e^{inx} \right)^* = \sum \bar{c}_n \cdot e^{-inx} \\ \text{(Genauer } & = \sum d_n e^{inx} \text{ mit } d_n = \bar{c}_{-n}.) \end{aligned}$$

Sprechweise Eine formale Fourier-Reihe heißt eine **reelle formale Fourier-Reihe**, wenn sie mit ihrer komplex konjugierten formalen Fourier-Reihe identisch ist, d.h. wenn $\bar{c}_{-n} = c_n$ für alle $n \in \mathbb{Z}$. Eine reelle formale Fourier-Reihe hat also die Gestalt

$$\frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos kx + b_k \sin kx) \quad \text{mit } a_k, b_k \text{ reell.}$$

Aufgabe Man schreibe die Formel für die schwingende Saite

$$u(x, t) = \sum_{k=1}^{\infty} d_k \cdot \sin(kx) \cdot \cos(\alpha \cdot kt)$$

in komplexer Form.

Absolute Konvergenz

Definition Eine Fourier-Reihe $\sum c_n e^{inx}$ heißt absolut konvergent, wenn $\sum |c_n| < \infty$.

Satz Eine absolutkonvergente Fourier-Reihe stellt eine stetige 2π -periodische Funktion dar. Gleichmäßig für $x \in \mathbb{R}$ gilt

$$f(x) = \lim_{N \rightarrow \infty} \sum_{n=-N}^{-N} c_n e^{inx}.$$

Die Koeffizienten sind durch $f(\cdot)$ eindeutig bestimmt. Es gilt nämlich

$$c_n = \frac{1}{2\pi} \int_0^{2\pi} f(x) e^{-inx} dx.$$

Beweis

$$\begin{aligned} 1) \quad & |f(x) - S_N(x)| \leq \sum_{n=N+1}^{\infty} |c_n| \\ 2) \quad & \frac{1}{2\pi} \int_0^{2\pi} e^{imx} \cdot e^{-inx} dx = \begin{cases} 1 & \text{für } m = n \\ 0 & \text{für } m \neq n \end{cases} \\ & \frac{1}{2\pi} \int_0^{2\pi} S_N(x) e^{-inx} dx = \begin{cases} c_n & \text{für } |n| \leq N \\ 0 & \text{sonst} \end{cases} \\ & \frac{1}{2\pi} \int_0^{2\pi} f(x) e^{-inx} dx = c_n \quad \text{für alle } n. \end{aligned}$$

Bemerke

$$\frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos kx + b_k \sin kx)$$

ist absolutkonvergent genau dann, wenn $\sum |a_k| < \infty$ und $\sum |b_k| < \infty$. Aus der dargestellten Funktion $f(x)$ erhalten wir die Koeffizienten

$$\begin{aligned} a_k &= \frac{1}{\pi} \int_0^{2\pi} f(x) \cos kx dx \\ b_k &= \frac{1}{\pi} \int_0^{2\pi} f(x) \sin kx dx \end{aligned}$$

Satz Wenn $f(\cdot)$ und $g(\cdot)$ absolutkonvergente Fourier-Reihen besitzen, dann auch das punktweise Produkt $h(\cdot) = (f \cdot g)(\cdot)$.

Beweis

$$f(x) \cdot g(x) = \left(\sum c_n e^{inx} \right) \left(\sum d_m e^{imx} \right) = \sum p_k e^{ikx}$$

mit $p_k = \sum_{m+n=k} c_n \cdot d_m = \sum_n c_n \cdot d_{k-n}$

$$\sum_k |p_k| \leq \sum |c_n| \cdot \sum |d_m| .$$

Sprechweise Wenn die Koeffizientenfolge

$$P = (\dots, p_{-1}, p_0, p_1, p_2, \dots)$$

in dieser Weise aus

$$C = (\dots, c_{-1}, c_0, c_1, c_2, \dots)$$

$$D = (\dots, d_{-1}, d_0, d_1, d_2, \dots)$$

hervorgeht, dann heißt P das Faltungsprodukt von C und D und man notiert

$$P = C * D .$$

Hinweis Eine Folge

$$P = (\dots, p_{-1}, p_0, p_1, p_2, \dots)$$

mit $p_k \geq 0$ für alle k und $\sum p_k = 1$ heißt in der Wahrscheinlichkeitstheorie eine Wahrscheinlichkeitsgewichtung auf $\mathbb{Z}$. Die Funktion $\sum p_k e^{ikx}$ heißt die charakteristische Funktion dieser Gewichtung.

3.) Die Fourier-Koeffizienten einer 2π -periodischen Funktion

Definition Wenn $f(\cdot)$ auf $[0, 2\pi]$ integrabel ist, dann definiert man die Fourier-Koeffizienten

$$c_n := \frac{1}{2\pi} \int_0^{2\pi} f(x) e^{-inx} dx \quad \text{für } n \in \mathbb{N} .$$

Man nennt $\sum c_n e^{inx}$ die formale Fourier-Reihe zu f und notiert traditionellerweise

$$f \sim \sum_{-\infty}^{+\infty} c_n e^{inx} .$$

(Beachte: Das Zeichen $\sim$ soll nicht suggerieren, daß es um eine Äquivalenz geht; f und $\sum c_n e^{inx}$ sind mathematische Objekte verschiedener Art.)

Beachte Wenn man die (nicht notwendigerweise stetige) Funktion $f(\cdot)$ auf einer Nullmenge abändert, dann ändert das die Fourier-Koeffizienten nicht. Man kann auf der anderen Seite zeigen, daß die Abänderung auf einer Nicht-Nullmenge die formale Fourier-Reihe tatsächlich ändert; der Beweis braucht aber Vorbereitungen.

Satz

$$(i) \quad f \sim \sum c_n e^{inx}, \quad g \sim \sum d_n e^{inx} \implies f + g \sim \sum (c_n + d_n) e^{inx}$$

$$(ii) \quad \lambda f \sim \sum (\lambda c_n) e^{inx} \quad \text{für alle Skalare } \lambda$$

$$(ii) \quad \overline{f} \sim \left(\sum c_n e^{inx} \right)^*$$

Der Beweis ist trivial.

Beachte Wir reden hier von integrierbaren Funktionen auf $[0, 2\pi]$. Wir könnten ebenso gut von 2π -periodischen Funktionen reden, die auf jedem endlichen Intervall integrierbar sind.

Satz Sei $f(\cdot)$ integrierbar auf $[0, 2\pi]$ mit $c_0 = \frac{1}{2\pi} \int_0^{2\pi} f(x) dx = 0$. Für die Stammfunktion

$F(\cdot)$ mit $\int_0^{2\pi} F(x) dx = 0$ gilt dann

$$F \sim \sum \left(\frac{-i}{n} \right) c_n e^{inx}.$$

Wenn $f(\cdot)$ reellwertig ist mit

$$f \sim \sum_{k=1}^{\infty} (a_k \cos kx + b_k \sin kx),$$

dann haben wir

$$F \sim \sum_1^{\infty} \left(\frac{1}{k} a_k \sin kx - \frac{1}{k} b_k \cos kx \right).$$

Beweis (mit Hilfe partieller Integration)

$$\frac{1}{2\pi} \int_0^{2\pi} F(x) e^{-inx} dx = \left(-\frac{i}{n} \right) \cdot \frac{1}{2\pi} \int_0^{2\pi} f(x) e^{-inx} dx.$$

Ebenso der reelle Fall.

Beispiel Für die folgende 2π -periodische Funktion $f(\cdot)$ und ihre Stammfunktion $F(\cdot)$ studieren wir die formalen Fourier-Reihen:

$$\begin{aligned} f(x) &= \begin{cases} +1 & \text{für } 0 < x < \pi \\ -1 & \text{für } -\pi < x < 0 \end{cases} \quad (\text{ungerade}) \\ F(x) &= |x| \quad \text{für } |x| < \pi \quad (\text{gerade}) \end{aligned}$$

Wir haben für $k = 1, 2, \dots$

$$\begin{aligned} \frac{1}{\pi} \int_{-\pi}^{+\pi} f(x) \sin kx \, dx &= -\frac{1}{k\pi} [\cos kx]_0^\pi + \frac{1}{k\pi} [\cos kx]_{-\pi}^0 \\ &= \frac{1}{k\pi} [2 \cos 0 - 2 \cos k\pi] \end{aligned}$$

$$\begin{aligned} f &\sim \frac{4}{\pi} \left[\sin x + \frac{1}{3} \sin 3x + \frac{1}{5} \sin 5x + \dots \right] \\ F &\sim \frac{4}{\pi} \left[\cos x + \frac{1}{3^2} \cos 3x + \frac{1}{5^2} \cos 5x + \dots \right] \end{aligned}$$

Die Fourier-Reihe für $F(\cdot)$ ist absolutkonvergent. Es stellt sich heraus, daß sie tatsächlich $F(\cdot)$ darstellt.

$$|x| = \frac{\pi}{2} - \frac{4}{\pi} \left[\cos x + \frac{1}{3^2} \cos 3x + \frac{1}{5^2} \cos 5x + \dots \right] \quad \text{für } |x| \leq \pi.$$

Zusatz Die Folge der Partialsummen

$$\frac{4}{\pi} \sin x, \quad \frac{4}{\pi} \left[\sin x + \frac{1}{3} \sin 3x \right], \quad \frac{4}{\pi} \left[\sin x + \frac{1}{3} \sin 3x + \frac{1}{5} \sin 5x + \dots \right]$$

ist im Buch von FORSTER, Analysis 1, S. 198 graphisch dargestellt. In jedem abgeschlossenen Intervall $\subseteq (0, \pi)$ herrscht gleichmäßige Konvergenz gegen die Konstante 1.

Man kann das ähnlich wie in der Übungsaufgabe zeigen, indem man beachtet, daß

$$s(x) := \frac{4}{\pi} \left[\sin x + \frac{1}{3} \sin 3x + \dots + \frac{1}{2N-1} \sin(2N-1)x \right]$$

die Stammfunktion ist für das elementar summierbare trigonometrische Polynom

$$\begin{aligned} c(x) &= \frac{4}{\pi} [\cos x + \cos 3x + \dots + \cos(2N-1)x] \\ &= \frac{2}{\pi} [e^{(-2N+1)ix} + e^{(-2N+3)ix} + \dots + e^{(2N-1)ix}] \\ &= \frac{2}{\pi} \frac{e^{Nix} - e^{-Nix}}{e^{ix} - e^{-ix}} = \frac{2}{\pi} \frac{\sin 2Nx}{\sin x} \quad \text{für } 0 < x < \pi \\ s(x) - s\left(\frac{\pi}{2}\right) &= \int_{\pi/2}^x c(t) \, dt = \int_{\pi/2}^x \frac{2}{\pi} \frac{\sin 2Nt}{\sin t} \, dt. \end{aligned}$$

Partielle Integration zeigt, daß diese Funktion in jedem abgeschlossenen Intervall $[a, b] \subseteq (0, \pi)$ gleichmäßig nach 0 konvergiert. Daraus ergibt sich

$$\int_{\pi/2}^x [s(y) - s(\pi/2)] dy \longrightarrow 0 \quad \text{gleichmäßig für } x \in [a, b] \subseteq (0, \pi) .$$

$$\text{Aus } \pi/4 = \arctan 1 = \lim_{N \rightarrow \infty} \left(1 - \frac{1}{3} + \frac{1}{5} - \dots + (-1)^{N-1} \frac{1}{2N-1} \right)$$

$$s(\pi/2) = \frac{4}{\pi} \left[1 - \frac{1}{3} + \frac{1}{5} - \dots \pm \frac{1}{2N-1} \right] \longrightarrow 1 \quad \text{für } N \rightarrow \infty .$$

Daraus ergibt sich leicht die Behauptung

$$\cos x + \frac{1}{3^2} \cos 3x + \frac{1}{5^2} \cos 5x + \dots = \frac{\pi^2}{8} - \frac{\pi}{4} |x| \quad \text{für } |x| \leq \pi .$$

Rechnungen dieser Art passen ins 18. Jahrhundert. Die modernere Theorie greift die Darstellbarkeit von Funktionen durch ihre Fourier-Reihe viel allgemeiner und wirksamer an.

4.) Anhang für Informatikstudenten:

Die $\mathcal{Z}$ -Transformierte

In der Elektrotechnik treten Folgen

$$C = (\dots, c_{-1}, c_0, c_1, \dots)$$

als Folgen von Signalen auf oder auch als die Abtastwerte einer Funktion $h(t)$ zu äquidistanten Zeitpunkten. Zum Zwecke der Analyse betrachtet man dazu die sog. $\mathcal{Z}$ -Transformierte

$$\mathcal{Z}(C) := \sum_{n=-\infty}^{+\infty} c_n z^{-n} .$$

Dies ist zunächst nur ein formaler Ausdruck, mit dem man Rechnungen anstellen kann.

Beispiel Man sagt, die Folge D entsteht aus der Folge C durch einen linearen Filter mit den Filterkoeffizienten $(a_0, a_1, \dots, a_N)$, wenn für alle $n \in \mathbb{Z}$

$$d_n = a_0 c_n + a_1 c_{n-1} + a_2 c_{n-2} + \dots + a_N c_{n-N} .$$

Die $\mathcal{Z}$ -Transformierte ist offenbar

$$\mathcal{Z}(D) = \left[a_0 + a_1 \frac{1}{z} + \dots + a_N \left(\frac{1}{z} \right)^N \right] \cdot \mathcal{Z}(C) .$$

Als Mathematiker bemerken wir: Wenn $c_n = 0$ für alle $n < 0$, dann ist die $\mathcal{Z}$ -Transformierte $\mathcal{Z}(C)$ eine formale Potenzreihe in der Unbestimmten z^{-1} .

Eine solche Potenzreihe hat in günstigen Fällen Werte für gewisse komplexe Werte z . Wie wir wissen, gibt es in diesem Fall ein $r > 0$, so daß $\mathcal{Z}(C)$ Werte hat im Bereich

$$D = \{z : |z| > r\}$$

(unter Umständen auch noch auf einem Teil des Randes von D).

Wenn $c_n = 0$ für alle $n > 0$, dann ist $\mathcal{Z}(C)$ eine formale Potenzreihe in der Unbestimmten z . Sie hat in günstigen Fällen Werte im Innern eines Konvergenzkreises $\{z : |z| < R\}$ (unter Umständen auch auf einem Teil des Randes).

In günstigen Fällen (wenn die $|c_n|$ für $n \rightarrow \pm\infty$ genügend schnell fallen) liefert die $\mathcal{Z}$ -Transformierte eine Funktion $F(z)$ in einem Kreisring

$$D = \{z : r < |z| < R\}.$$

In anderen (nicht so günstigen) Fällen hat $\mathcal{Z}(C)$ vielleicht immerhin Werte auf einem Kreis

$$K = \{z : |z| = R\}.$$

In solchen Fällen bietet sich an, die Unbestimmte umzubenennen

$$z = R \cdot e^{-ix}.$$

Die $\mathcal{Z}$ -Transformierte hat dann die Gestalt einer formalen Fourier-Reihe

$$\mathcal{Z}(C) = \sum_{-\infty}^{+\infty} c_n z^{-n} = \sum_{-\infty}^{+\infty} (c_n \cdot R^{-n}) e^{inx}.$$

Diese Fourier-Reihe (die ja nach Annahme Werte hat) liefert uns eine 2π -periodische Funktion $f(x)$, aus welcher man die Folge C zurückgewinnen kann. Manchmal kann man aus dem Studium dieser Funktion nützliche Schlüsse ziehen.

Viele der in der Praxis (der Elektrotechnik) auftretenden $\mathcal{Z}$ -Transformierten haben die Gestalt einer gebrochenrationalen Funktion $F(z) = \frac{\mathcal{Z}(z)}{\mathcal{N}(z)}$ mit Polynomen $\mathcal{Z}(\cdot)$ und $\mathcal{N}(\cdot)$. Ein probates Mittel zum Studium der Koeffizienten von $F(\cdot)$ ist die Partialbruchzerlegung.

Die Frankfurter Informatiker erwarten von ihren Studenten zwar keine Kenntnisse in irgendeiner Theorie der $\mathcal{Z}$ -Transformierten; sie erwarten aber, daß die Studenten nicht erschrecken, wenn irgendwo (z.B. in der technischen Informatik) einfache Beispiele auftreten. Wir erwähnen hier drei Beispiele und verweisen im übrigen auf Tabellen, wie man sie z.B. im bekannten „Teubner Taschenbuch der Mathematik“ („Bronstein“ oder Springers Mathematische Formeln, L. Røade u. Bertil Westergren) findet.

$$1) \ c_n = \begin{cases} 0 & \text{für } n < 0 \\ a^n & \text{für } n \geq 0 \end{cases}$$

$$\begin{aligned} \mathcal{Z}(C) &= \sum_{n=0}^{\infty} a^n z^{-n} = 1 + \frac{a}{z} + \left(\frac{a}{z}\right)^2 + \dots \\ &= \frac{1}{1 - \frac{a}{z}} = \frac{z}{z - a} \quad \text{für } |z| > |a|. \end{aligned}$$

Wenn man im Falle $a = 1$ eine Folge $D = (\dots, d_{-1}, d_0, d_1, \dots)$ mit diesem C faltet, erhält man die Folge

$$(D \star C)_k = \sum_{n \leq k} d_n,$$

d.h. die Folge der Partialsummen und dazu die $\mathcal{Z}$ -Transformierte

$$\mathcal{Z}(D \star C) = \frac{z}{z - 1} \cdot \mathcal{Z}(D).$$

2) Für reelle Ω

$$\begin{aligned} c_n &= \begin{cases} 0 & \text{für } n < 0 \\ \cos(\Omega n) & \text{für } n \geq 0 \end{cases} \\ \mathcal{Z}(C) &= \sum_{n=0}^{\infty} \cos(\Omega n) z^{-n} = \frac{1}{2} \sum_{n=0}^{\infty} (e^{i\Omega n} + e^{-i\Omega n}) z^{-n} \\ &= \frac{1}{2} \left(\frac{z}{z - e^{i\Omega}} + \frac{z}{z - e^{-i\Omega}} \right) \\ &= \frac{z^2 - z \cos \Omega}{z^2 + 1 - 2z \cos \Omega} \quad \text{für } |z| > 1 \end{aligned}$$

$$3) \ c_n = \begin{cases} 0 & \text{für } n < 0 \\ \sin(\Omega n) & \text{für } n \geq 0 \end{cases}.$$

$$\mathcal{Z}(C) = \frac{z \sin \Omega}{z^2 + 1 - 2z \cos \Omega} \quad \text{für } |z| > 1.$$

Mit diesen Andeutungen hören wir hier auf, denn die Angelegenheit kann nur durch echte Anwendungen wirklich spannend gemacht werden.

5.) Geschichtliches: Impulse, die von Fragen über Fourier-Reihen ausgegangen sind

- a) BERNHARD RIEMANN (1826–1866) legte 1854 seine Habilitationsschrift vor: „Über die Darstellbarkeit einer Funktion durch eine trigonometrische Reihe“. Er schreibt, daß durch Dirichlets Untersuchungen zwar wohl alle in der Natur auftretenden Fälle erledigt seien, hebt dann aber hervor, daß „die Anwendbarkeit der Fourier’schen Reihen nicht auf physikalische Untersuchungen beschränkt [ist]; sie ist jetzt auch in einem Gebiete der reinen Mathematik, der Zahlentheorie, angewandt, und hier schienen gerade diejenigen Funktionen, deren Darstellbarkeit durch eine trigonometrische Reihe Dirichlet nicht untersucht hat, von Wichtigkeit zu sein.“ Eine der Bedingungen bei Dirichlet forderte, daß die Funktion „integrabel“ sein muß. Was bedeutet das aber? Cauchy und Dirichlet hatten schon gewisse Antworten gegeben; Riemann setzte an ihre Stelle seine Konstruktion, die in den heutigen Lehrbüchern als das „Riemannsche Integral“ abgehandelt wird. Riemann zeigte u.a., daß durch Fourier-Reihen definierte Funktionen solche Eigenschaften haben können, wie das Auftreten unendlich vieler Maxima und Minima, die von älteren Mathematikern bei der Definition einer Funktion nicht zugelassen worden wären. In seinen Vorlesungen gab Riemann ein Beispiel einer stetigen Funktion ohne Ableitung.
- b) Auch Weierstraß benützte 1861 Fouriersche Reihen für sein Beispiel für eine stetige Funktion, die **stetig aber nirgends differenzierbar** ist. Sein Beispiel war die Funktion

$$f(x) = \sum_{n=0}^{\infty} b^n \cdot \cos(a^n \pi x) ,$$

wo a eine ungerade ganze Zahl ist und $b \in (0, 1)$ mit

$$ab > 1 + \frac{3\pi}{2} .$$

Solche Funktionen wurden von den meisten Zeitgenossen als „pathologische“ Funktionen abgelehnt. Die Beispiele machten aber klar, daß die Grundlagen der Analysis neu bedacht werden mußten.

Aufgabe Für die Weierstraßsche Funktion $f(\cdot)$ ist zu zeigen

$$\exists M > 0 \quad \forall \varepsilon > 0 \quad \exists x_1, x_2 :$$

$$\begin{aligned} (|x_1 - x| < \varepsilon) \quad \wedge \quad (|x_2 - x| < \varepsilon) \\ \wedge \quad \left(\frac{f(x_1) - f(x)}{x_1 - x} > M \right) \\ \wedge \quad \left(\frac{f(x_2) - f(x)}{x_2 - x} > -M \right) \end{aligned}$$

- c) Das Riemannsche Integral kann heute nur noch als eine Vorstufe zum allgemeinen **Integralbegriff von Lebesgue** gelten. Für jede Lebesgue-integrierte Funktion $f(x)$ definiert man die Koeffizienten

$$c_n = \frac{1}{2\pi} \int_{-\pi}^{+\pi} f(x) e^{-inx} dx \quad \text{für } n \in \mathbb{Z}$$

und die trigonometrischen Polynome

$$S_N(x) = \sum_{-N}^{+N} c_n e^{inx}.$$

Generationen von Mathematikern haben sich mit der Frage beschäftigt, in welchem Sinne die Folge $S_N(\cdot)$ konvergiert. Es wird heute bezweifelt, daß es eine erschöpfende Antwort, die alle integrierbaren $f(\cdot)$ erfaßt, je geben wird. Ein Aufsehen erregendes Resultat ist 1966 LENNART CARLESON gelungen. Carleson hat bewiesen, daß für alle $f(\cdot)$ mit

$$\int |f(x)|^p dx < \infty \quad \text{für } p > 1$$

die Folge $(S_N(\cdot))_N$ **fast überall** gegen $f(\cdot)$ konvergiert. Der Beweis ist außerordentlich schwierig. Es ist bekannt, daß etwas Besseres als (im Lebesgueschen Sinne) Fast überall-konvergenz nicht zu erwarten ist. Es gibt stetige Funktionen, für welche $S_N(x)$ nicht für alle x konvergiert. (Heuser sagt mehr zur Problematik auf S. 154 seines Lehrbuchs der Analysis, Teil 2).

Anhang: Zum Kontext bei Dirichlet, Riemann und Weierstraß

Die Fourier-Reihen lieferten, wie gesagt, im ganzen 19. Jahrhundert immer wieder Anlässe zur Frage, was man sich nun eigentlich unter einer "ganz allgemeinen" Funktion vorstellen sollte. Besonders P.L. DIRICHLET und B. RIEMANN gaben dazu wichtige Impulse. Zum Wirken von Dirichlet und Riemann in weiten Bereichen der Mathematik schreibt Struik in seinem „Abriß der Geschichte der Mathematik“ (S. 181 ff):

Peter Lejeune Dirichlet war sowohl mit Gauß und Jacobi als auch mit den französischen Mathematikern eng verbunden. Er lebte von 1822–1827 als Privatlehrer und traf mit Fourier zusammen, dessen Buch er studierte; er machte sich auch mit den „Disquisitiones arithmeticae“ von Gauß vertraut. Er lehrte später an der Universität Breslau und wurde 1855 Nachfolger von Gauß in Göttingen. Seine persönliche Bekanntschaft sowohl mit französischen und deutschen Mathematikern als auch mit der mathematischen Wissenschaft beider Länder ließen ihn zum geeigneten Mann werden, um als Interpret von Gauß zu wirken und die Fourierreihen einer tieferschürfenden Analyse zu unterziehen. Dirichlets ausgezeichnete „Vorlesungen über Zahlentheorie“ (1863 veröffentlicht) bilden noch immer eine der besten Einführungen in die Gaußschen Untersuchungen zur Zahlentheorie. Sie enthalten auch

viele neue Ergebnisse. In einer Arbeit aus dem Jahre 1840 lehrte Dirichlet, wie man die volle Kraft der Theorie der analytischen Funktionen auf Probleme der Zahlentheorie anwendet; eben in diesen Untersuchungen führte er die „Dirichletschen“ Reihen ein. Er erweiterte auch den Begriff der quadratischen Irrationalitäten zu dem der allgemeinen algebraischen Rationalitätsbereiche (Körper).

Dirichlet gab als erster einen strengen Konvergenzbeweis für Fourierreihen und trug auf diese Weise zum richtigen Verständnis des Wesens einer Funktion bei. Außerdem führte er das sogenannte Dirichletsche Prinzip in die Variationsrechnung ein, das die Existenz einer Funktion v als gegeben annahm, die das Integral $\int (v_x^2 + v_y^2 + v_z^2) d\tau$ bei vorgegebenen Randbedingungen zum Minimum macht. Es war eine Abänderung eines Prinzips, das Gauß in seiner Potentialtheorie von 1839/40 eingeführt hatte, und diente später Riemann als kraftvolles Hilfsmittel bei der Lösung von Problemen der Potentialtheorie. Es wurde schon erwähnt, daß schließlich Hilbert die Gültigkeit dieses Prinzips streng nachwies (S. 164).

13. Mit Bernhard Riemann, dem Nachfolger Dirichlets in Göttingen, kommen wir zu dem Mann, der mehr als irgendein anderer den Weg der modernen Mathematik beeinflusst hat. Riemann war der Sohn eines Landpfarrers und studierte an der Universität Göttingen, wo er 1851 den Doktorgrad erwarb. Im Jahre 1854 wurde er Privatdozent und 1859 Professor an der gleichen Universität. Er war wie Abel kränklich und verbrachte seine letzten Tage in Italien, wo er 1866 im Alter von 40 Jahren starb. In seinem kurzen Leben hat er nur eine verhältnismäßig kleine Anzahl von Arbeiten veröffentlicht, aber jede von ihnen war — und ist es noch — bedeutend, und einige von ihnen haben ganz neue und fruchtbare Gebiete eröffnete. Im Jahre 1851 erschien Riemanns Doktordissertation über die Theorie der komplexen Funktionen $u + iv = f(x + iy)$. Wie d'Alembert und Cauchy wurde Riemann von hydrodynamischen Vorstellungen beeinflusst. Er bildete die x, y -Ebene konform auf die u, v -Ebene ab und wies die Existenz einer Funktion nach, die die Transformation eines beliebigen einfach zusammenhängenden Gebiets der einen Ebene in ein beliebiges einfach zusammenhängendes Gebiet der anderen Ebene leistete. Das führte zur Idee der Riemannschen Fläche, die topologische Betrachtungen in die Analysis hineinrug. In jener Zeit war die Topologie noch ein fast unberührtes Gebiet, über das J.B. Listing in den „Göttinger Studien“ von 1847 eine Arbeit veröffentlicht hatte. Riemann zeigte ihre zentrale Bedeutung für die Theorie der komplexen Funktionen. Diese Dissertation erläuterte auch die Riemannsche Definition einer komplexen Funktion: Real- und Imaginärteil müssen in einem gegebenen Gebiet den „Cauchy-Riemannschen“ Differentialgleichungen genügen: $u_x = v_y$, $u_y = -v_x$ und müssen außerdem gewisse Bedingungen bezüglich des Randes und der Singularitäten erfüllen.

Riemann wendete seine Ideen auf hypergeometrische und abelsche Funktionen (1857) an, wobei er ausgiebig von dem (von ihm so bezeichneten) Dirichletschen Prinzip Gebrauch machte. Unter seinen Ergebnissen findet sich die Entdeckung des Geschlechts einer Fläche als einer topologischen Invariante und eines Hilfsmittels zur Klassifizierung abelscher Funktionen. In einer posthum veröffentlichten Arbeit werden diese Ideen auf Minimalflächen angewendet (1867). Zu diesem Teil der Riemannschen Tätigkeit gehören auch seine Untersuchungen über elliptische Modulfunktionen und Thetareihen von p unabhängigen Veränderlichen ebenso wie die über lineare Differentialgleichungen mit algebraischen Koeffizienten. Riemann wurde im Jahre 1854 Privatdozent, indem er nicht weniger als zwei grundlegende Arbeiten vorlegte, eine über trigonometrische Reihen und die Grundlagen der Analysis, die andere über die Grundlagen der Geometrie. In der ersten Arbeit wurden die Dirichletschen Bedingungen für die Entwickelbarkeit einer Funktion in eine Fourierreihe untersucht. Eine dieser Bedingungen forderte, daß die Funktion „integrabel“ sein muß. Was bedeutet das aber? Cauchy und Dirichlet hatten schon gewisse Antworten gegeben; Riemann ersetzte sie durch seine eigene umfassende Antwort. Er gab jene Definition, die wir heute als „Riemannsches Integral“ kennen und die erst im zwanzigsten Jahrhundert durch das Lebesguesche Integral ergänzt wurde. Riemann zeigte, daß durch Fourierreihen definierte

Funktionen solche Eigenschaften besitzen können wie das Auftreten unendlich vieler Maxima und Minima, die von älteren Mathematikern bei der Definition einer Funktion nicht zugelassen worden wären. Der Funktionsbegriff begann sich ernstlich von der „curva quaecunque libero manus ductu descripta“³⁾ von Euler zu befreien. In seinen Vorlesungen gab Riemann ein Beispiel einer stetigen Funktion ohne Ableitung; ein Beispiel einer solchen Funktion, das Weierstraß angegeben hatte, wurde im Jahre 1875 veröffentlicht. Die Mathematiker weigerten sich, derartige Funktionen besonders ernst zu nehmen und nannten sie „pathologische“ Funktionen; die moderne Analysis hat nachgewiesen, wie naturgemäß solche Funktionen sind und wie Riemann hier wieder zu einem grundlegenden Gebiet der Mathematik vorgestoßen war.

Die andere Arbeit aus dem Jahre 1854 behandelte die Hypothesen, welche der Geometrie zugrunde liegen. Der Raum wurde als topologische Mannigfaltigkeit von beliebiger Dimensionszahl eingeführt; in einer solchen Mannigfaltigkeit wurde mittels einer quadratischen Differentialform eine Metrik definiert. Während Riemann in der Analysis eine komplexe Funktion durch ihr lokales Verhalten definiert hatte, definierte er in dieser Arbeit den Charakter des Raumes in derselben Weise, Riemanns vereinheitlichendes Prinzip ermöglichte ihm nicht nur, alle vorhandenen Formen der Geometrie zu klassifizieren, einschließlich der noch sehr undurchsichtigen nichteuklidischen Geometrie, sondern gestattete ihm auch die Schaffung einer beliebigen Anzahl von neuen Raumtypen, von denen seither viele einen nützlichen Platz in der Geometrie und mathematischen Physik gefunden haben. Riemann veröffentlichte seine Arbeit ohne irgendwelche analytischen Hilfsmittel, wodurch es schwierig war, seinen Ideen zu folgen. Später erschienen einige der Formeln in einer Preisschrift über die Wärmeverteilung in einem festen Körper, die Riemann der Pariser Akademie eingereicht hatte (1861). Hierin ist eine Skizze der Transformationstheorie der quadratischen Formen enthalten.

Die letzte Arbeit von Riemann, die erwähnt werden muß, ist seine Diskussion der Anzahl $F(x)$ der Primzahlen, die unterhalb einer gegebenen Zahl x liegen (1859). Es handelte sich um eine Anwendung der Theorie der komplexen Zahlen auf das Problem der Verteilung der Primzahlen, und es wurde die Gaußsche Vermutung untersucht, daß $F(x)$ durch „den Integrallogarithmus“ $\int_2^x (\log t)^{-1} dt$ angenähert werden kann. Diese Arbeit ist berühmt, weil sie die sogenannte Riemannsche Vermutung enthält, daß die Eulersche Zetafunktion $\zeta(s)$ — diese Bezeichnung stammt von Riemann —, wenn man sie für komplexe Werte $s = x + iy$ betrachtet, sämtliche nicht reellen Nullstellen auf der Geraden $x = \frac{1}{2}$ besitzt. Diese Vermutung ist bis jetzt weder bewiesen noch als falsch nachgewiesen worden.⁴⁾

14. Riemanns Auffassung der Funktion einer komplexen Veränderlichen ist oft mit der von Weierstraß verglichen worden. Karl Weierstraß war viele Jahre hindurch Lehrer an einem preußischen Gymnasium und wurde 1856 Mathematikprofessor an der Universität Berlin, wo er dreißig Jahre lang lehrte. Seine stets auf sorgfältigste vorbereiteten Vorlesungen erfreuten sich wachsender Berühmtheit; hauptsächlich durch diese Vorlesungen sind die Gedanken von Weierstraß zum Gemeingut der Mathematiker geworden.

In seiner Gymnasiallehrerzeit schrieb Weierstraß mehrere Arbeiten über hyperelliptische Integrale, abelsche Funktionen und algebraische Differentialgleichungen. Seine am meisten bekannt gewordene Leistung ist die Begründung der Theorie der komplexen Funktionen durch die Methode der Potenzreihen. Das war in gewissem Sinne eine Rückkehr zu Lagrange, allerdings mit dem Unterschied, daß Weierstraß in der komplexen Ebene arbeitete und durchweg völlige Strenge erzielte. Die Werte der Potenzreihe innerhalb ihres Konvergenzkreises ergeben das „Funktionselement“, das dann, falls es

³⁾ Eine Kurve, die man mit der freien Hand zeichnen kann. Aus *Institutiones Calculi Integralis* III § 301.

⁴⁾ R. Courant, *Bernhard Riemann und die Mathematik der letzten hundert Jahre*, Naturwissenschaften 14, 813–818 (1926).

möglich ist, mit Hilfe der sogenannten analytischen Fortsetzung erweitert wird. Weierstraß studierte insbesondere ganze Funktionen und solche, die durch unendliche Produkte definiert sind. Seine elliptische Funktion $\wp(u)$ ist ebenso maßgeblich geworden wie die älteren $\operatorname{sn} u$, $\operatorname{cn} u$, $\operatorname{dn} u$ von Jacobi.

Der Ruhm von Weierstraß beruht auf seinen mit höchster Sorgfalt ausgeführten Schlußweisen, auf der „Weierstraßschen Strenge“, die nicht nur in seiner Funktionentheorie zutage tritt, sondern auch in seiner Variationsrechnung. Er klärte die Begriffe des Minimums, der Funktion und der Ableitung völlig auf und beseitigte damit die noch vorhandene Unbestimmtheit der Ausdrucksweise in den grundlegenden Begriffen der Infinitesimalrechnung. Er war das mathematische Gewissen schlechthin, sowohl in methodischer als auch in logischer Beziehung. Ein anderes Beispiel seiner peinlich genauen Denkweise ist seine Entdeckung der gleichmäßigen Konvergenz. Mit Weierstraß begann jene Zurückführung der Prinzipien der Analysis auf die einfachsten arithmetischen Begriffe, die wir die Arithmetisierung der Mathematik nennen.

„Wenn heute in Verfolgung der Schlußweisen, die auf dem Begriff der Irrationalzahl und überhaupt des Limes beruhen, in der Analysis volle Übereinstimmung und Sicherheit herrscht und in den verwickeltsten Fragen, die die Theorie der Differential- und Integralgleichungen betreffen, trotz der kühnsten und mannigfaltigsten Kombinationen unter Anwendung von Über-, Neben- und Durcheinanderhäufung der Limites doch Einhelligkeit aller Ergebnisse statthat, so ist dies wesentlich ein Verdienst der wissenschaftlichen Tätigkeit von Weierstraß.“⁵⁾)

...

⁵⁾ D. Hilbert, *Über das Unendliche*, Math. Annalen 95, 161–190 (1926).

Kapitel B

ANALYSIS I

B.1 Das Programm der Arithmetisierung der Analysis

Im 18. Jahrhundert wurde die Mathematik als eine Hilfswissenschaft der Naturwissenschaften angesehen. Eine mathematische Formel mußte in wahrnehmbarer, direkter Weise eine Interpretation in den Naturwissenschaften zulassen. Das Hauptkriterium für eine Sinnhaftigkeit einer mathematischen Aussage lag auf seiten des Gegenstands, da ihre Übereinstimmung mit den Objekten der Naturwissenschaft als entscheidend für die Gültigkeit der Aussage betrachtet wurde; in der Mathematik herrschte eine weitgehende Methodenfreiheit. — Der A-Zug dieser Vorlesung versucht einen Eindruck von dieser Auffassung zu vermitteln.

Im 19. Jahrhundert setzte ein Theoretisierungsprozeß ein, der gegen Ende des 19. Jahrhunderts zur Entstehung der „Reinen Mathematik“ führte. In den abgeschlankten, zur Prüfungsvorbereitung wohlgeeigneten Lehrbüchern (besonders beliebt ist O. FORSTER *Analysis I*) erfährt der Student leider kaum etwas davon, daß die Loslösung von empirischen Sachverhalten und die relative Eigenständigkeit mathematischer Theorie im 19. Jahrhundert getragen ist von erfolgreichen Versuchen, diese Theorien nicht nur auf die traditionellen Gebiete Mechanik und Astronomie zu beziehen, sondern die Anwendung auch auf andere Bereiche auszudehnen, die den ursprünglichen, bei der Formulierung der Theorien intendierten Anwendungen fremd sind. Eine der ersten bedeutenden Arbeiten dieser Art stellt die im Jahre 1822 von J.B.J. FOURIER vorgelegte „Theorie analytique de la Chaleur“ dar, in der die Theorie der Wärme hauptsächlich mit den Methoden der Analysis entwickelt wurde.

In dem im 19. Jahrhundert sich herausbildenden Verständnis von Mathematik wird die Wahrheit einer mathematischen Aussage nach und nach weniger an der Widerspiegelung eines physikalischen oder empirischen Sachverhalts gemessen. Statt dessen erhalten theorieinterne Kriterien zur Feststellung der Wahrheit einer Theorie eine immer größere Bedeutung.

Am Wendepunkt zur Mathematikauffassung des 19. Jahrhunderts steht C.F. GAUSS (1777–1832). K. REIDEMEISTER (1893–1971), (als Mathematiker insbesondere durch seine Beiträge zur Knotentheorie bekannt), hat in einem Aufsatz über Gauß (Springer-Verlag 1957) aus philosophischer Sicht dargestellt, was die ganz wesentlich von Gauß geprägte Mathematik

des 19. Jahrhunderts von der älteren Mathematik unterscheidet.

Das 17. und 18. Jahrhundert ist die Epoche der mathematischen Erfindung. Sie steht unter der Vorherrschaft des Kalküls und versteht das exakte Denken als einen Prozeß, der durch den rechten Kalkül befestigt und vorwärts getrieben werden soll. Dieser Gedanke bestimmt sowohl die Wendung Descartes' zur analytischen Geometrie, mit deren Hilfe die Schlüsse Euklids durch den Kalkül der Algebra ersetzt werden, wie Leibniz' Konstruktion der Differentialrechnung, durch die eine Fülle von Einzelüberlegungen methodisch zu einem Kalkül zusammengefaßt wird. Das Imaginäre und die unendlichen Reihen fügen sich in diesem Zusammenhang so ohne Anstoß ein, weil sich mit ihnen Rechnungen durchführen lassen, welche die Regelmäßigkeit eines Kalküls besitzen.

Diesen Erfindungen stellt Reidemeister die „Entdeckungen“ von Gauß in der Zahlentheorie gegenüber. (Die Zahlen sind uns direkt, und nicht auf irgendeinem Umweg über die Naturwissenschaften gegeben; und im Reich der Zahlen gibt es Entdeckungen zu machen). Nur am Rande erkennt Reidemeister an, daß die Vielseitigkeit und die Intensität der Durchdringung exakter und praktischer Interessen, die sich in Gauß' Arbeiten über Geodäsie, Astronomie und Physik bekundet, kaum ihresgleichen kennt.

Als Advokaten der Angewandten Mathematik geben wir Reidemeister insofern recht, daß die Leistungen von Gauß in der reinen Mathematik (insbesondere in der Zahlentheorie) sehr viel schneller Nachfolger angeregt (und somit eine neue Epoche eingeleitet) haben als seine Leistungen in der angewandten Mathematik. Die Ansätze von Gauß zu einer eigenständigen Angewandten Mathematik (etwa in der Wahrscheinlichkeitstheorie) fielen nicht sofort auf fruchtbaren Boden. Zu den Schwierigkeiten, die sich im 19. Jahrhundert einer neuen Auffassung von Angewandter Mathematik entgegenstellten, siehe G. RICHENHAGEN: „*Carl Runge (1856–1927), von der reinen Mathematik zur Numerik*“, Vandenhoeck & Ruprecht, 1985.

Es scheint, daß Gauß mit seinen Beiträgen zur Angewandten Mathematik seinen Zeitgenossen noch weiter voraus war als mit seinen zahlentheoretischen Entdeckungen.

Bei den alten Griechen war die Mathematik eng mit der Philosophie verbunden. Den Griechen galt das mathematische Wissen wegen seiner **Begründbarkeit** aus Axiomen als ein Wissen ganz besonderer Art. Im Gefolge der Entdeckung des Irrationalen („Die Diagonale des Quadrats hat kein gemeinsames Maß mit der Seite“) waren sie zum Schluß gekommen, daß Arithmetik und Geometrie als getrennte Bereiche zu sehen sind. Das die Arithmetik beherrschende Diskrete („Alles ist Zahl“) sahen sie als unverträglich an mit dem Kontinuierlichen; sie hatten große philosophische Probleme mit der Realität der Veränderungen. Im 17. Jahrhundert waren die Autoren bereit, die archimedische Strenge zugunsten von Gedankengängen aufzugeben, die auf unstrengen Voraussetzungen beruhend schnell zu Ergebnissen führten (vgl. Struik, S. 106 ff). Im 18. Jahrhundert kam die Entwicklung zu einer gewissen Sättigung. Im 19. Jahrhundert stellten sich die Mathematiker nun die Aufgabe, die bei der Beschreibung von Bewegungen so erfolgreiche Infinitesimalrechnung zu begründen, und zwar im Sinne einer Arithmetisierung. Dazu mußte nicht nur der Zahlbegriff erweitert werden; es mußten auch die Vorstellungen von der Existenzweise (von der Anschaulichkeit und von

der Beschreibbarkeit) der mathematischen Objekte revidiert werden. Im Aufsatz „*Eine Begründung der Infinitesimalrechnung*“ (in der Sammlung „Raum und Zahl“, Springer 1957) schreibt Reidemeister:

„Die Hauptschwierigkeit der Infinitesimalrechnung ist die Einführung der reellen Zahlen. Dort liegt das Kernproblem der mathematisch-logischen Begründung, dort liegt die erkenntnistheoretische Schwierigkeit, die man kurz das Problem der reinen Anschauung nennt, dort liegt aber auch die Schwierigkeit für den Lernenden und Lehrenden: Wie genau soll ein Beweis geführt werden, an welcher Stelle soll der kritische Sinn geweckt werden? Was ist bekannt und sicher und was nur scheinbar bekannt und unsicher?

Die Entdeckung des Differentialquotienten und des Integrals erfolgte zu einer Zeit, als man in der Mathematik nicht „more geometrico“ dachte. Das Vertrauen zu der neuen Methode stützte sich auf ihren Erfolg; die logischen und erkenntnistheoretischen Zweifel beschieden sich und überließen der formalen Phantasie die Zügel. Erst nach etwa 100 Jahren tritt hierin eine Wandlung ein. CAUCHY und DEDEKIND bezeichnen die zwei klassischen Etappen der kritischen Besinnung; CAUCHY untersucht den Grenzwertbegriff, DEDEKIND führt den Begriff des Schnitts – ganz ähnlich wie es EUKLID schon getan hatte – ein; alsbald entsteht aber jene Krise der mathematischen Logik, von der heute so viel gesprochen wird und die auch den DEDEKINDschen Begriff der reellen Zahlen erschütterte.“ ...

„Die Grundelemente der Infinitesimalrechnung sind die Funktionen von einer Veränderlichen $y = f(x)$. Der anschauliche Anlaß, solche Funktionen zu konstruieren, sind Kurven, die in einem rechtwinkligen Koordinatensystem festgelegt werden sollen. Das beruht auf der Vorstellung, daß ein Punkt durch Abszisse und Ordinate bestimmt ist und diese durch Maßzahlen gemessen werden können. Kurz, wir haben zur Definition einer Funktion zunächst einen Bereich von Zahlen nötig. Beispiele von Zahlen sind bekannt; die natürlichen Zahlen, die rationalen Zahlen – wir lassen dahingestellt, ob es noch weitere Zahlen gibt, und schließen diese Möglichkeit jedenfalls nicht aus. Für die rationalen Zahlen gibt es zwei Rechenoperationen, die Addition und Multiplikation, die einigen Rechenregeln genügen, welche aus dem sogenannten Buchstabenrechnen bekannt sind. Ferner lassen sich die rationalen Zahlen anordnen, und die Größenrelation ist durch die sogenannten Monotoniegesetze mit Addition und Multiplikation verknüpft. Diese Tatsachen sprechen wir kurz so aus: Die rationalen Zahlen bilden einen geordneten Zahlkörper. Unsere erste Grundannahme sei nun:

Es gibt einen geordneten Zahlkörper, welcher die rationalen Zahlen enthält und hinreicht, um die anschaulichen Kurven durch geeignete streng definierte Gebilde zu ersetzen.“

Das, was Reidemeister hier als Grundannahme bezeichnet, ist die Quintessenz aus Einsichten aus der zweiten Hälfte des 19. Jahrhunderts. Die Menge $\mathbb{R}$ der reellen Zahlen, genauer gesagt, der (im Sinne von DEDEKIND) vollständig angeordnete Körper der reellen Zahlen

$(\mathbb{R}, 0, 1, +, \cdot, \leq)$ wird axiomatisch bestimmt. Zur Bestimmung des Gegenstandsbereichs der mathematischen Analysis hören wir nochmals auf Reidemeister (im Artikel „Geometrie und Zahlentheorie“):

„Die Gegenstände, die der mathematischen Analysis als analytisches Äquivalent von Kurven in der Ebene und Flächen im Raum zuwachsen, sind die Funktionen, die als Zuordnung von Zahlen zu Zahlen oder von Zahlen zu Zahlenpaaren analytisch erklärt und unter recht allgemeinen Voraussetzungen über den Charakter dieser Zuordnung den Methoden der mathematischen Analysis unterworfen werden können. Das kommt dann wieder der Geometrie zugute, die sich nun nicht mehr auf die Geraden, Kreise und Kegelschnitte zu beschränken braucht, sondern auch Kurven und Flächen behandeln kann, die den allgemeinen Funktionen der Analysis entsprechen. So entsteht z.B. die Differentialgeometrie der Flächen und daraus dann die Theorie der gekrümmten Räume, die das mathematische Gerüst für die allgemeine Relativitätstheorie ist.“

Wir halten fest: Um eine mathematisch strenge Theorie der anschaulichen Kurven zu entwickeln, braucht man eine **Erweiterung des Zahlbegriffs** über den Begriff der rationalen Zahlen hinaus. Die für die mathematische Analysis passende Erweiterung kann streng gefaßt werden durch das Vollständigkeitsaxiom von Dedekind. Dieses Axiom stützt sich auf den Begriff des Schnitts in der linear angeordneten Menge der rationalen Zahlen. Dedekinds Zugang soll hier kurz skizziert werden. (Der Begriff der Vollständigkeit wird uns in anderer Form im Kapitel über die Vervollständigung eines metrischen Raums wieder begegnen.)

Bekanntlich gilt für jedes Paar rationaler Zahlen a, b genau eine der Relationen $a < b$ oder $b < a$ oder $a = b$. Wir schreiben $a \leq b$, wenn $a < b$ oder $a = b$ gilt, so daß also $a = b \iff (a \leq b \text{ und } b \leq a)$, $a < b \iff (a \leq b \text{ und } a \neq b)$.

Definition Ein **Dedekindscher Schnitt** ist eine Partition von $\mathbb{Q}$ in zwei nichtleere Mengen $(\mathbb{Q} = M_u \cup M_0, M_u \cap M_0 = \emptyset)$, so daß für jedes $a \in M_u$ und jedes $b \in M_0$ gilt $a \leq b$.

Bemerke Der Begriff des Dedekindschen Schnitts kann auf beliebige linear angeordnete Mengen $(M, \leq)$ übertragen werden.

Definition Eine zweistellige Relation „ $\leq$ “ auf einer Menge $(M, =)$ heißt eine *partielle Ordnung*, wenn

$$\begin{aligned} a &\leq a && \text{für alle } a \in M \\ a \leq b, b \leq c &&& \text{impliziert } a \leq c \\ a \leq b, b \leq a &&& \text{impliziert } a = b. \end{aligned}$$

Man spricht von einer *linearen Ordnung*, wenn für $a \neq b$ entweder $a \leq b$ oder $b \leq a$ gilt.

Definition Sei $(M, \leq)$ eine lineargeordnete Menge. Ein *Schnitt* ist eine Partition

$$M = M_u + M_0 \text{ mit } a \leq b \text{ für alle } a \in M_u, b \in M_0.$$

Definition Eine lineargeordnete Menge $(M, \leq)$ heißt *vollständig im Sinne von Dedekind*, wenn es zu jedem Schnitt $M = M_u + M_0$ ein Element $a^* \in M$ gibt, so daß

$$a \leq a^* \text{ für alle } a \in M_u, \quad a^* \leq b \text{ für alle } b \in M_0.$$

Bemerke

- 1) Die Menge $\mathbb{Z}$ der ganzen Zahlen mit der üblichen Ordnung ist vollständig im Sinne von Dedekind.
- 2) Die Menge $\mathbb{Q}_+$ ist nicht vollständig. Sei z.B.

$$M_0 = \{r : r \text{ rational, } r^2 \geq 2\}, \quad M_u = \mathbb{Q}_+ \setminus M_0.$$

Dann existiert keine rationale Zahl a^* , die diesen Schnitt erzeugt; $\sqrt{2}$ ist eben keine rationale Zahl.

- 3) Jeder, der sich die Menge $\mathbb{R}$ der reellen Zahlen als die Menge der Punkte auf der Zahlengeraden vorstellt, wird zugeben, daß $\mathbb{R}$ vollständig ist im Sinne von Dedekind. Man kann die reellen Zahlen in der Tat identifizieren mit der Menge aller Schnitte in $(\mathbb{Q}, \leq)$. Es ist das Verdienst von Dedekind, daß er dieser intuitiven Einsicht eine scharfe Fassung und damit dem System der reellen Zahlen eine axiomatische Fundierung gegeben hat.

Wir zitieren nun aus C.H. EDWARDS, JR.: *The historical development of the calculus*. Springer Verlag 1979. (Unsere freie Übersetzung)

Im Kapitel „The arithmetization of Analysis“ heißt es:

Der Kalkül von Newton und Leibniz war ein Kalkül geometrischer Variabler, von Größen, die explizit mit geometrischen Kurven verknüpft sind, und ihre Analyse beruhte weitgehend auf intuitiven geometrischen Begriffen. Euler, Lagrange und Cauchy versuchten in ihrer Begründung der Analysis die geometrische Intuition durch die Prinzipien der Arithmetik zu ersetzen; in ihren Büchern über Infinitesimalrechnung findet sich kein einziges geometrisches Diagramm. Diese ersten Versuche einer Arithmetisierung des Gebiets waren jedoch nur teilweise erfolgreich, weil bis zum Ende des 19. Jahrhunderts die reellen Zahlen selbst nur intuitiv verstanden wurden.

Seit dem 17. Jahrhundert hatten die Mathematiker die Irrationalzahlen (wie etwa $\sqrt{2}$) pragmatisch und unkritisch verwendet, ohne ernsthaft nach ihrem präzisen Wesen zu fragen, wobei sie für das Rechnen davon ausgingen, daß jede Irrationalzahl beliebig genau durch rationale Zahlen approximiert

werden kann (z.B. $\sqrt{2} = 1.41421\dots$). Dabei nahm man nicht nur an, daß die Irrationalzahlen, die man für die üblichen Zwecke brauchte, existieren, sondern auch, daß für sie die von den Rationalzahlen her bekannten Regeln der algebraischen Operationen gelten. Es gab aber, wie RICHARD DEDEKIND (1831 – 1916) in einem Aufsatz 1872 bemerkte, keine strenge Begründung für so einfache Tatsachen wie etwa $\sqrt{2} \cdot \sqrt{3} = \sqrt{6}$.

Da das System der reellen Zahlen nicht völlig verstanden wurde, konnte man den Kalkül nicht solide begründen. ...“

Edwards führte als Beispiele für Beweislücken, die hierauf zurückzuführen sind, auf:

- i) Bolzanos Beweis des Zwischenwertsatzes,
- ii) Cauchys Kriterium für die Konvergenz von Zahlenfolgen,
- iii) die von Cauchy und Riemann implizit verwendeten Annahmen bei Beweisen, daß Funktionen unter geeigneten Voraussetzungen ein Integral besitzen.

Grundlegende Eigenschaften der reellen Zahlen wurden nicht verifiziert, sondern aus geometrischen Gründen für evident gehalten. Edwards führt weiter aus:

„Diese Vagheit in der Begründung der Analysis führte nicht nur zu Lücken in der Beweisführung, sondern gelegentlich auch zu falschen Vorstellungen. Z.B. glaubte man im frühen 19. Jahrhundert ganz allgemein, daß jede stetige Funktion allenfalls in isolierten Punkten nichtdifferenzierbar sein könne (wie z.B. die Funktion $f(x) = |x|$ im Nullpunkt nicht differenzierbar ist). In der Tat taten einige Analysistexte so, als ob sie diese falsche Aussage bewiesen. Es kam daher als heilsamer Schock, als KARL WEIERSTRASS (1815–1897) 1861 in seinen Berliner Vorlesungen eine Funktion angab, die überall stetig aber nirgends differenzierbar ist. ...“

Zwar hatte Bolzano schon 1834 eine stetige nichtdifferenzierbare Funktion angegeben; das war aber nicht beachtet worden. Erst das Beispiel von Weierstraß machte klar, daß die Grundlagen der Analysis neu zu überdenken waren. Insbesondere wurde klar, daß man das System der reellen Zahlen nicht als gegeben ansehen darf, sondern daß man es in einer Weise konstruieren oder definieren muß, daß man die Eigenschaften der Irrationalzahlen streng beweisen kann. ...“

R. DEDEKIND (1831–1916) kann als Vollender des Programms der Arithmetisierung der Analysis gelten. Ihm ist die Einsicht zu verdanken, daß die Analysis keinerlei Anleihen bei der geometrischen Intuition nötig hat, wenn man die Vollständigkeit des Systems der reellen Zahlen als Axiom einführt, d.h. wenn man zu den Axiomen, die einen angeordneten Körper charakterisieren, die Vollständigkeit als weiteres Axiom hinzufügt. Dedekinds Einsicht soll durch die Beweise, die wir im B-Zug unserer Einführung führen werden, plausibel gemacht werden. Es kann aber nicht unser Ziel sein, eine streng arithmetisierte Analysis von Grund auf

zu entwickeln. Wir halten es mit Dedekind, der aus didaktischen Gründen den gelegentlichen Rekurs auf die Anschauung empfiehlt: Um nicht allzuviel Zeit zu verlieren, müsse man wohl an die geometrische Intuition appellieren, wenn man in einer ersten Einführung solche Tatsachen vermitteln will, wie z.B. daß jede monoton steigende Zahlenfolge, die nicht über alle Grenzen wächst, einen Grenzwert besitzt.

Nach dem Gesagten sollte klar sein, inwiefern der folgende (hier nicht bewiesene) Satz als das Fundament des B-Zugs unserer Vorlesung anzusehen ist.

Axiomatische Kennzeichnung des Körpers der reellen Zahlen Es existiert (bis auf anordnungserhaltende Isomorphie) genau ein im Sinne von Dedekind vollständiger angeordneter Körper

$$(\mathbb{R}, \leq, 0, 1, +, \cdot) .$$

Wir zitieren aus einigen bekannten Lehrbüchern detaillierte Axiomensysteme für diesen vollständigen angeordneten Körper, den Körper der reellen Zahlen.

Gesetze der Gleichheit

- G_1 (Reflexivität) $\bigwedge_a a = a$
 G_2 (Komparativität) $\bigwedge_{a,b,c} (a = c \wedge b = c) \rightarrow a = b$

Gesetze der Addition

- A_1 (Verknüpfung) $\bigwedge_{a,b} \bigvee_c^1 a + b = c$
 A_2 (Konsistenz) $\bigwedge_{a,b,c,d} (a = b \wedge c = d) \rightarrow a + c = b + d$
 A_3 (Assoziativität) $\bigwedge_{a,b,c} (a + b) + c = a + (b + c)$
 A_4 (Streichregel) $\bigwedge_{a,b,c} a + c = b + c \rightarrow a = b$
 A_5 (Kommutativität) $\bigwedge_{a,b} a + b = b + a$
 A_6 (Neutrales Element 0 oder Null) $\bigvee_0^1 \bigwedge_a a + 0 = 0 + a = a$
 A_7 (Subtraktion) $\bigwedge_a \bigvee_b^1 a + b = 0 = b + a$. Dieses wird mit $-a$ bezeichnet.

Anmerkung Ein Teil der vorstehenden Gesetze ist in verkürzter Form geschrieben. Z.B. ist $a + b = b + a$ in A_5 keine Grundaussage. Im Sinne des eingangs erklärten Rechnens wäre A_5 vielmehr in der Form $\bigwedge_{a,b,c} a + b = c \rightarrow b + a = c$ zu schreiben, worin nur Grundaussagen vorkommen. Analog für andere Fälle.

Gesetze der Multiplikation

- M_1 (Verknüpfung) $\bigwedge_{a,b} \bigvee_c^1 ab = c$
 M_2 (Konsistenz) $\bigwedge_{a,b,c,d} (a = b \wedge c = d) \rightarrow ac = bd$
 M_3 (Assoziativität) $\bigwedge_{a,b,c} (ab)c = a(bc)$
 M_4 (Streichregel) $\bigwedge_{a,b,c} (ac = bc \wedge c \neq 0) \rightarrow a = b$
 M_5 (Kommutativität) $\bigwedge_{a,b} ab = ba$
 M_6 (Neutrales Element 1 oder Eins) $\bigvee_1^1 \bigwedge_a a1 = 1a = a$
 M_7 (Division) $\bigwedge_a a \neq 0 \rightarrow \bigvee_b^1 ab = 1 = ba$. Dieses b wird mit a^{-1} bezeichnet
 M_8 (Distributivität) $\bigwedge_{a,b,c} a(b+c) = ab+ac \wedge (b+c)a = ba+ca$

Anmerkung 1 Statt ab schreibt man auch $a \cdot b$

Anmerkung 2 Summen bzw. Produkte aus n gleichen Summanden bzw. Faktoren a werden mit na oder $n \cdot a$ bzw. mit a^n bezeichnet (also $na := a_1 + \dots + a_n$ mit $\bigwedge_\nu a_\nu = a$ usw.)

Gesetze der Ordnung (Relation $<$). Es wird gesetzt $c > d :\leftrightarrow d < c$. Dann sei $a \leq b :\leftrightarrow a < b \vee a = b$ und $c \geq d :\leftrightarrow d \leq c$ sowie $b \not\leq a :\leftrightarrow \neg(b < a) \leftrightarrow a \not\leq b$

- U_1 (Konsistenz) $\bigwedge_{a,b,c,d} (a = b \wedge c = d \wedge a < c) \rightarrow b < d$
 U_2 (Einsinnigkeit) $\bigwedge_{a,b} a \leq b \rightarrow b \not\leq a$
 U_3 (Transitivität) $\bigwedge_{a,b,c} (a < b \wedge b < c) \rightarrow a < c$
 U_4 (Linearität) Für beliebige a, b liegt genau einer der drei Fälle vor:
 $a = b$ oder $a < b$ oder $b < a$, d.h. $\bigwedge_{a,b} a \neq b \rightarrow (a < b \vee b < a)$
 U_5 (Monotonie der Summe) $\bigwedge_{a,b,c} a < b \rightarrow a + c < b + c$
 U_6 (Monotonie des Produktes) $\bigwedge_{a,b,c} (a < b \wedge 0 < c) \rightarrow ac < bc$
 U_7 (Archimedisches Axiom) $\bigwedge_{a,b} (0 < a \wedge 0 < b) \rightarrow (\bigvee_n n \in \mathbb{N}^+ \wedge b < na)$
 U_8 (Vollständigkeit (im Sinne der Ordnung) von M nach oben)

Jede nicht leere nach oben beschränkte Teilmenge A von M besitzt in M eine kleinste obere Schranke, d.h. es existiert in M das Supremum $\sup A$ von A , in Zeichen:

$$\bigwedge_{A \subset M} \bigwedge_{b \in M} (A \neq \emptyset \wedge A \leq b) \rightarrow \bigvee_{c \in M} (A \leq c) \wedge \bigwedge_{d \in M} (A \leq d \rightarrow c \leq d)$$

Aus Aumann–Haupt: Einführung in die reelle Analysis, de Gruyter, 1974.

Zusammenstellung der Axiome der reellen Zahlen

Körperaxiome: Es sind zwei Verknüpfungen (Addition und Multiplikation) definiert, so daß folgende Axiome erfüllt sind.				
Axiome der Addition	Axiome der Multiplikation			
Assoziativgesetz	Assoziativgesetz			
Kommutativgesetz	Kommutativgesetz			
Existenz der Null	Existenz der Eins ($\neq 0$)			
Existenz des Negativen	Existenz des Inversen (zu Elementen $\neq 0$)			
Distributivgesetz				
Anordnungsaxiome : Es sind gewisse Elemente als positiv ausgezeichnet ($x > 0$), so daß folgende Axiome erfüllt sind.				
Für jedes Element x gilt genau eine der Beziehungen $x > 0$, $x = 0$, $-x > 0$				
$x > 0$ & $y > 0 \Rightarrow x + y > 0$				
$x > 0$ & $y > 0 \Rightarrow xy > 0$				
Archimedisches Axiom: Zu $x > 0$, $y > 0$ existiert eine natürliche Zahl n mit $nx > y$				
Vollständigkeitsaxiom: Jede Cauchy-Folge konvergiert				

Die Axiome bis einschließlich Distributivgesetz definieren einen Körper. Die weiteren ohne die beiden letzten Zeilen einen angeordneten Körper.

Zur Charakterisierung des Körpers der reellen Zahlen kommen noch die beiden letzten Axiome hinzu. Die reellen Zahlen bilden einen vollständigen archimedisch angeordneten Körper.

Bemerkung In vielen Lehrbüchern wird das sog. Dedekindsche Schnittaxiom verwendet. Dieses ist gleichwertig mit dem Archimedischen und dem Vollständigkeits-Axiom und kann deshalb anstelle dieser beiden Axiome verwendet werden. Ebenfalls gleichwertig mit dem Archimedischen und Vollständigkeits-Axiom ist der Satz, daß jede nicht-leere, nach oben beschränkte Menge von reellen Zahlen eine kleinste obere Schranke besitzt. Also kann man auch diesen Satz statt der beiden Axiome selbst als Axiom verwenden.

Aus O. Forster: Analysis 1, Vieweg-Verlag

B.2 Wurzelziehen

Mit Hilfe des Dedekindschen Schnittaxioms sieht man leicht, daß es zu jeder positiven Zahl x genau eine positive Zahl $\sqrt{x} = \sup\{r : r^2 < x\}$ gibt, deren Quadrat gleich x ist. Es gibt viele Algorithmen, diese Zahl beliebig genau zu approximieren. Wir diskutieren hier einen beliebten Algorithmus, der auf der Einsicht beruht, daß das geometrische Mittel zweier positiver Zahlen a, b nie größer ist als das arithmetische Mittel, und daß Gleichheit nur herrscht, wenn $a = b$.

Satz $\sqrt{a \cdot b} \leq \frac{1}{2}(a + b)$ für alle $a, b > 0$.

Beweis Wir müssen zeigen $a \cdot b \leq \frac{1}{4}(a + b)^2$. Dies ergibt sich aus

$$\frac{1}{4}(a + b)^2 - a \cdot b = \frac{1}{4}(a - b)^2 \geq 0.$$

$\sqrt{x}$ ist das geometrische Mittel von x und 1 oder auch, für beliebiges $y > 0$, das geometrische Mittel von $\frac{x}{y}$ und y .

Algorithmus Beginne mit irgendeiner Zahl $y_0 > 0$ und setze

$$y_1 = \frac{1}{2} \left(y_0 + \frac{x}{y_0} \right), \quad y_{n+1} = \frac{1}{2} \left(y_n + \frac{x}{y_n} \right) \quad \text{für } n = 1, 2, \dots$$

Es gilt dann

$$y_1 \geq y_2 \geq \dots \quad \text{und} \quad \lim \searrow y_n = \sqrt{x}.$$

Beweis

- 1) Der Beweis beruht auf einer Idee, auf die der Anfänger wohl kaum selber kommen würde. Man betrachtet (für ein festes $x > 0$) die Funktion

$$g(y) := \frac{1}{2} \left(y + \frac{x}{y} \right) \quad \text{für } y > 0.$$

- 2) Wir beweisen einige Eigenschaften dieser Funktion. Es gilt $g(y) \geq \sqrt{x}$ und $g(y) \leq y$ für alle $y \geq \sqrt{x}$; denn

$$g(y) - y = \frac{1}{2} \left(-y + \frac{x}{y} \right) = \frac{1}{2y} (-y^2 + x).$$

Also haben wir

$$y_1 \geq y_2 \geq \dots \quad \text{mit} \quad y_n \geq \sqrt{x} \quad \text{für alle } n$$

$$\frac{x}{y_1} \leq \frac{x}{y_2} \leq \dots \quad \text{mit} \quad \frac{x}{y_n} \leq \sqrt{x} \quad \text{für alle } n$$

3) Sei $y^* = \inf\{y_n : n \in \mathbb{N}\}$. Dann gilt einerseits

$$\left(y_n - \frac{x}{y_n}\right) \searrow y^* - \frac{x}{y^*} \geq 0$$

und andererseits aufgrund der Konstruktion von y_{n+1}

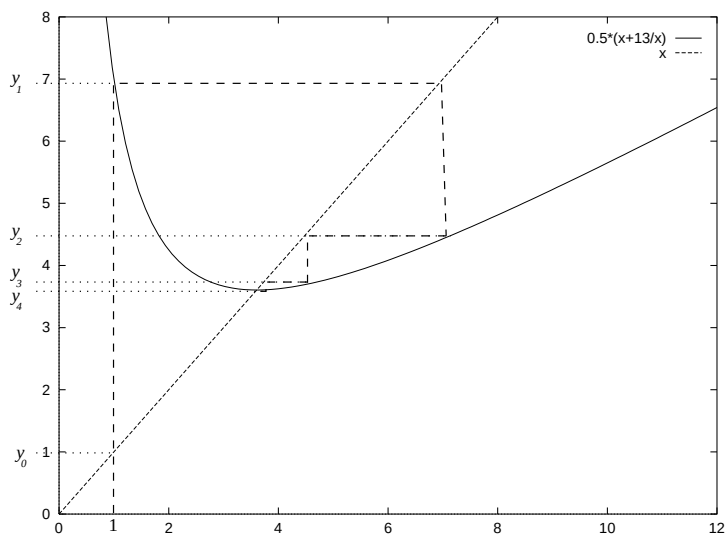
$$\left(y_{n+1} - \frac{x}{y_{n+1}}\right) = \frac{1}{2} \left(y_n - \frac{x}{y_n}\right) \searrow \frac{1}{2} \left(y^* - \frac{x}{y^*}\right).$$

Daraus folgt $y^* = \frac{x}{y^*}$. ■

Figur B2.1 Approximation von $\sqrt{13}$ mit $y_0 = 1$

$$y_0 = 1, y_1 = \frac{1}{2} \left(y_0 + \frac{13}{y_0}\right), \dots, y_{n+1} = \frac{1}{2} \left(y_n + \frac{13}{y_n}\right), \dots$$

$$g(y) = \frac{1}{2} \left(y + \frac{13}{y}\right) \quad \text{für } y > 0$$



Approximation der Wurzelfunktion

Uns interessieren nicht nur Zahlenfolgen, sondern auch **Funktionenfolgen**. Wir wollen die Wurzelfunktion $\sqrt{x}$ approximieren. Betrachten wir also x als eine Variable mit Werten > 0 .

Beginnen wir z.B. mit der Funktion $f_0(x) \equiv 1$ und setzen wir

$$f_1(x) = \frac{1}{2} \left(f_0(x) + \frac{x}{f_0(x)}\right), f_{n+1}(x) = \frac{1}{2} \left(f_n(x) + \frac{x}{f_n(x)}\right).$$

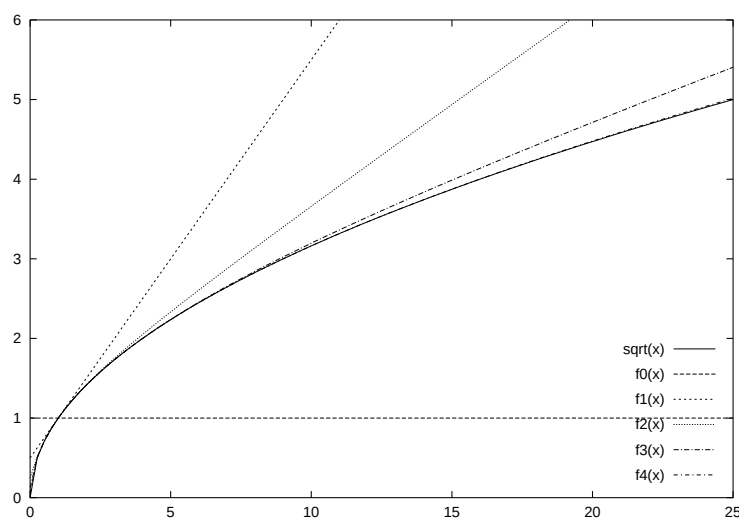
Es gilt dann $f_n(x) \geq \sqrt{x}$ für alle n und

$$f_1(x) \geq f_2(x) \geq \dots \quad \text{mit} \quad \inf\{f_n(x) : n \in \mathbb{N}\} = \sqrt{x}.$$

Man schreibt auch: $\sqrt{x} = \lim \searrow f_n(x)$.

Figur B2.2 Approximation von $\sqrt{x}$ für $0 < x \leq 25$

$$f_0(x) = 1, \quad f_1(x) = \frac{1}{2}(1+x), \quad f_2(x) = \frac{1}{2}\left(f_1(x) + \frac{x}{f_1(x)}\right), \dots$$



Aufgabe 1 Sei $x > 0$. Setze $y_0 = 1, y_{n+1} = \frac{1}{3}\left(2y_n + \frac{x}{y_n^2}\right)$. Studiere das Verhalten der Folge.

Aufgabe 2

a) Zeige, daß die Zahlenfolge

$$\sqrt{2}, \sqrt{2 + \sqrt{2}}, \sqrt{2 + \sqrt{2 + \sqrt{2}}}, \dots$$

monoton ansteigt und berechne ihr Supremum.

b) Für $x > 0$ betrachte

$$f_1(x) = \sqrt{x}, \quad f_{n+1}(x) = \sqrt{x + f_n(x)} \quad \text{für } n = 1, 2, \dots$$

Zeige $f_1 \leq f_2 \leq f_3 \leq \dots$ und berechne die Grenzfunktion $f^*(\cdot)$.

Aufgabe 3 Betrachte die für $x > 0$ definierten Funktionen

$$f_1(x) = x, \quad f_2(x) = x + \frac{1}{f_1(x)}, \quad f_3(x) = x + \frac{1}{f_2(x)}, \quad \dots$$

Zeige

$$f_1 \leq f_3 \leq f_5 \leq \dots \leq f_6 \leq f_4 \leq f_2$$

und berechne

$$\lim_{n \rightarrow \infty} \uparrow f_{2n+1}(x), \quad \lim_{n \rightarrow \infty} \downarrow f_{2n}(x).$$

(In der Theorie der Kettenbrüche notiert man die Grenzfunktion

$$f^*(x) = x + \frac{1}{\left| \frac{1}{x} \right|} + \frac{1}{\left| \frac{1}{x} \right|} + \frac{1}{\left| \frac{1}{x} \right|} + \dots)$$

Die Folge der Fibonaccischen Zahlen $0, 1, 1, 2, 3, 5, 8, 13, 21, \dots$ entstand aus folgender

Aufgabe 4 Wieviele Kaninchenpaare können in einem Jahr von einem einzigen Paar erzeugt werden, wenn a) jedes Paar in jedem Monat ein neues Paar erzeugt, das vom zweiten Monat an selbst wieder neue Paare erzeugt, b) keine Todesfälle vorkommen?

In moderner Notation lautet die

Aufgabe 4a) $y_0 = 0, y_1 = 1, y_{n+1} = y_n + y_{n-1}$ für $n = 1, 2, \dots$
Wie schnell steigt die Folge an?

Lösung Setze $x_1 = 1, x_n = \frac{y_{n+1}}{y_n}$ für $n = 1, 2, \dots$

Es gilt dann $x_1 = 1, x_2 = 2, x_3 = \frac{3}{2}, x_4 = \frac{5}{3}, \dots$

$$x_{n+1} = 1 + \frac{1}{x_n} \quad \text{für } n = 1, 2, \dots$$

$$x_{n+2} = 1 + \frac{1}{1 + \frac{1}{x_n}}$$

$$x_1 < x_3 < x_5 < \dots < x_6 < x_4 < x_2.$$

(Man beweise das durch vollständige Induktion !)

a) $x_{2n+1} < x_{2n}$ für alle n

b) $x_{2n+2} = 1 + \frac{1}{1 + \frac{1}{x_{2n+1}}} > 1 + \frac{1}{1 + \frac{1}{x_{2n}}} = x_{2n+1}.$

Es gilt $\sup\{x_{2n+1}\} = \inf\{x_{2n}\} = x^*$ mit

$$x^* = 1 + \frac{1}{x^*} \quad x^* = \frac{1}{2} (1 + \sqrt{5})$$

Dies ist die Zahl des Goldenen Schnitts: Eine Strecke soll so geteilt werden, daß der kleinere Abschnitt zum größeren im selben Verhältnis steht wie der größere zur ganzen Strecke.

$$\frac{a}{b} = \frac{b}{a+b}; \quad \frac{a+b}{b} = \frac{b}{a}.$$

Mit $y = \frac{b}{a}$ also $y = 1 + \frac{1}{y}.$

Anhang : Zur Geschichte der Zahlen $\sqrt{2}$, π und e

- 1) Wir zitieren aus D.J. STRUIK: Abriß der Geschichte der Mathematik. VEB Deutscher Verlag der Wissenschaften, Berlin 1963, S. 39:

Die wichtigste den Pythagoreern zugeschriebene Entdeckung war aber die Entdeckung des Irrationalen mit Hilfe von inkommensurablen Strecken. Diese Entdeckung kann das Ergebnis ihres Interesses für das geometrische Mittel $a : b = b : c$ gewesen sein, das als Symbol der Aristokratie galt. Wie groß war das geometrische Mittel von 1 und 2, von zwei heiligen Symbolen? Das führte zum Studium des Verhältnisses von Seite und Diagonale eines und desselben Quadrates, und man fand, daß dieses Verhältnis nicht durch „Zahlen“ ausgedrückt werden konnte — d.h. durch positive ganze Zahlen und ihre Verhältnisse, die in der Pythagoreischen Arithmetik allein zugelassenen Begriffe.

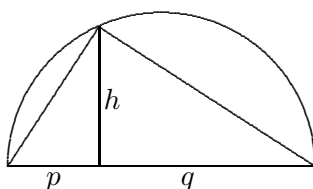
Angenommen, dieses Verhältnis sei $p : q$, wobei man stets p und q als teilerfremd voraussetzen kann. Daraus folgt $p^2 = 2q^2$, also ist p^2 und damit auch p eine gerade Zahl, etwa $p = 2r$. Dann müßte q ungerade sein, aber wegen $q^2 = 2r^2$ müßte es zugleich gerade sein. Dieser Widerspruch wurde nicht, wie im Orient oder in Europa im Zeitalter der Renaissance, durch eine Verallgemeinerung des Zahlbegriffs aufgelöst, sondern unter Ablehnung einer auf Zahlen beruhenden Theorie durch den Versuch einer geometrischen Synthese. Diese Entdeckung, die die wohlabgestimmte Harmonie zwischen Arithmetik und Geometrie zerstörte, wurde wahrscheinlich in den letzten Jahrzehnten des fünften Jahrhunderts v.u.Z. gemacht.

- 2) Auf Seite 91 schreibt Struik über das Interesse an den Wurzeln im Mittelalter:

Leonardo (von Pisa), auch Fibonacci („Sohn des Bonaccio“) genannt, reiste als Kaufmann in den Orient. Nach seiner Rückkehr schrieb er das „Liber Abaci“ (1202), das arithmetische und algebraische Unterweisungen enthält, die er auf seinen Reisen gesammelt hatte. In der „Practica Geometriae“ (1220) beschrieb Leonardo in ähnlicher Weise alles, was er in Geometrie und Trigonometrie in Erfahrung gebracht hatte. Er dürfte zudem auch ein selbständiger Forscher gewesen sein; denn seine Bücher enthalten viele Beispiele, die keine genauen Vorlagen in der arabischen Literatur zu haben scheinen. Er zitiert jedoch besonders Al-Khwarizmi, z.B. in der Diskussion der Gleichung $x^2 + 10x = 39$. Das Problem, das zur Folge der Fibonacci'schen Zahlen 0, 1, 1, 2, 3, 5, 8, 13, 21, ... führt, von denen jede die Summe der beiden vorhergehenden ist, scheint neu zu sein, ebenso wie sein bemerkenswert durchdachter Beweis, daß die Wurzel der Gleichung $x^3 + 2x^2 + 10x = 20$ nicht mit Hilfe von Euklidischen Irrationalitäten der Form $\sqrt{a + \sqrt{b}}$ ausgedrückt (und daher nicht unter alleiniger Verwendung von Zirkel und Lineal konstruiert)

werden können. Leonardo führte den Beweis, indem er jeden der fünfzehn Fälle bei Euklid nachprüfte und berechnete danach die positive Wurzel dieser Gleichung näherungsweise auf sechs Sexagesimalstellen.

- 3) Das Konstruieren mit Zirkel und Lineal hat eine große Geschichte. Das geometrische Mittel $h = \sqrt{p \cdot q}$ gewinnt man mit Hilfe des Satzes des Thaleskreises



Man hat hier ähnliche Dreiecke, also $\frac{p}{h} = \frac{h}{q}$.

Die Entdeckung der Konstruierbarkeit des regelmäßigen Siebzehnecks, die Gauß am 30. März 1796 in der Morgenfrühe gemacht hat, soll ihn im Alter von 19 Jahren dazu bestimmt haben, sich zum Studium der Mathematik zu entschließen. Und die Mathematikhistoriker sind sich einig, daß es Gauß war, der am Anfang des 19. Jahrhunderts den Blick auf einen „neuen Kontinent mathematischer Sachverhalte“ eröffnet hat. „*Mathematicorum princeps*“ lautet die Inschrift der Medaille, die der König von Hannover zum Andenken an Carl Friedrich Gauß in dessen Todesjahr 1855 schließen ließ. (vgl. Reidemeister: Raum und Zahl. Kapitel VIII) (vgl. auch Struik, S. 156 und S. 160 ff).

Man kann also sagen, daß die Quadratwurzeln und insbesondere die Irrationalzahl $\sqrt{2}$ eine zentrale Rolle in der von den Griechen begründeten Linie der Mathematik gespielt haben.

- 4) Eine vergleichbare Bedeutung kommt wohl nur der Zahl π und dem Problem der Quadratur des Kreises zu. Mit der Zahl π haben sich schon die alten Ägypter befaßt. Für sie war π das Verhältnis des Kreisumfangs zum Durchmesser; daß $\pi/4$ auch das Verhältnis der Kreisfläche zur Fläche des umbeschriebenen Quadrats ist, war ihnen anscheinend nicht bekannt (vgl. Struik S. 17).

Charakterisierungen der Zahl π , die keine Anleihen bei der Geometrie machen, werden uns im Verlauf der Vorlesung immer wieder einmal begegnen. Großes Aufsehen erregte um 1882 der Beweis der Transzendenz von π durch Lindemann (Cantor war der Gutachter für die Aufnahme dieser Arbeit in die Mathematischen Annalen). π spielt eine wichtige Rolle in der Theorie der trigonometrischen Reihen. Diese Theorie war in der historischen Entwicklung mehrfach Ausgangspunkt für die Schaffung tieflyingender neuer Begriffe. Auch für Cantors Weg als Wissenschaftler waren Probleme der Eindeutigkeit der Fourierreihe von entscheidender Bedeutung. Es führt also ein direkter Weg vom Interesse an π in die abstrakte Analysis des 20. Jahrhunderts.

- 5) Das Interesse an der Zahl e , der Basis des natürlichen Logarithmus, ist jüngerem Datum. Die Erfindung der Logarithmen geht auf praktische Interessen zurück. Zwar erschienen Logarithmen auf der Grundlage der Funktion $y = e^x$ fast gleichzeitig mit den Briggschen Logarithmen (um 1618 veröffentlichte E. WRIGHT einige natürliche Logarithmen, J. SPEIDEL 1619, aber danach wurden bis 1770 keine Tafeln dieser Logarithmen mehr veröffentlicht); ihre fundamentale Bedeutung wurde aber nicht eher erkannt, als bis die Infinitesimalrechnung besser verstanden worden war (siehe Struik, S. 104).

J. LIOUVILLE (1809–1882) war in Frankreich in den vierziger bis sechziger Jahren der führende Mathematiker. Er wies die Existenz von transzendenten Zahlen nach und bewies 1844, daß weder e noch e^2 Wurzel einer quadratischen Gleichung mit rationalen Koeffizienten sein kann. 1873 zeigte Hermite die Transzendenz von e . Nach Struiks Meinung war HERMITE (1822–1901) nach dem Tode von Cauchy im Jahre 1857 der führende Vertreter der Analysis in Frankreich. Sein Werk bewegte sich ebenso wie das von Liouville in der Tradition von Gauß und Jacobi; es zeigte auch eine gewisse Verwandtschaft mit Riemann und Weierstraß (Struik, S. 212).

- 6) Bemerkenswert ist der Widerstand, den hochangesehene und mächtige Mathematiker den Irrationalzahlen entgegengesetzt haben. L. KRONECKER (1823–1891) ragt heraus; Hilbert hat ihn „den klassischen Verbotsdiktator“ genannt. G. CANTOR hat schwer unter der Machtausübung von Kronecker gelitten (siehe Purkert, insbesondere S. 72, S. 76/77). Auf S. 51 heißt es

„Kronecker vertrat bezüglich der Grundlegung der Mathematik Ansichten, die in vielen Punkten die Konzeptionen des Intuitionismus vorwegnahmen. Er sah es als das höchste Ziel an, die gesamte Mathematik zu arithmetisieren, d.h. sie auf die sogenannte allgemeine Arithmetik zurückzuführen. Unter allgemeiner Arithmetisierung verstand er die Theorie der ganzen Zahlen und die Theorie der Polynomringe mit ganzzahligen Koeffizienten. ‚Die ganzen Zahlen hat der liebe Gott gemacht, alles andere ist Menschenwerk‘ war sein Wahlspruch. Er war davon überzeugt, daß ‚alle Ergebnisse der tiefstnigsten mathematischen Forschung schließlich in jenen einfachen Formen ganzer Zahlen ausdrückbar sein müssen‘. Insbesondere wollte er ‚die Hinzunahme der irrationalen sowie der kontinuierlichen Größen wieder abstreifen ...‘. Von den Schlußweisen, die zum Aufbau der Mathematik auf der so fixierten Grundlage erforderlich sind, verlangte Kronecker, daß sie finit und konstruktiv sind.

Alle Definitionen müssen entscheidbar sein, d.h. es muß ein Verfahren geben, welches es ermöglicht, in endlich vielen Schritten zu entscheiden, ob ein Objekt der gegebenen Definition entspricht oder nicht. Existenz bedeutet Konstruierbarkeit durch ein finites Verfahren. Es ist Kronecker gelungen, einige der von ihm entwickelten Theorien in diesem Sinne konstruktiv zu gestalten, z.B. die Idealtheorie in algebraischen Zahlkörpern. Im allgemei-

nen jedoch gelang das nicht, insbesondere hat er sich bei seinen berühmten Arbeiten über elliptische Funktionen völlig auf die traditionelle Analysis gestützt. Es ist festzuhalten, daß eine Beschränkung der zulässigen Hilfsmittel, wie Kronecker sie vornahm, auch positive Seiten hat und zu interessanten methodischen Fortschritten führen kann. ...“

B.3 Konstruktion der Exponentialfunktion

Aus der Sicht der arithmetisierten Analysis ist es nicht zulässig, die Exponentialfunktion und die trigonometrischen Funktionen aus geometrischen Gegebenheiten abzuleiten. Ihre Definition darf sich nur auf die Axiome stützen. Wir werden die Konstruktion in aller Ausführlichkeit durchführen. Das Ziel ist das

Theorem

a) Für jede komplexe Zahl z konvergiert die Folge

$$\left(1 + \frac{z}{n}\right)^n$$

gegen eine Zahl $e(z)$.

b) Für die Grenzfunktion $e(\cdot)$ gilt

$$e(z + w) = e(z) \cdot e(w) \quad \text{für alle } z, w.$$

c) Setzen wir

$$e_n(z) = 1 + z + \frac{1}{2!} z^2 + \dots + \frac{1}{n!} z^n$$

so gilt

$$e_n(z) \longrightarrow e(z) \quad \text{für alle } z.$$

Wir studieren zuerst reelle Argumente x , insbesondere $x \geq 0$. Die Beweise sind ein hervorragendes Übungsfeld für den Umgang mit Ungleichungen. Die Resultate werden wir dann benutzen, um die trigonometrischen Funktionen ohne Anleihen bei geometrischen Konstruktionen einzuführen.

Lemma 1 Für beliebiges $x \in \mathbb{R}$ ist die Zahlenfolge

$$\left(1 + \frac{x}{n}\right)^n$$

beschränkt und für $x \geq 0$ gilt

$$\left(1 + \frac{x}{n}\right)^n \leq \left(1 + \frac{x}{n+1}\right)^{n+1}.$$

Als Vorbereitung beweisen wir (durch vollständige Induktion nach n) zwei Aussagen, die uns auch sonst noch von Nutzen sein werden.

Bernoullische Ungleichung Für alle $y > -1$ und alle $n \in \mathbb{N}$ gilt

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

Man beweist leicht, daß sich für natürliche n die Einträge des Pascalschen Dreiecks ergeben. Was die verallgemeinerten Binomialkoeffizienten mit den „gebrochenen Potenzen“ $(1+x)^\alpha$ zu tun haben, werden wir später sehen.

Lemma 2 Sei $x \geq 0$ und $m, n \in \mathbb{N}$ mit $x \leq m$. Dann gilt

$$\left(1 + \frac{x}{n}\right)^{n+m} \geq \left(1 + \frac{x}{n+1}\right)^{n+1+m}.$$

Beweis

$$\begin{aligned} \left(\frac{1 + \frac{x}{n}}{1 + \frac{x}{n+1}}\right)^{n+1+m} &= \left(\frac{n(n+1) + (n+1)x}{n(n+1) + nx}\right)^{n+1+m} \\ &= \left(1 + \frac{x}{n(n+1+x)}\right)^{n+1+m} \\ &\geq 1 + \frac{n+1+m}{n+1+x} \cdot \frac{x}{n} \geq 1 + \frac{x}{n}. \end{aligned}$$

Beweis von Lemma 1

$$\begin{aligned} \left(1 + \frac{x}{n}\right)^n &= 1 + x + \binom{n}{2} \left(\frac{x}{n}\right)^2 + \binom{n}{3} \left(\frac{x}{n}\right)^3 \\ &\quad + \dots + \binom{n}{n} \left(\frac{x}{n}\right)^n \\ &= 1 + x + \frac{1}{2!} \left(1 - \frac{1}{n}\right) x^2 + \frac{1}{3!} \left(1 - \frac{1}{n}\right) \cdot \left(1 - \frac{2}{n}\right) x^3 \\ &\quad + \dots + \frac{1}{n!} \left(1 - \frac{1}{n}\right) \cdot \dots \cdot \left(1 - \frac{n-1}{n}\right) x^n \\ &\leq 1 + x + \frac{1}{2!} \left(1 - \frac{1}{n+1}\right) x^2 + \frac{1}{3!} \left(1 - \frac{1}{n+1}\right) \cdot \left(1 - \frac{2}{n+1}\right) x^3 \\ &\quad + \dots + \frac{1}{n!} \left(1 - \frac{1}{n+1}\right) \cdot \dots \cdot \left(1 - \frac{n-1}{n+1}\right) x^n \\ &\leq \left(1 + \frac{x}{n+1}\right)^{n+1} \end{aligned}$$

Spezialfall $x = 1$

Die Folge $\left(1 + \frac{1}{n}\right)^n$ ist monoton steigend,

die Folge $\left(1 + \frac{1}{n}\right)^{n+1}$ ist monoton fallend,

die Folgen haben den gemeinsamen Grenzwert

$$\lim \uparrow \left(1 + \frac{1}{n}\right)^n = e = \lim \downarrow \left(1 + \frac{1}{n}\right)^{n+1}.$$

Ebenso ergibt sich der

Satz 1 Für jedes $x \geq 0$ und jedes $m \geq x$

$$\sup \left\{ \left(1 + \frac{x}{n}\right)^n : n \in \mathbb{N} \right\} = \inf \left\{ \left(1 + \frac{x}{n}\right)^{n+m} : n \in \mathbb{N} \right\} .$$

Den gemeinsamen Grenzwert bezeichnen wir mit $e(x)$.

Satz 2 Sei $e(x)$ für $x \geq 0$ wie in Satz 1 und $e(-x) = \frac{1}{e(x)}$.

Es gilt dann für alle $x, y \in \mathbb{R}$

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

Beweis

1) Es genügt die Funktionalgleichung für $x \geq 0, y \geq 0$ zu beweisen.

Sei nämlich $z = x + y$ mit irgendwelchen x, y, z

$$\begin{aligned} -z &= (-x) + (-y), \\ x &= z + (-y), \quad -x = (-z) + y, \\ y &= z + (-x), \quad -y = (-z) + x . \end{aligned}$$

o.B.d.A. nehmen wir an, daß z dem Betrag nach am größten ist, also $|z| = |x| + |y|$.
Wenn wir

$$e(|x|) \cdot e(|y|) = e(|z|)$$

bewiesen haben, dann haben wir auch

$$e(x) \cdot e(y) = e(z)$$

bewiesen.

2) Sei $x \geq 0, y \geq 0, z = x + y$. Es gilt

$$\left(1 + \frac{z}{n}\right)^n \leq \left(1 + \frac{x}{n}\right)^n \left(1 + \frac{y}{n}\right)^n$$

also $e(z) \leq e(x) \cdot e(y)$.

3) Setze für $x \geq 0$

$$e_n(x) = 1 + x + \frac{1}{2!} x^2 + \dots + \frac{1}{n!} x^n .$$

Im Beweis von Lemma 1 haben wir gesehen

$$\left(1 + \frac{x}{n}\right)^n \leq e_n(x) \leq e(x) .$$

Also gilt $e(x) = \lim \uparrow e_n(x)$.

4) Für $x, y, \geq 0$ haben wir

$$e_n(x) \cdot e_n(y) \leq e_{2n}(x+y);$$

denn Ausmultiplizieren liefert nach der binomialen Formel

$$\begin{aligned} \left(\sum_{k=0}^n \frac{1}{k!} x^k \right) \cdot \left(\sum_{\ell=0}^n \frac{1}{\ell!} x^\ell \right) &\leq \sum_{m=0}^{2n} \frac{1}{m!} \left(\sum_{k+\ell=m} \frac{m!}{k!\ell!} x^k y^\ell \right) \\ &= \sum_{m=0}^{2n} \frac{1}{m!} (x+y)^m. \end{aligned}$$

5) $e(x) \cdot e(y) \geq e_{2n}(x+y)$ für alle n , also mit 2)

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

Didaktischer Hinweis Die elementare Analysis bietet eine Unzahl von Gelegenheiten, wo der Student seine Findigkeit zeigen kann. Identitäten und Abschätzungen, die zunächst mit dem vorgegebenen Ziel nichts zu tun zu haben scheinen, müssen vermutet und (häufig durch vollständige Induktion) bewiesen werden. Wer wenig Erfahrung hat, muß viel ausprobieren; es ist fast wie bei Schachproblemen. Wer mehr weiß, tut sich leichter mit den Vermutungen. Der Student staunt oft, wie wohl der Dozent auf seine Vermutungen gekommen ist, die er da beweist.

Wie könnte er etwa auf die Vermutung gekommen sein, daß

$$\left(1 - \frac{x}{n}\right)^n \geq \left(1 - \frac{x}{n+1}\right)^{n+1} \quad \text{für } 0 < x < n?$$

Der Dozent hat eben das Bild der Funktion

$$\frac{1}{y} \ln(1+y) \quad (\text{für } y > -1)$$

vor Augen. Er weiß, daß sie fallend ist; wegen $\left(-\frac{x}{n}\right) > \left(-\frac{x}{n+1}\right)$ gilt

$$\left(-\frac{n}{x}\right) \ln\left(1 - \frac{x}{n}\right) \leq \left(-\frac{n+1}{x}\right) \ln\left(1 - \frac{x}{n+1}\right).$$

Wer nichts vom Logarithmus weiß, hat es schwerer.

Den Beweis von

$$\left(1 - \frac{x}{n}\right)^n \searrow e(-x) = \frac{1}{e(x)} \quad \text{für } x > 0$$

überlassen wir dem Leser.

Leichter zu vermuten und zu beweisen sind die Abschätzungen

$$\begin{aligned} e_{2n}(-x) &\leq e_{2n+2}(-x) \leq e(-x), & \text{wenn } 0 < x < 2n \\ e_{2n-1}(-x) &\geq e_{2n+1}(-x) \geq e(-x), & \text{wenn } 0 < x < 2n+1 \\ |e_n(x) - e(x)| &\leq |e_n(|x|) - e(|x|)| & \text{für alle } x. \end{aligned}$$

Mühevoll für einen, der keine Theorie kennt, sind die folgenden Abschätzungen zu beweisen:
Für $x > 0$ gilt

$$\begin{aligned} |e(-x) - e_n(-x)| &\leq \frac{x^{n+1}}{(n+1)!}, \\ |e(x) - e_n(x)| &\leq e^x \cdot \frac{x^{n+1}}{(n+1)!}. \end{aligned}$$

Eine wirklich wichtige Abschätzung ist die folgende

Satz $e(x) \geq 1 + x$ für alle $x \in \mathbb{R}$.
(Beweis mit Hilfe der Bernoullischen Ungleichung !)

Korollar Die Funktion $e(\cdot)$ ist überall stetig und

$$\frac{1}{h} (e(x+h) - e(x)) \longrightarrow e(x) \quad \text{für } h \rightarrow 0.$$

Beweis $e(x+h) - e(x) = e(x) [e(h) - 1]$.

Es genügt also zu beweisen

$$\frac{1}{h} (e(h) - 1) \longrightarrow 1.$$

Wegen $e^h \geq 1 + h$, $e^{-h} \leq \frac{1}{1+h}$, $e^h \leq \frac{1}{1-h}$, $e^{-h} \geq 1 - h$ gilt

$$e^h - 1 \longrightarrow 0 \quad \text{für } h \rightarrow 0$$

$$\begin{aligned} e^{2h} - 1 &= e^h(e^h - e^{-h}) \geq e^h \left(1 + h - \frac{1}{1+h}\right) \\ &= e^h \frac{1}{1+h} 2h \left(1 + \frac{h}{2}\right) \\ e^{2h} - 1 &\leq e^h \left(\frac{1}{1-h} - (1-h)\right) = e^h \frac{1}{1-h} 2h \left(1 - \frac{h}{2}\right). \end{aligned}$$

Daraus folgt die Behauptung.

Bemerkung Für jede konvergente Folge $x_n \rightarrow x^*$ gilt

$$\left(1 + \frac{x_n}{n}\right) \longrightarrow e(x^*).$$

Für jedes $\varepsilon > 0$ gilt nämlich schließlich

$$x^* - \varepsilon < x_n < x^* + \varepsilon$$

und daher

$$\left(1 + \frac{x^* - \varepsilon}{n}\right)^n \leq \left(1 + \frac{x_n}{n}\right)^n < \left(1 + \frac{x^* + \varepsilon}{n}\right)^n.$$

Der begriffliche Standpunkt (Vorgriff auf weitere Begriffsbildungen)

Ist nun die oben konstruierte Funktion $e(\cdot)$ wirklich dieselbe Funktion, die andere Autoren auf andere Weise konstruieren und dann die Exponentialfunktion nennen? Manche Autoren konstruieren zuerst den natürlichen Logarithmus als Stammfunktion der Funktion $\frac{1}{x}$ und nennen dessen Umkehrfunktion die Exponentialfunktion. Andere führen die Exponentialfunktion ein als die (eindeutig bestimmte) Lösung der Differentialgleichung

$$f'(x) = f(x) \text{ mit } f(0) = 1.$$

Die einen brauchen also etwas Integrationstheorie, die anderen ein wenig Theorie der Differentialgleichungen. In einer solchen Situation suchen die modernen Mathematiker nach möglichst einfachen intrinsischen Charakterisierungen des Gesuchten, d.h. Befreiung von den Eigenarten der speziellen Konstruktion.

Cauchy hat in seinen „Cours d'analyse“ von 1821 den folgenden Charakterisierungssatz bewiesen:

Satz Die stetigen Lösungen der Funktionalgleichung

$$f(x+y) = f(x) \cdot f(y) \quad \text{für alle } x, y \in \mathbb{R}$$

sind die Funktionen $f_\alpha(x) = e^{\alpha x}$ (α reell). Andere stetige Lösungen gibt es nicht.

Im 20. Jahrhundert konnte dann im Zuge des Aufkommens der Mengenlehre bewiesen werden, daß es auch keine weiteren borelmeßbaren Lösungen der „Cauchy-schen Funktionalgleichung“ gibt, daß jedoch die Annahme, daß es nichtmeßbare Lösungen gibt, nicht zum Widerspruch geführt werden kann (wenn nicht überhaupt die klassische Analysis in sich widerspruchsvoll ist).

Fazit Jeder, der mit irgendwelchen Mitteln eine stetige Lösung der Cauchyschen Funktionalgleichung gewonnen hat, hat eine der Funktionen $f_\alpha(\cdot)$ gewonnen.

Cauchys Anstrengungen wurden von prominenten Zeitgenossen kritisiert. M. OHM schrieb 1829 in seiner Rezension:

„Sieht nicht der Anfänger sogleich, daß der Verfasser das Gesuchte bereits im Sacke mit sich herumträgt und nur gegen seinen Begleiter so tut, als suche er es? ... allein wenn die ganze Untersuchung sich bloß auf die Auffindung der Funktionen A^x beschränkt, so erscheint uns das Ganze, wenn man dem Anfänger von den Dingen, die er kennen lernen soll, nur die Schatten zeigen wollte; und dem Geübten sind dergleichen Wendungen in seinen Fragen nichts Neues.“ (zitiert nach H.N. JAHNKE: „Motive und Probleme der Arithmetisierung der Mathematik in der ersten Hälfte des 19. Jahrhunderts“, Archive for History of Exact Sciences, Vol. 37 Nr. 2, 1987, pp. 101–182).

Ohm hat andererseits in seinem Lehrbuch von 1822, als einer der ersten Autoren in Deutschland eine dem modernen Verständnis entsprechende Definition des Stetigkeitsbegriffs gehabt. Jahnke kommentiert, daß es für die Tatsache, ob ein Autor wirklich über einen bestimmten Begriff verfügt, nicht so sehr seine verbale Definition entscheidend ist, sondern vor allem die Art und Weise, wie er damit arbeitet. Während „stetig“ bei Ohm eine unter vielen möglichen Eigenschaften seiner (durch Potenzreihen konkret vorgegebenen) Funktionen darstellte, hatte der Stetigkeitsbegriff in Cauchys Werk die fundamentale Aufgabe, den bis dahin algebraischen Funktionenbegriff *begrifflich zu verallgemeinern*. (Jahnke, S. 155).

Der natürliche Logarithmus

Die Logarithmen wurden um 1600 herum erfunden, nicht lange nach der allgemeinen Einführung der indisch-arabischen Zahlenschreibweise und nachdem sich die Dezimalbrüche durchgesetzt hatten. SIMON STEVIN hatte sich darum 1585 verdient gemacht. Mehrere Mathematiker des sechzehnten Jahrhunderts hatten mit der Möglichkeit gespielt, arithmetische und geometrische Reihen zu verknüpfen, hauptsächlich zu dem Zweck, das Arbeiten mit den komplizierten trigonometrischen Tafeln zu erleichtern. (Trigonometrie gab es schon länger, siehe Struik S. 95 ff).

JOHN NEPER (oder Napier), ein schottischer Gutsbesitzer, hat 1614 einen wesentlichen Beitrag geleistet. Sein Bewunderer Briggs veröffentlichte 1624 eine erste Logarithmentafel, die 1627 von EZECHIEL DE DECKER, einem holländischen Feldmesser vervollständigt wurde (Struik S. 103/104, Walter S. 57). Die Bedeutung der natürlichen Logarithmen wurde erst mit dem Aufkommen der Infinitesimalrechnung erkannt (siehe auch Walter S. 132). GREGORIUS A SANTO VICENTIO hat 1647 entdeckt, daß man Logarithmen als Flächen zwischen einer Hyperbel und ihrer Asymptote ansehen kann. MERCATOR veröffentlichte 1668 die logarithmische Reihe

$$\ln(1+a) = \int_0^a \frac{1}{1+x} dx = a - \frac{a^2}{2} + \frac{a^3}{3} - \frac{a^4}{4} + \dots$$

Wir haben oben die Exponentialfunktion direkt vom Prinzip der oberen Grenze (einer monotonen Zahlenfolge) her konstruiert. In ganz ähnlicher Weise kann man auch die Logarithmusfunktion konstruieren. (Wenn man den Logarithmus als Umkehrfunktion der Exponentialfunktion einführen will, braucht man den Zwischenwertsatz, was als weniger direkt gilt.) Die Idee der direkten Konstruktion beruht auf der (für viel allgemeinere Zusammenhänge nützlichen) Einsicht

$$\lim_{h \rightarrow 0} \frac{1}{h} (x^h - 1) = \ln x .$$

In Barner-Flohr, Analysis I, 1974 wird das Programm auf S. 130 ff durchgeführt:

$$h_k(x) := 2^k \left(\sqrt[k]{x} - 1 \right)$$

ist monoton fallend für alle $x > 0$

$$g_k(x) := 2^k \left(1 - \frac{1}{\sqrt[k]{x}} \right)$$

ist monoton steigend. $h_k(x) - g_k(x) \geq 0$.

Der gemeinsame Grenzwert ist eine monotone Funktion $f(x)$ mit

$$(i) \quad f(xy) = f(x) + f(y) \quad \text{für alle } x, y > 0$$

$$(ii) \quad f(x) \leq x - 1 \quad \text{für alle } x > 0$$

Man zeigt, daß es nur eine Funktion mit diesen Eigenschaften gibt; man nennt sie den natürlichen Logarithmus. Von diesem Eindeutigkeitssatz her findet man die Brücke zu allen anderen direkten Konstruktionen des natürlichen Logarithmus. Bei jeder Konstruktion kann man etwas lernen. Eine Frage, die wir nicht behandeln wollen, ist die nach einer effizienten Berechnung der Logarithmen. Das ist Sache der numerischen Mathematik; die Mittel sind da das Entscheidende.

B.4 Konstruktion der trigonometrischen Funktionen

Definition Für $y \in \mathbb{R}$ und $n \in \mathbb{N}$ setze

$$\begin{aligned} c_{2n}(y) &= 1 - \frac{1}{2!} y^2 + \frac{1}{4!} y^4 - \dots + (-1)^n \cdot \frac{1}{2n!} y^{2n} \\ s_{2n-1}(y) &= y - \frac{1}{3!} y^3 + \frac{1}{5!} y^5 - \dots + (-1)^{n-1} \cdot \frac{1}{(2n-1)!} y^{2n-1}. \end{aligned}$$

Weiterhin

$$\begin{aligned} c(y) &= \lim_{n \rightarrow \infty} c_{2n}(y) = \sum_{k=0}^{\infty} (-1)^k \cdot \frac{1}{2k!} y^{2k} \\ s(y) &= \lim_{n \rightarrow \infty} s_{2n-1}(y) = \sum_{k=0}^{\infty} (-1)^k \cdot \frac{1}{(2k+1)!} y^{2k+1} \\ e(iy) &= c(y) + is(y) \end{aligned}$$

Theorem

- (i) $|e(iy)| = 1$ für alle y
- (ii) $e(iy_1) \cdot e(iy_2) = e(i(y_1 + y_2))$ für alle y_1, y_2 .
- (iii) $e(i\pi) = -1$

Für den Beweis benötigen wir einige Vorarbeiten. Insbesondere müssen wir die Ludolfsche Zahl π mit rein analytischen Mitteln einführen.

Lemma 1 Die Folgen $c_{2n}(y)$ und $s_{2n}(y)$ konvergieren für jedes y .

Beweis

- 1) Setze für $n = 1, 2, \dots$

$$\begin{aligned} e_{2n}(iy) &= c_{2n}(iy) + is_{2n-1}(iy) \\ e_{2n-1}(iy) &= c_{2n-2}(iy) + is_{2n-1}(iy). \end{aligned}$$

Es gilt für $N = 1, 2, \dots$

$$e_N(iy) = 1 + (iy) + \frac{1}{2!} (iy)^2 + \frac{1}{3!} (iy)^3 + \dots + \frac{1}{N!} (iy)^N.$$

- 2) Zu jedem $y \in \mathbb{R}$, $\varepsilon > 0$ existiert ein N , so daß für $N \leq n < m$

$$|e_m(iy) - e_n(iy)| < \varepsilon.$$

Es gilt nämlich

$$|e_m(iy) - e_n(iy)| \leq \sum_{k=N+1}^{\infty} \frac{1}{k!} |y|^k = e(|y|) - e_N(|y|)$$

mit den oben für reelle Argumente definierten Funktion $e(\cdot)$.

- 3) Aus der Konvergenz der Folge $e_N(iy)$ folgt die Konvergenz des Real- und des Imaginärteils. ■

Satz Für jedes $z \in \mathbb{C}$ konvergiert die Folge $e_n(z)$. Die Grenzfunktion $e(z)$ ist stetig und es gilt

$$\left| \frac{e(z) - e(w)}{z - w} - e(w) \right| \longrightarrow 0 \quad \text{für } |z - w| \rightarrow 0$$

gleichmäßig für $|z|, |w|$ beschränkt.

Beweis

- 1) Für $|z| \leq C$ und $N \leq n$ gilt

$$|e_n(z) - e_n(w)| \leq \sum_{k=N+1}^{\infty} \frac{1}{k!} C^k = e(C) - e_n(C).$$

Die $e_n(z)$ bilden also eine Cauchyfolge, die gleichmäßig in $|z| \leq C$ konvergiert.

- 2) $z^k - w^k = (z - w) (z^{k-1} + z^{k-2}w + \dots + zw^{k-2} + w^{k-1})$

$$\frac{z^k - w^k}{z - w} - kw^{k-1} = (z^{k-1} - w^{k-1}) + (z^{k-2} - w^{k-2})w + \dots + (z - w)w^{k-2}$$

$$\left| \frac{1}{z - w} \left[\frac{z^k - w^k}{z - w} - kw^{k-1} \right] \right| \leq \frac{(k-1)k}{2} \cdot C^{k-2} \quad \text{für } k = 2, 3, \dots$$

$$\left| \frac{e_n(z) - e_n(w)}{z - w} - e_{n-1}(w) \right| \leq |z - w| \cdot \frac{1}{2} \sum_{k=0}^{n-2} \frac{1}{k!} \leq |z - w| e(C).$$

Für festes $z \neq w$ wählen wir n so groß, daß

$$\left| \frac{e_n(z) - e_n(w)}{z - w} - \frac{e(z) - e(w)}{z - w} \right| < \frac{\varepsilon}{4} |z - w|$$

und $|e_{n-1}(w) - e(w)| < \frac{\varepsilon}{4} |z - w|$.

Dann haben wir

$$\left| \frac{e_n(z) - e_n(w)}{z - w} - e(w) \right| \leq |z - w| \frac{1}{2} [e(C) + \varepsilon].$$

■

Hinweis Das Argument wird uns wieder begegnen, wenn wir beweisen, daß man eine konvergente Potenzreihe gliedweise differenzieren darf.

Satz Für jedes $z \in \mathbb{C}$ gilt

$$\lim_{n \rightarrow \infty} \left(1 + \frac{z}{m}\right)^m = e(z) = \lim_{n \rightarrow \infty} e_n(z).$$

Die Konvergenz ist gleichmäßig für $|z| \leq C$ (C beliebig $< \infty$).

Beweis

1) Für $n < m$ gilt

$$\begin{aligned} \left(1 + \frac{z}{m}\right)^m &= 1 + z + \frac{1}{2!} \left(1 - \frac{1}{m}\right) z^2 + \dots \\ &\quad + \frac{1}{n!} \left(1 - \frac{1}{m}\right) \cdot \dots \cdot \left(1 - \frac{n-1}{m}\right) z^n + R^{(n,m)}(z) \end{aligned}$$

mit

$$\begin{aligned} R^{(n,m)}(z) &= \frac{1}{(n+1)!} \left(1 - \frac{1}{m}\right) \cdot \dots \cdot \left(1 - \frac{n}{m}\right) z^{n+1} + \dots \\ &\quad + \frac{1}{m!} \left(1 - \frac{n}{m}\right) \cdot \dots \cdot \left(1 - \frac{m-1}{m}\right) z^m \\ |R^{(n,m)}(z)| &\leq \sum_{k=n+1}^{\infty} \frac{1}{k!} C^k \quad \text{für } |z| \leq C \end{aligned}$$

$$\begin{aligned} 2) \quad \left(1 + \frac{z}{m}\right)^m - R^{(n,m)}(z) - e_n(z) &= \sum_{k=2}^n \alpha_k^{(m)} \cdot \frac{1}{k!} z^k \\ &=: P^{(n,m)}(z) \end{aligned}$$

$$\text{mit } \alpha_k^{(m)} = -1 + \left(1 - \frac{1}{m}\right) \cdot \left(1 - \frac{2}{m}\right) \cdot \dots \cdot \left(1 - \frac{k-1}{m}\right)$$

$P^{(n,m)}(z)$ ist ein Polynom vom Grad n .

Die Koeffizienten streben nach 0 für $m \rightarrow \infty$.

3) Wählen wir n so groß, daß für alle $|z| \leq C$

$$\begin{aligned} \text{(i)} \quad |e_n(z) - e(z)| &< \frac{\varepsilon}{3} \\ \text{(ii)} \quad \sum_{k=n+1}^{\infty} \frac{1}{k!} C^k &< \frac{\varepsilon}{3}. \end{aligned}$$

Wählen wir sodann m so groß, daß

$$\text{(iii)} \quad |P^{(n,m)}(z)| < \frac{\varepsilon}{3} \quad \text{für alle } |z| \leq C.$$

Wir haben dann

$$\left|\left(1 + \frac{z}{m}\right)^m - e(z)\right| < \varepsilon.$$

■

Korollar Für $z = x + iy$ gilt $|e(z)| = e(x)$.

Insbesondere gilt für alle $y \in \mathbb{R}$

$$c^2(y) + s^2(y) = 1.$$

Beweis

$$\begin{aligned} \left|1 + \frac{z}{m}\right|^2 &= \left|1 + \frac{x}{m} + i \frac{y}{m}\right|^2 = \left(1 + \frac{x}{m}\right)^2 + \frac{y^2}{m^2} \\ &= 1 + \frac{2x}{m} + \frac{1}{m^2}(x^2 + y^2) \\ \left|1 + \frac{z}{m}\right|^m &= \left(1 + \frac{2x}{m} + \frac{1}{m^2}(x^2 + y^2)\right)^{m/2} \longrightarrow e(x). \end{aligned}$$

■

Satz

$$e(z + w) = e(z) \cdot e(w) \quad \text{für alle } z, w \in \mathbb{C}.$$

Beweis Wir brauchen hier, daß $\left(1 + \frac{z}{m}\right)^m$ in einer Umgebung des Nullpunkts gleichmäßig konvergiert

$$\begin{aligned} \frac{\left(1 + \frac{z}{m}\right)\left(1 + \frac{w}{m}\right)}{1 + \frac{z+w}{m}} &= 1 + \frac{1}{m} z_m \quad \text{mit } z_m = \frac{zw}{m + (z + w)} \longrightarrow 0 \\ \left(1 + \frac{1}{m} z_m\right)^m &\longrightarrow e(0) = 1 \end{aligned}$$

■

Beweis des Theorems Die Abbildung

$$\mathbb{R} \ni y \longmapsto e(iy)$$

bildet die reelle Achse auf den Einheitskreis ab. Wenn man die reelle Achse mit gleichförmiger Geschwindigkeit durchläuft, dann durchläuft der Bildpunkt den Einheitskreis mit gleichförmiger Winkelgeschwindigkeit; denn bei der Multiplikation komplexer Zahlen vom Betrag 1 addieren sich die Winkel.

Bezeichne $\frac{\pi}{2}$ die kleinste positive Nullstelle der Funktion $c(y) = \Re e(e(iy))$. $e\left(i \frac{\pi}{2}\right) = i$. Wir zeigen mit Mitteln, die später systematisch entwickelt werden (Variablensubstitution bei Integralen), daß das so definierte π das Volumen des Einheitskreises ist. (Die Idee, $\frac{\pi}{2}$ als die erste Nullstelle durch die Potenzreihe definierte Cosinusfunktion zu definieren, wurde 1933 als undeutsche Unnatürlichkeit gebrandmarkt (siehe SCHAPPACHER u. KNESER in der Festschrift

zum 100jährigen Bestehen der Deutschen Mathematikervereinigung, 1990, S. 58).

$$e\left(i\left(y + \frac{\pi}{2}\right)\right) = ie(y), \quad e(-iy) = \overline{e(iy)} = \frac{1}{e(iy)}$$

$$c\left(y + \frac{\pi}{2}\right) = -s(y), \quad s\left(y + \frac{\pi}{2}\right) = c(y)$$

$$c(-y) = c(y), \quad s(-y) = -s(y)$$

$$u = s(y) \text{ ist isoton für } y \in \left[-\frac{\pi}{2}, +\frac{\pi}{2}\right], \quad s\left(-\frac{\pi}{2}\right) = -1, \quad s\left(\frac{\pi}{2}\right) = 1$$

$$c(y) = \sqrt{1 - s^2(y)} \quad \text{für } y \in \left[-\frac{\pi}{2}, +\frac{\pi}{2}\right].$$

$c^2(\cdot)$ und $s^2(\cdot)$ sind π -periodische Funktionen

$$\int_{-\pi/2}^{+\pi/2} c^2(y) dy = \frac{1}{2} \int_{-\pi/2}^{+\pi/2} (c^2(y) + s^2(y)) dy = \frac{\pi}{2}.$$

Auf der anderen Seite liefert Variablensubstitution $u = \sin y$

$$\int_{-\pi/2}^{+\pi/2} c^2(y) dy = \int_{-1}^{+1} \sqrt{1 - u^2} du = \text{Volumen des Halbkreises}.$$

Damit ist die Verbindung zwischen der Exponentialfunktion im Komplexen und der Ludolf'schen Zahl π hergestellt. ■

Hinweis Die oben konstruierte Funktion

$$e(z) = \sum_{k=0}^{\infty} \frac{1}{k!} z^k = \lim_{n \rightarrow \infty} \left(1 + \frac{z}{n}\right)^n \quad \text{für } z \in \mathbb{C}$$

bezeichnet man gewöhnlich mit

$$e^z \quad \text{oder} \quad \exp(z).$$

Die Funktionen

$$\cos z = \frac{1}{2} (e^{iz} + e^{-iz}), \quad \sin z = \frac{1}{2i} (e^{iz} - e^{-iz})$$

heißen die Cosinusfunktion bzw. Sinusfunktion im Komplexen. Wir haben damit auch für komplexe Argumente z

$$e^z = \cos z + i \cdot \sin z.$$

Es gilt auch für komplexe z

$$\begin{aligned}\cos z &= 1 - \frac{1}{2!} z^2 + \frac{1}{4!} z^4 - \frac{1}{6!} z^6 + \dots \\ \sin z &= z - \frac{1}{3!} z^3 + \frac{1}{5!} z^5 - \frac{1}{7!} z^7 + \dots\end{aligned}$$

Die Additionstheoreme gelten auch für komplexe Argumente

$$\begin{aligned}\cos z \cdot \cos w - \sin z \cdot \sin w &= \frac{1}{4} (e^{iz} + e^{-iz}) (e^{iw} + e^{-iw}) - \frac{1}{4} (e^{iz} - e^{-iz}) (e^{iw} - e^{-iw}) \\ &= \frac{1}{2} e^{iz} \cdot e^{iw} + \frac{1}{2} e^{-iz} \cdot e^{-iw} \\ &= \cos(z + w) .\end{aligned}$$

Ebenso

$$\sin(z + w) = \sin z \cdot \cos w + \cos z \cdot \sin w .$$

Die Funktionen $\cos z$ und $\sin z$ sind komplex differenzierbar mit dem Ergebnis

$$\frac{d}{dz} \cos z = -\sin z , \quad \frac{d}{dz} \sin z = \cos z .$$

Schließlich gilt auch für komplexe z

$$\cos^2 z + \sin^2 z = 1 .$$

Aufgabe Zeige

$$\sin z \neq 0 \quad \text{für} \quad \frac{z}{\pi} \notin \mathbb{Z} .$$

B.5 Konvergente Potenzreihen

Ein formaler Ausdruck der Art

$$S(z) : a_0 + a_1z + a_2z^2 + \dots$$

heißt eine **formale Potenzreihe**. Die a_n sind komplexe Zahlen und heißen die Koeffizienten. Die Unbestimmte z wird auch häufig mit anderen Buchstaben bezeichnet, z.B. mit x .

Das formale Rechnen mit Potenzreihen wird uns in der nächsten Stunde beschäftigen. Solange man nur addiert, subtrahiert und multipliziert, ergibt sich kaum Neues gegenüber dem formalen Rechnen mit Polynomen. Interessant wird es, wenn man auch zu dividieren versucht (mit Vorsicht!) und wenn man eine formale Potenzreihe in eine andere einsetzt (mit Vorsicht).

Hier interessiert uns schon einmal die „formale Ableitung“

$$S'(z) : a_1 + 2a_2z + 3a_3z^2 + \dots$$

$$S^{(2)}(z) : 2a_2 + 3 \cdot 2a_3z + 4 \cdot 3 \cdot a_4z^2 + \dots$$

...

$$S^{(n)}(z) : n!a_n + [n+1]_n a_{n+1}z + [n+2]_n a_{n+2}z^2 + \dots$$

...

(Wir benützen hier die Notation aus der Kombinatorik

$$[n]_k := n(n-1) \cdot \dots \cdot (n-k+1) \quad \text{„}n \text{ untere Faktorielle } k\text{“}$$

Für jede formale Potenzreihe $S(z)$ kann man den „Wert im Nullpunkt“ $S(0) = a_0$ definieren und die „Werte der Ableitungen im Nullpunkt“, $S^{(n)}(0) := n!a_n$. So etwas wie den „Wert in einem Punkt $\tilde{z} \neq 0$ “ gibt es jedoch nur für Potenzreihen besonderer Art. Diese Potenzreihen interessieren uns hier.

Definition Man sagt von einer Potenzreihe

$$S(z) : a_0 + a_1z + a_2z^2 + \dots,$$

daß sie im Punkt $\tilde{z}$ einen Wert hat, wenn die Folge $(a_0 + a_1\tilde{z} + \dots + a_N\tilde{z}^N)_N$ konvergiert.

Beispiel Die Potenzreihe

$$s(z) : z - \frac{z^2}{2} + \frac{z^3}{3} - \frac{z^4}{4} + \dots$$

hat einen Wert im Punkt $\tilde{z} = 1$; sie hat aber keinen Wert im Punkt $\tilde{z} = -1$. (Die harmonische Reihe divergiert!)

Satz 1 Wenn die formale Potenzreihe $S(z)$ in $\tilde{z} \neq 0$ einen Wert $f(\tilde{z})$ hat, dann haben die formalen Ableitungen $S^{(n)}(z)$ einen Wert $f^{(n)}(\tilde{z})$ für alle z mit $|z| < |\tilde{z}|$ ($n = 0, 1, 2, \dots$).

(Beweis unten!)

Sprechweisen Wenn eine formale Potenzreihe $S(z)$ in keinem Punkt $\tilde{z} \neq 0$ einen Wert hat, dann sagt man, sie hat den Konvergenzradius 0.

Man sagt, $S(z)$ hat den Konvergenzradius R , wenn $S(z)$ einen Wert hat für alle z mit $|z| < R$.

Beispiele

a) Die folgenden Potenzreihen haben den Konvergenzradius $R = \infty$

$$\begin{aligned} 1 + z + \frac{1}{2!} z^2 + \frac{1}{3!} z^3 + \dots \\ 1 - \frac{1}{2!} z^2 + \frac{1}{4!} z^4 - \frac{1}{6!} z^6 + \dots \\ z - \frac{1}{3!} z^3 + \frac{1}{5!} z^5 - \frac{1}{7!} z^7 + \dots \end{aligned}$$

b) Die folgenden Potenzreihen haben den Konvergenzradius $R = 1$

$$\begin{aligned} z - \frac{1}{2} z^2 + \frac{1}{3} z^3 - \frac{1}{4} z^4 + \dots \\ z - \frac{1}{3} z^3 + \frac{1}{5} z^5 - \frac{1}{7} z^7 + \dots \end{aligned}$$

(Wir werden sehen, daß der Wert dieser letzten Reihe für reelle z mit $|z| < 1$ gleich $\arctan z$ ist.)

Satz 2 Wenn die formale Potenzreihe $S(z)$ in $\tilde{z} \neq 0$ einen Wert hat, dann konvergiert die Folge der Partialsummen gleichmäßig in jedem Kreis $B_r(0) = \{z : |z| \leq r < |\tilde{z}|\}$

$$f(z) = \lim_{N \rightarrow \infty} (a_0 + a_1 z + \dots + a_N z^N) \quad \text{gleichmäßig in } B_r(0).$$

Hinweis Der Satz impliziert zusammen mit Satz 1, daß die Folge der Partialsummen der formalen Ableitungen ebenfalls in jedem $B_r(0)$ gleichmäßig konvergiert. Wir werden sehen, daß der Grenzwert gleich der Ableitung (in dem von der Schule her bekannten Sinn) ist.

Beweis der Sätze 1 und 2

1) Aus der Konvergenz der Zahlenfolge

$$\left(\sum_{n=0}^N a_n \tilde{z}^n \right)_N$$

folgt, daß es eine Zahl C gibt mit

$$|a_n \tilde{z}^n| \leq C \quad \text{für alle } n.$$

In der Tat gilt sogar $|a_n \tilde{z}^n| < \varepsilon$ für alle genügend großen n .

2) Wähle $q < 1$. Für $|z| \leq q|\tilde{z}|$ haben wir

$$|a_n \tilde{z}^n| \leq C q^n.$$

Für alle $m, n \geq N$ gilt somit

$$\left| \sum_{k=0}^m a_k z^k - \sum_{k=0}^n a_k z^k \right| \leq \sum_{k=N}^{\infty} |a_k| |z^k| \leq C q^N (1 + q + q^2 + \dots)$$

Dies ist $< \varepsilon$, wenn N so groß ist, daß $C q^N \frac{1}{1-q} < \varepsilon$. Die Folge der Partialsummen konvergiert also gleichmäßig in $B_r(0)$.

3) Wir wissen zwar nicht, ob die Folge $|n a_n \tilde{z}^{n-1}|$ beschränkt ist. Jedoch gilt für jedes $\tilde{\tilde{z}}$ mit $|\tilde{\tilde{z}}| < |\tilde{z}|$

$$\begin{aligned} |n a_n \tilde{\tilde{z}}^{n-1}| & \text{ beschränkt} \\ |n(n-1) a_n \tilde{\tilde{z}}^{n-2}| & \text{ beschränkt} \\ \dots \end{aligned}$$

Mit dem Schluß in 2) ergibt sich die gleichmäßige Konvergenz der Folge der Partialsummen für $S'(z), S''(z), \dots$ in jedem $B_r(0)$, $r < |\tilde{z}|$.

Die Begriffe und Aussagen übertragen sich sofort auf Potenzreihen mit der Unbestimmten $z - z_0$

$$S(z - z_0) : b_0 + b_1(z - z_0) + b_2(z - z_0)^2 + \dots$$

Wenn die Folge der Partialsummen in $\tilde{z} \neq z_0$ konvergiert, dann konvergiert die Folge der Polynome

$$(b_0 + b_1(z - z_0) + \dots + b_N(z - z_0)^N)_N$$

gleichmäßig in jedem Kreis

$$B_r(z_0) = \{z : |z - z_0| \leq r\} \quad \text{mit } r < |\tilde{z} - z_0|.$$

Wir gewinnen als Limes eine Funktion $f(z)$ im Innern des Konvergenzkreises $\{z : |z - z_0| < R\}$.

Auch die Folge der Partialsummen von

$$S'(z - z_0) : b_1 + 2b_2(z - z_0) + 3b_3(z - z_0)^2 + \dots$$

konvergiert gleichmäßig in jedem Kreis

$$B_r(z_0) := \{z : |z - z_0| \leq r\}.$$

Wir haben gesehen, daß man von einer formalen Potenzreihe $S(z-z_0)$, d.h. durch Vorgabe der Koeffizientenfolge (im Falle von einem positiven Konvergenzradius) zu einer Funktion $f(z)$ in einem Kreis $\{z : |z-z_0| < R\}$ gelangt. Wir gehen nun von den Funktionen aus. Wir fragen, ob man ein vorgegebenes $f(\cdot)$ (in einem kleinen Kreis um z_0) durch eine Potenzreihe $S(z-z_0)$ darstellen kann. Da wir hier noch keine überzeugenden (d.h. sowohl einfachen als auch allgemeinen) Kriterien nennen können, erledigen wir die Frage durch die

Definition Von einer Funktion $f(\cdot)$, definiert in $D \subset \mathbb{C}$ (offen) sagt man, sie sei holomorph in $z_0 \in D$, wenn es ein $R > 0$ gibt und Zahlen $b_0, b_1, \dots$ so, daß

$$f(z) = \lim_{N \rightarrow \infty} (b_0 + b_1(z-z_0) + \dots + b_N(z-z_0)^N)$$

für alle z mit $|z-z_0| < R$.

Beispiel $f(z) = \frac{1}{1-z}$ definiert für $z \in D = \mathbb{C} \setminus \{1\}$ ist in allen $z_0 \in D$ holomorph.

Beweis

$$\frac{1}{1-z} = \frac{1}{(1-z_0) - (z-z_0)} = \frac{1}{1-z_0} \left(\frac{1}{1-w} \right) \quad \text{mit } w = \frac{z-z_0}{1-z_0}$$

Gleichmäßig in $\{w : |w| \leq r < 1\}$ gilt

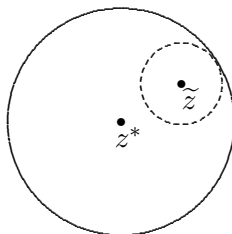
$$\frac{1}{1-w} = 1 + w + w^2 + \dots$$

In $\{z : |z-z_0| \leq r|1-z_0|\}$ gilt also gleichmäßig

$$\begin{aligned} \frac{1}{1-z} &= \frac{1}{1-z_0} \left(1 + \frac{z-z_0}{1-z_0} + \left(\frac{z-z_0}{1-z_0} \right)^2 + \dots \right) \\ &= b_0 + b_1(z-z_0) + b_2(z-z_0)^2 + \dots \end{aligned}$$

mit $b_0 = \frac{1}{1-z_0}$, $b_1 = \frac{1}{(1-z_0)^2}$, $\dots$, $b_n = \frac{1}{(1-z_0)^{n+1}}$, $\dots$

Theorem Wenn $f(\cdot)$ in z^* holomorph ist, dann ist $f(\cdot)$ in einer vollen Umgebung von z^* holomorph.



Es genügt, das Theorem für $z^* = 0$ zu beweisen. Betrachten wir also eine Potenzreihe

$$S(z) : a_0 + a_1 z + a_2 z^2 + \dots,$$

welche in $|z| < R$ die Funktion $f(z)$ darstellt.

Satz Die von $S'(z)$ dargestellte Funktion $g(z)$ hat die folgende Eigenschaft:
Für z, w mit $|z|, |w| \leq r < R$ gilt

$$\left| \frac{f(w) - f(z)}{w - z} - g(z) \right| \leq |w - z| \cdot \frac{1}{2} \sum n(n-1) |a_n| r^{n-2}.$$

Beweis

$$\begin{aligned} w^n - z^n &= (w - z)(w^{n-1} + w^{n-2}z + \dots + wz^{n-2} + z^{n-1}) \\ \frac{w^n - z^n}{w - z} - nz^{n-1} &= (w^{n-1} - z^{n-1}) + (w^{n-2} - z^{n-2})z + \dots + (w - z)z^{n-2} \\ &= (w - z) \left[(w^{n-2} + w^{n-3}z + \dots + z^{n-2}) \right. \\ &\quad \left. + (w^{n-3} + \dots + z^{n-3})z + \dots + (w + z)z^{n-3} + z^{n-2} \right] \end{aligned}$$

Die Anzahl der Summanden ist $(n-1) + (n-2) + \dots + 2 + 1 = \frac{n(n-1)}{2}$.
Jeder ist im Absolutbetrag $\leq r^{n-2}$

$$\begin{aligned} \left| \frac{1}{w - z} (a_n w^n - a_n z^n) - n a_n z^{n-1} \right| &\leq |w - z| |a_n| \frac{1}{2} n(n-1) r^{n-2} \\ \left| \frac{1}{w - z} (f(w) - f(z)) - g(z) \right| &\leq |w - z| \frac{1}{2} \sum n(n-1) |a_n| r^{n-2} \end{aligned}$$

■

Die Abschätzung impliziert eine Eigenschaft von $f(\cdot)$, welche man die in $B_r(0)$ gleichmäßige Differenzierbarkeit nennen könnte.

Satz Wenn $f(\cdot)$ im Nullpunkt holomorph ist, dann existiert eine Funktion $g(\cdot)$ so, daß

$$\forall \varepsilon > 0 \exists \delta > 0 \forall z, w \in B_r(0) \quad \left(|w - z| < \delta \curvearrowright \left| \frac{f(w) - f(z)}{w - z} - g(z) \right| < \varepsilon \right)$$

Diese Funktion $g(\cdot)$ wird durch die formale Ableitung $S'(z)$ dargestellt.

Damit ist der Zusammenhang zwischen der Ableitung einer holomorphen Funktion $f(\cdot)$ und der formalen Ableitung der darstellenden Potenzreihe hergestellt.

Durch Induktion beweist man

Satz Die k -te Ableitung der holomorphen Funktion $f(\cdot)$ wird durch die k -te Ableitung der $f(\cdot)$ darstellenden Potenzreihe dargestellt

$$\frac{1}{k!} f^{(k)}(z_0) = \frac{1}{k!} S^{(k)}(z_0) = \sum_{m=0}^{\infty} \frac{[m+k]_k}{k!} a_{m+k} z_0^m \quad \text{für } |z_0| < R.$$

Satz Für jedes z_0 im Konvergenzkreis und $r < R - |z_0|$ haben wir gleichmäßig in

$$\begin{aligned} B_r(z_0) &= \{|z - z_0| \leq r\} \\ f(z) &= \sum_{k=0}^{\infty} \frac{1}{k!} f^{(k)}(z_0)(z - z_0)^k. \end{aligned}$$

Beweis

$$\begin{aligned} \sum_{k=0}^{\infty} \frac{1}{k!} f^{(k)}(z_0)(z - z_0)^k &= \sum_{k,m=0}^{\infty} \frac{(m+k)!}{m!k!} a_{m+k} z_0^m (z - z_0)^k \\ &= \sum_{n=0}^{\infty} a_n \sum_{k=0}^{\infty} \frac{n!}{k!(n-k)!} z_0^{n-k} (z - z_0)^k \\ &= \sum_{n=0}^{\infty} a_n z^n \end{aligned}$$

Bemerkung Die Annahme $|z_0| + |z - z_0| \leq r < R$ garantiert, daß das Umsummieren der unendlich vielen Summanden unproblematisch ist. Es ist aber durchaus möglich, daß die Potenzreihe

$$\sum_{k=0}^{\infty} \frac{1}{k!} f^{(k)}(z_0)(z - z_0)^k$$

in einem weitaus größeren Kreis als $B_r(z_0)$ gleichmäßig konvergiert. Man findet also u.U. eine holomorphe Fortsetzung der durch $\sum a_n z^n$ gegebenen holomorphen Funktion.

Dieses Phänomen ist der Ausgangspunkt der Theorie der Riemannschen Flächen. Eine wunderschöne Darstellung dieser bezaubernden Theorie ist H. WEYL: “*Die Idee der Riemannschen Fläche*“.

Wer nur einfach einmal die Anfangsgründe der Theorie erlernen will, sei verwiesen auf die Göschenbändchen von KNOPP „*Funktionentheorie*“. Diese Anfangsgründe gehören im gegenwärtigen Lehrbetrieb üblicherweise zum Stoffkanon der Vorlesung *Analysis III*.

B.6 Formale Potenzreihen

Definition Eine Potenzreihe ist ein formaler Ausdruck

$$a_0 + a_1(x - x_0) + a_2(x - x_0)^2 + \dots \quad \text{oder} \quad \sum_0^{\infty} a_n(x - x_0)^n .$$

Die Zahlen a_n heißen die Koeffizienten der Potenzreihe.

Über die Rolle und Bedeutung von $(x - x_0)$ sagen wir vorerst nichts; ebensowenig fragen wir nach dem „Wert“ der Potenzreihe.

Formale Potenzreihen

- a) Die formalen Potenzreihen erfreuen sich heute in der Algebra allgemeiner Wertschätzung. Formale Potenzreihen sind Rechengrößen von der Gestalt

$$a_0 + a_1 \cdot x + a_2 \cdot x^2 + \dots \quad \text{oder} \quad \sum_0^{\infty} a_n \cdot x^n .$$

Die Koeffizienten a_n sind irgendwelche Elemente aus einem Körper. Die Bezeichnung der Unbestimmten tut nichts zur Sache.

Man kann formale Potenzreihen **addieren** und mit Körperelementen multiplizieren. Die Gesamtheit der Potenzreihen ist also ein Vektorraum.

Formale Potenzreihen kann man aber auch **multiplizieren**. Mit dieser Multiplikation gewinnt man eine Algebra. Die Multiplikationsregel ist die folgende

$$\left(\sum a_k x^k \right) \cdot \left(\sum b_\ell x^\ell \right) = \sum c_n x^n ,$$

$$\text{wenn } c_0 = a_0 \cdot b_0 , \quad c_1 = a_0 b_1 + a_1 b_0 , \dots , c_n = \sum_{k+\ell=n} a_k b_\ell .$$

Das Einselement bzgl. der Multiplikation ist

$$1 = 1 + 0 \cdot x + 0 \cdot x^2 + 0 \cdot x^3 + \dots$$

Die Polynome sind spezielle formale Potenzreihen, die nämlich, wo nur endlich viele Koeffizienten $\neq 0$ sind.

- b) Eine spezielle Rolle spielen diejenigen formalen Potenzreihen, deren nullter Koeffizient verschwindet; man kann sie nämlich in eine beliebige formale Potenzreihe **einsetzen**.

Man definiert für

$$S(x) = a_0 + a_1 x + a_2 x^2 + \dots$$

$$T(y) = b_1 y + b_2 y^2 + \dots$$

$$\begin{aligned}
(S \circ T)(y) = S(T(y)) &= a_0 \\
&\quad + a_1(b_1 y + b_2 y^2 + \dots) \\
&\quad + a_2(b_1 y + b_2 y^2 + \dots)^2 + \dots \\
&= a_0 + a_1 b_1 y + (a_1 b_2 + a_2 b_1^2) y^2 \\
&\quad + [a_1 b_3 + a_2 2b_1 b_2 + a_3 b_1^3] y^3 + \dots
\end{aligned}$$

Man rechnet leicht nach

$$\begin{aligned}
(S_1 + S_2)(T(y)) &= S_1(T(y)) + S_2(T(y)) \\
(S_1 \cdot S_2)(T(y)) &= S_1(T(y)) \cdot S_2(T(y)) .
\end{aligned}$$

Man kann nochmals einsetzen $S(T(U(z)))$; und man rechnet leicht nach

$$S(T \circ U) = (S \circ T) \circ U = S \circ T \circ U ,$$

wobei natürlich zu beachten ist, daß in T und U der nullte Koeffizient verschwindet.

- c) Eine spezielle Rolle spielen auch die formalen Potenzreihen, deren nullter Koeffizient *nicht* verschwindet. Zu einer solchen Potenzreihe gibt es nämlich eine **multiplikative Inverse**. Es genügt, den Fall zu untersuchen, wo der führende Koeffizient = 1 ist. Z.B. gilt

$$(1 - x)(1 + x + x^2 + \dots) = 1 = (1 + x + x^2 + \dots)(1 - x)$$

oder allgemeiner für ein T mit verschwindendem führenden Koeffizienten

$$\begin{aligned}
(1 + T(x))(1 - T(x) + T^2(x) - \dots) &= 1 \\
&= (1 - T(x) + T^2(x) - \dots)(1 + T(x)) .
\end{aligned}$$

Damit haben wir die multiplikative Inverse von $1 + T(x)$. Beispielsweise

$$(1 - a_1 x - a_2 x^2)(1 + a_1 x + (a_2 + a_1^2)x^2 + \dots) = 1 .$$

In $T^k(x)$ verschwinden offenbar die ersten k Koeffizienten; in der Summe $1 + T(x) + T^2(x) + \dots$ tragen daher nur endliche viele Summanden zum Koeffizienten von x^n bei.

- d) Für formale Potenzreihen definiert man die **formale Ableitung**

$$\begin{aligned}
S(x) &= a_0 + a_1 x + a_2 x^2 + a_3 x^3 + \dots \\
S'(x) &= a_1 + a_2 x + 3a_3 x^2 + \dots
\end{aligned}$$

Man verifiziert leicht die Produktregel und die Kettenregel

$$\begin{aligned}
(S \cdot T)' &= S' \cdot T + S \cdot T' \\
(S \circ T)'(y) &= S'(T(y)) \cdot T'(y)
\end{aligned}$$

Bei der Kettenregel ist zu beachten, daß der führende Koeffizient des eingesetzten $T(y)$ verschwindet.

- e) **Umkehrung** Für jedes $S(x)$ mit verschwindendem nullten Koeffizienten und nicht verschwindendem ersten Koeffizienten eine Umkehrung existiert; es existiert genau ein $T(y)$ (ebenfalls mit verschwindendem nullten Koeffizienten, so daß

$$S(T(y)) = y ,$$

und für dieses T gilt auch $T(S(x)) = x$.

Man findet diese Umkehrung durch „Koeffizientenvergleich“

$$\begin{aligned} S(x) &= a_1 x + a_2 x^2 + a_3 x^3 + \dots \quad \text{mit } a_1 \neq 0 \\ T(y) &= b_1 y + b_2 y^2 + b_3 y^3 + \dots \\ y &\stackrel{!}{=} S(T(y)) \\ &= a_1 b_1 y + (a_1 b_2 + a_2 b_1^2) y^2 + [a_1 b_3 + a_2 2b_1 b_2 + a_3 b_1^3] y^3 + \dots \end{aligned}$$

Man muß also wählen

$$\begin{aligned} b_1 &\text{ so, daß } a_1 b_1 = 1 \\ b_2 &\text{ so, daß } a_1 b_2 + a_2 b_1^2 = 0 \\ b_3 &\text{ so, daß } a_1 b_3 + a_2 2b_1 b_2 + a_3 b_1^3 = 0 \\ &\dots \end{aligned}$$

Bemerke, daß man aus den b_n die a_n zurückgewinnen kann

$$\begin{aligned} a_1 &\text{ so, daß } b_1 a_1 = 1 \\ a_2 &\text{ so, daß } b_1 a_2 + b_2 a_1^2 = 0 \\ &\dots \end{aligned}$$

Der formale Beweis für „ $S \circ T = \text{id} \iff T \circ S = \text{id}$ “ geht so:

$$S \circ T = \text{id}(y) \implies S(x) = (S \circ T)(S(x)) = S(T \circ S(x)) .$$

Die einzige formale Potenzreihe, die in S eingesetzt S selbst liefert, ist die Identität $\text{id}(x)$; also $(T \circ S)(x) = x$.

Bemerke $S(T(y)) = y \implies 1 = S'(T(y)) \cdot T'(y)$

d.h. $T'(y)$ ist die multiplikative Inverse von $S'(\cdot)$ ausgewertet in $T(y)$.

- f) Den nullten Koeffizienten a_0 der formalen Potenzreihe

$$S(x) = a_0 + a_1 x + a_2 x^2 + \dots$$

nennt man auch den Wert von S im Nullpunkt

$$S(0) = a_0 .$$

Weitere „Werte“ kann man den formalen Potenzreihen im allg. nicht zuweisen. Man kann die Wertzuweisung im Nullpunkt aber dazu benützen um die Koeffizienten von S zu identifizieren. $n!a_n$ ist der Wert im Nullpunkt für n -fach iterierte formale Ableitung von S

$$S(0) = a_0, \quad S'(0) = a_1, \quad S''(0) = 2a_2, \quad \dots, \quad S^{(n)}(0) = n!a_n, \quad \dots$$

$$a_n = \frac{1}{n!} S^{(n)}(0) \quad \text{für } n = 0, 1, 2, \dots$$

$$S(x) = S(0) + \frac{1}{1!} S'(0) \cdot x + \frac{1}{2!} S''(0) \cdot x^2 + \dots$$

Beispiele

a) Eine bemerkenswerte formale Potenzreihe ist

$$\exp(x) := 1 + x + \frac{x^2}{2!} + \frac{x^3}{3!} + \dots$$

Offenbar ist $\exp(x)$ die einzige formale Potenzreihe $f(x)$ mit

$$f'(x) = f(x), \quad f(0) = 1.$$

b) Man definiert

$$\ln(1+x) := x - \frac{1}{2} x^2 + \frac{1}{3} x^3 - \frac{1}{4} x^4 + \dots$$

Die formale Ableitung ist $(\ln(1+x))' = 1 - x + x^2 - x^3 + \dots = \frac{1}{1+x}$.

Der nullte Koeffizient verschwindet. Man kann $\ln(1+x)$ also einsetzen, insbesondere in die formale Potenzreihe $\exp(\cdot)$. Die Rechnung ergibt

$$\exp(\ln(1+x)) = 1+x.$$

Man beweist das am einfachsten mit Hilfe der Kettenregel

$$(\exp(\ln(1+x)))' = \exp(\ln(1+x)) \cdot (\ln(1+x))'$$

Für $\exp(\ln(1+x)) = \sum a_n x^n$ haben wir also

$$(1+x) \sum n a_n x^{n-1} = \sum a_n x^n.$$

Das ergibt $a_0 = 1 = a_1$ und $0 = a_2 = a_3 = \dots$

Setzen wir $x = \exp(y) - 1$, dann erhalten wir

$$\begin{aligned} \exp(\ln \exp(y)) &= \exp(\ln(1 + \exp(y) - 1)) \\ &= 1 + (\exp(y) - 1) = \exp(y) \\ \ln(\exp(y)) &= y. \end{aligned}$$

Beachte Von Werten ist hier nicht die Rede.

c) Für beliebiges reelles α definiert man

$$(1+x)^\alpha := 1 + \alpha x + \frac{\alpha(\alpha-1)}{2!} x^2 + \binom{\alpha}{3} x^3 + \dots = \sum_0^\infty \binom{\alpha}{k} x^k.$$

Diese Notation ist verträglich mit der Definition bei ganzzahligen Exponenten.

$$\begin{aligned} (1+x)^2 &= 1 + 2x + \frac{2 \cdot 1}{2!} x^2 + 0 \cdot x^3 + \dots \\ (1+x)^n &= \sum \binom{n}{k} x^k \quad \text{für } n = 3, 4, 5, \dots \\ (1+x)^{-1} &= 1 - x + \frac{(-1)(-2)}{2!} x^2 + \frac{(-1)(-2)(-3)}{3!} x^3 + \dots \\ &= 1 - x + x^2 - x^3 + \dots \\ (1+x)^{-n+1} &= (1+x)^{-n}(1+x) \quad \text{für } n = 2, 3, \dots \end{aligned}$$

Wir beweisen allgemeiner

$$(1+x)^\alpha \cdot (1+x) = (1+x)^{\alpha+1}$$

$$\begin{aligned} (1+x)^\alpha \cdot (1+x) &= 1 + \alpha x + \frac{\alpha(\alpha-1)}{2!} x^2 + \frac{\alpha(\alpha-1)(\alpha-2)}{3!} x^3 + \dots \\ &\quad + x + \alpha x^2 + \frac{\alpha(\alpha-1)}{2!} x^3 + \dots \\ &= 1 + (\alpha+1)x + \frac{(\alpha+1)\alpha}{2!} x^2 + \frac{(\alpha+1)\alpha(\alpha-1)}{3!} x^3 + \dots \end{aligned}$$

■

Satz Für alle α, β gilt

$$(1+x)^\alpha (1+x)^\beta = (1+x)^{\alpha+\beta}.$$

Beweis

(i) Formales Differenzieren liefert zunächst einmal für alle α

$$\begin{aligned} ((1+x)^{\alpha+1})' &= (\alpha+1) + \frac{(\alpha+1)\alpha}{2!} 2x + \frac{(\alpha+1)\alpha(\alpha-1)}{3!} x^2 + \dots \\ &= (\alpha+1) \left[1 + \alpha x + \frac{\alpha(\alpha-1)}{2!} x^2 + \dots \right] \\ &= (\alpha+1)(1+x)^\alpha. \end{aligned}$$

(ii) Wir benützen die Produktregel

$$\begin{aligned} & \left((1+x)^\alpha (1+x)^\beta \right)' (1+x) \\ &= \left[(1+x)^\alpha \cdot \beta (1+x)^{\beta-1} + \alpha (1+x)^{\alpha-1} (1+x)^\beta \right] (1+x) \\ &= (\alpha + \beta) (1+x)^\alpha (1+x)^\beta \end{aligned}$$

(iii) Andererseits gilt

$$\left((1+x)^{\alpha+\beta} \right)' = (\alpha + \beta) (1+x)^{\alpha+\beta-1}$$

Da beide formale Potenzreihen dieselbe formale Ableitung haben, ergibt der Koeffizientenvergleich die Behauptung.

Satz

$$(1+x)^\alpha = \exp(\alpha \cdot \ln(1+x)) .$$

Beweis mit Hilfe der Kettenregel

$$(\exp(\alpha \cdot \ln(1+x)))' = \exp(\alpha \cdot \ln(1+x)) \cdot \alpha \cdot \frac{1}{1+x}$$

Für $\exp(\alpha \cdot \ln(1+x)) = \sum a_n x^n$ haben wir also $a_0 = 1$ und

$$(1+x) \sum n a_n x^{n-1} = \alpha \sum a_n x^n ,$$

$$n a_n + (n+1) a_{n+1} = \alpha a_n , \quad a_{n+1} = \frac{\alpha - n}{n+1} a_n ;$$

$$a_1 = \alpha a_0 = \alpha , \quad a_2 = \frac{\alpha - 1}{2} a_1 = \frac{\alpha(\alpha - 1)}{2!} \quad \text{usw.}$$

Zur Geschichte der Potenzreihen

Die Potenzreihen waren zunächst eine Spezialität der Nachfolger von Newton, insbesondere B. TAYLOR (1685–1731) und C. MACLAURIN (1698–1746). Taylor wendete die Reihen dieser Art auf die Integration einiger Differentialgleichungen an. In den Herleitungen von Taylor gab es keine Konvergenzbetrachtungen; es war Maclaurin, der mit solchen Betrachtungen anfangte.

LAGRANGE (1736–1813) lehnte die Theorie der Grenzwerte ab, wie sie von NEWTON (1643–1727) angedeutet und von D’ALEMBERT (1717–1783) formuliert worden war. Er konnte sich keine genaue Vorstellung davon machen, was dann geschah, wenn $\frac{\Delta y}{\Delta x}$ seinen Grenzwert annahm. L. CARNOT (1753–1823), der sich ebenso wenig mit Newtons Methode der infinitesimalen Größen befreunden konnte, hat seine Bedenken so formuliert: „*Jene Methode hat die große Unannehmlichkeit, Größen gerade in dem Zustande zu betrachten, in dem sie sozusagen aufhören, Größen zu sein; denn obgleich wir das Verhältniß von zwei Größen immer sehr gut verstehen können, solange sie endlich bleiben, liefert dieses Verhältniß dem Verstand keinen klaren und genauen Begriff, sobald ihre beiden Bestandteile zugleich verschwinden.*“

Die Methode von Lagrange unterschied sich von der seiner Vorgänger. Er ging von der Taylorschen Reihe aus, die er mit Restglied ableitete, wobei er in recht naiver Weise zeigte, daß eine „beliebige“ Funktion $f(x)$ mit Hilfe eines rein algebraischen Prozesses in eine derartige Reihe entwickelt werden kann. Dann wurden die Ableitungen $f'(x), f''(x)$ usw. als Koeffizienten von $h, h^2, \dots$ in der Taylorentwicklung $f(x+h)$ nach Potenzen von h definiert. [Die Schreibweise $f'(x), f''(x)$ stammt von Lagrange].

$$f(x+h) = f(x) + h \cdot f'(x) + \frac{h^2}{2!} f''(x) + \dots$$

Die von Lagrange propagierte abstrakte Behandlung einer Funktion war ein beträchtlicher Schritt vorwärts. Lagrange wurde damit zu einem Gründervater der Theorie der Funktionen einer reellen Veränderlichen. Er wandte seine Methode höchst erfolgreich auf eine Vielzahl von Problemen der Algebra, der Geometrie und vor allem der analytischen Mechanik an.

Die „algebraische“ Methode der Begründung der Infinitesimalrechnung von Lagrange war der Ausgangspunkt für viele Mathematiker des 19. Jahrhunderts, die sich gegen Cauchys Auffassung stellten. Cauchys Insistieren auf Konvergenzbetrachtungen schien ihnen allzu schwerfällig. Bei der herrschenden bequemeren Praxis des formalen Rechnens kamen kaum Fehler vor. Für das Auftreten von Paradoxien fand man immer wieder Argumente, die die Prinzipien der kombinatorischen Herangehensweise nicht antasteten. Es war schon klar, daß „einige“ Funktionen nicht in eine Taylor-Reihe entwickelt werden konnten. Man war aber der Meinung, daß sich das in jedem Einzelfall zeigen würde, indem dann die Taylor-Reihe nicht hingeschrieben werden kann bzw. eine „unbrauchbare oder einen Widerspruch anzeigende Form“ annimmt. (Siehe Jahnke S. 156).

Die auf Potenzreihendarstellungen gegründete „algebraische Analysis“ wurde im 19. Jahrhundert als ein Gegenprogramm zum begrifflichen Vorgehen von Cauchy gepflegt. Ein wichtiger Vertreter der algebraischen Analysis in Deutschland war MARTIN OHM (1792–1872), der

Bruder des Physikers GEORG SIMON OHM (1789–1854). 1822, also gerade parallel zu Cauchys „Cour d’Analyse“, erschienen die ersten beiden Bände seines Hauptwerks „Versuch eines vollkommen consequenten Systems der Mathematik“. Es heißt da in der Einleitung: „... Die Vergleichung der Größen geschieht mittels absoluter ganzer Zahlen. Andere Zahlen gibt es nicht. Die Lehre der Zahlen muß der Lehre der Größen vorangehen, unabhängig vom Begriff der GröÙe. — Die Zahlen und ihre Verbindungen müssen bezeichnet werden (durch Buchstaben). Die Zahlenlehre wird deshalb eine Zahlzeichenlehre und existiert als solche selbständig, unabhängig von dem Begriff der GröÙe. — Man kann zunächst sieben Verbindungen der Zahlen betrachten und erhält dadurch zunächst 7 Zahlzeichen

$$a + b, a - b, a \cdot b, a : b, a^b, \sqrt[b]{a}, a?b,$$

welche bezüglich Summe, Differenz, Produkt, Quotient, Potenz, Wurzel und Logarithme genannte werden. ... Man kann nun die allgemeinsten Begriffe der Gleichung ... entwickeln Diese Elementarformeln ... bilden nun die Grundlage der gesamten Zahlenlehre (die wieder in Arithmetik, Algebra, Analysis, Differential-, Integral- und Variationskalkül etc. classificirt werden kann), insofern keine anderen Zahlverbindungen als die angeführten, und daher keine anderen Ausdrücke, als Summen, Differenzen, Produkte und Quotienten, etc etc vorkommen und zu verbinden seyn können.“ („?“ ist bei Ohm das Symbol für den allgemeinen Logarithmus, während er den natürlichen Logarithmus durch Log bezeichnete) (zitiert nach Jahnke S. 121).

Eine paradigmatische Bedeutung in der Theorie der Potenzreihen hatte die von Newton entdeckte Reihe (1665). In moderner Notation lautet sie

$$(1+x)^\alpha \sim 1 + \alpha \cdot x + \frac{\alpha(\alpha-1)}{2!} x^2 + \frac{\alpha(\alpha-1)(\alpha-2)}{3!} x^3 + \dots$$

wo α eine beliebige reelle Zahl ist. Von einer Bedingung $|x| > 1$ und von irgendwelchen Konvergenzbetrachtungen ist bei Newton nicht die Rede. Mit solchen Reihen rechneten die Anhänger der „algebraischen Analysis“ rein formal. Die Erfolge, die mit solchem formalen Rechnen erzielt wurden, setzten Cauchy und seine Anhänger unter erheblichen Druck.

Ohm hatte 1829 Cauchy kritisiert, weil dieser seiner Auffassung nach ein zuwenig allgemeines Vorgehen wählte und die existierende erfolgreiche Praxis der Mathematik zu sehr beschränkte: „...; dagegen würde es schlimm mit der Anwendung der Analysis aussehen, wenn wir mit Reihen nur dann rechnen könnten, sobald sie konvergent sind, wenn wir deshalb die Sätze mit derjenigen Beschränkung nur stattfinden lassen dürften, welche, um ein Beispiel zu geben, S. 122 für den Gebrauch der Binomialreihe für $(1+x)^\mu$ gegeben wird, die nur so gebraucht werden können für ein gebrochenes μ , wenn x zwischen den Grenzen $+1$ und -1 liegt. Der grösste, ja vielleicht der einzige Werth der Analysis besteht ja gerade darin, daß wir in ihr gerade diejenigen allgemeinen Betrachtungen anstellen, welche von den bestimmten speziellen Werthen der Größen ganz unabhängig sind, oder, um es praktisch auszudrücken, daß wir mit noch ganz unbestimmten Größen rechnen können. Wer möchte noch

analytische Rechnungen anstellen wollen, sobald er keinen Schritt mehr tun könnte, ohne erst alle möglichen Beschränkungen gehörig berücksichtigt zu haben? “ (Das Zitat stammt aus der Rezension der deutschen Übersetzung von Cauchy „Cours d’analyse“ in der Leipziger Literatur-Zeitung; hier zitiert nach Jahnke S. 151/2.)

In seinen Bemühungen, die Stetigkeit als das eigentliche Beweisprinzip der (reellen) Analysis zu etablieren, mußte Cauchy den Glauben in die Allgemeingültigkeit des Taylorschen Satzes erschüttern. Zu diesem Zweck hat er 1823 die Funktion

$$e^{-x^2} + e^{-1/x^2}$$

angegeben. Das Beispiel deutet darauf hin, daß von einer konvergenten Potenzreihe, die sich etwa als Lösung einer Differentialgleichung ergibt, nicht eindeutig auf die Lösungsfunktion geschlossen werden kann. Wie die Verfechter der kombinatorischen Analysis mit Cauchys Beispiel umgegangen sind, beschreibt Jahnke in einigem Detail (S. 157 ff).

Über die allgemeine binomische Reihe hat N. ABEL (1802–1827) eine wichtige Arbeit publiziert. Er führte dem mathematischen Publikum nach dem Vorbild Cauchys vor, wie man unter Verzicht auf den Begriff der allgemeinen (formalen) Potenzreihe mit rein numerisch verstandenen Reihen konkret und erfolgreich rechnen konnte. Abel versuchte aufzudecken, warum die herrschende Praxis recht erfolgreich war, obwohl sie seiner Auffassung nach grundsätzliche Mängel aufwies. Es gab auch Beispiele an, wo die algebraische Analysis zu falschen Aussagen führte (vgl. Jahnke S. 137).

Wirklich transparent wurde die Sachlage erst ganz allmählich durch die Arbeiten von Weierstraß. Eines der großen Verdienste von Weierstraß um die Theorie der Funktionen bestand darin, daß er die Bedeutung des Begriffs der (lokal) gleichmäßigen Konvergenz von Funktionenfolgen herausgearbeitet hat. Weierstraß sieht in seiner Funktionentheorie diejenigen Eigenschaften von Funktionen als wesentlich an, die beim Übergang von endlichen zu unendlichen arithmetischen Ausdrücken erhalten bleiben. Bei seinen Funktionen einer komplexen Veränderlichen, die durch Potenzreihen dargestellt werden, geht die Sache in eine ganz andere Richtung als bei den von ihm ebenfalls studierten Funktionen einer reellen Veränderlichen. In jedem Falle spielt für seine Funktionentheorie die (lokal) gleichmäßige Konvergenz eine entscheidende Rolle (vgl. Richtenhagen, insbesondere S. 17).

Die Techniken der formalen Potenzreihen sind erst viel später wieder zu Ehren gekommen, als die Mathematiker endlich gelernt hatten zwischen der syntaktischen Auffassung von Formeln für Funktionen und der semantischen Bedeutung, wo Gleichheit numerische Gleichheit bedeutet, scharf zu unterscheiden.

Der Zugang zur Taylorschen Formel, den wir im nächsten Kapitel ausführen, ist neueren Datums; er benützt den Begriff der Totalstetigkeit, der eigentlich in die Integrationstheorie von H. Lebesgue gehört. (Diese wird in Analysis II ausführlich behandelt.)

B.7 Totalstetige Funktionen; die Taylorsche Formel

Definition 1 Man sagt von einer Funktion $F(\cdot)$, sie sei auf dem Intervall (a, b) **totalstetig**, wenn eine Funktion $f(\cdot)$ existiert, so daß

$$F(x) = F(x_0) + \int_{x_0}^x f(y) dy \quad \text{für alle } x \in (a, b) .$$

(x_0 kann irgendwo in (a, b) gewählt werden.)

Beachte Der Integrand ist nicht ganz eindeutig bestimmt. Man kann ihn auf einer sog. Lebesgueschen Nullmenge beliebig abändern. Man muß den Integranden als eine Äquivalenzklasse von Funktionen auffassen. Je zwei Repräsentanten unterscheiden sich auf einer Lebesgueschen Nullmenge. (Über Nullmengen wird in Analysis II ausführlich gesprochen werden.)

Definition 2

a) $F(\cdot)$ heißt in (a, b) **stetig differenzierbar**, wenn es eine stetige Funktion $f(\cdot)$ gibt, so daß

$$F(x) = F(x_0) + \int_{x_0}^x f(y) dy \quad \text{für alle } x \in (a, b) .$$

b) Die in (a, b) totalstetige Funktion $F(\cdot)$ heißt in x_0 **stetig differenzierbar**, wenn man den Integranden in x_0 stetig wählen kann.

Satz Wenn $F(\cdot)$ totalstetig ist in einer Umgebung von x_0 und stetig differenzierbar in x_0 , dann existiert der Limes

$$\lim_{h \rightarrow 0} \frac{1}{h} (F(x_0 + h) - F(x_0)) .$$

Beweis Sei $f(\cdot)$ stetig in x_0 und

$$F(x) = F(x_0) + \int_{x_0}^x f(y) dy .$$

Wähle (zu $\varepsilon > 0$) δ so klein, daß

$$|y - x_0| < \delta \leadsto |f(y) - f(x_0)| < \varepsilon .$$

Dann gilt für $|h| < \delta$

$$\left| \frac{1}{h} (F(x_0 + h) - F(x_0)) - f(x_0) \right| = \left| \frac{1}{h} \int_{x_0}^{x_0+h} (f(y) - f(x_0)) dy \right| \leq \varepsilon .$$

■

Warnung Eine Funktion $G(\cdot)$, die in jedem Punkte $x_0 \in (a, b)$ differenzierbar ist, ist nicht notwendigerweise totalstetig. Als Beispiel betrachte man

$$G(x) = \frac{1}{2} x^2 \cdot \sin \frac{1}{x^2}.$$

$G(\cdot)$ ist in $(0, \infty)$ und in $(-\infty, 0)$ stetig differenzierbar, also auch totalstetig. $G(\cdot)$ ist aber nicht totalstetig in $\mathbb{R}$; denn

$$G'(x) = x \cdot \sin \frac{1}{x^2} - \frac{1}{x} \cos \frac{1}{x^2}$$

ist nicht integrierbar in einer Umgebung des Nullpunkts.

Beachte, daß $G(\cdot)$ im Nullpunkt sehr wohl differenzierbar ist:

$$\left| \frac{1}{h} (G(h) - G(0)) \right| \leq \left| \frac{1}{h} \cdot h^2 \right| = |h|.$$

Didaktische These

Bekanntlich bedeutet Stetigkeit von $F(\cdot)$ in einem Intervall die Stetigkeit in jedem Punkt. In der Theorie der Meßbarkeit (Analysis II) und in der Theorie der Differenzierbarkeit sollte man (im Interesse einer Entrümpelung) aber anders vorgehen. Es erscheint uns verfehlt, den Begriff der differenzierbaren Funktion einzuführen, indem man etwa festlegt, eine Funktion solle differenzierbar im Intervall (a, b) heißen, wenn sie in jedem $x_0 \in (a, b)$ differenzierbar ist, wenn also in jedem x_0 der Grenzwert

$$\lim_{h \rightarrow 0} \frac{1}{h} (G(x_0 + h) - G(x_0))$$

existiert. Es gibt nämlich keine nennenswerte Theorie der so umrissenen Funktionenklasse. (Die Freunde des Satzes von Rolle und des sog. Mittelwertsatzes der Differentialgleichung konnten uns nicht überzeugen.)

Die fruchtbaren Begriffsbildungen sind. u.E. einerseits „Totalstetigkeit in einem Intervall“ und andererseits „polynomiale Approximierbarkeit bis zu Ordnung n im Punkte x_0 “.

Definition 3 $F(\cdot)$ heißt polynomial approximierbar von der Ordnung n im Punkte x_0 , wenn es ein Polynom

$$p(x) = a_0 + a_1(x - x_0) + \dots + a_n(x - x_0)^n$$

gibt, so daß

$$\frac{1}{|x - x_0|^n} |f(x) - p(x)| \longrightarrow 0 \quad \text{für } x - x_0 \rightarrow 0,$$

(anders geschrieben: $f(x) - p(x) = o(|x - x_0|^n)$ für $x \rightarrow x_0$.)

Besonders wichtig ist natürlich der Fall $n = 1$. Da lautet die Bedingung einfach: $f(\cdot)$ ist in x_0 differenzierbar.

Zunächst ein Satz über ganz allgemeine totalstetige Funktionen.

Satz Wenn $F(\cdot)$ über (a, b) totalstetig ist, dann gibt es Lipschitzstetige Funktionen $F_M(\cdot)$, welche $F(\cdot)$ auf jedem $[c, d] \subseteq (a, b)$ beliebig genau gleichmäßig approximieren.

Beweis Es existiert $f(\cdot)$ so, daß für alle $x', x'' \in (a, b)$

$$F(x'') - F(x') = \int_{x'}^{x''} f(y) dy .$$

Wenn $f(\cdot)$ beschränkt ist, etwa $|f(x)| \leq M$, dann haben wir

$$|F(x'') - F(x')| \leq M \cdot |x'' - x'| ;$$

also ist $F(\cdot)$ Lipschitzstetig mit Streckungsfaktor $\leq M$.

Im allgemeinen Fall gibt es zu jedem $[c, d] \subseteq (a, b)$ ein M , so daß

$$\int_c^d (|f(y)| \wedge M) dy \geq \int_c^d |f(y)| dy - \varepsilon .$$

Für die „trunkierte“ Funktion

$$f_M(y) = \text{med}\{-M, f(y), M\} = \begin{cases} f(y) & \text{falls } |f(y)| \leq M \\ +M & \text{falls } f(y) > M \\ -M & \text{falls } f(y) < -M \end{cases}$$

liefert (mit $x_0 \in [c, d]$)

$$F_M(x) := F(x_0) + \int_{x_0}^x f_M(y) dy$$

eine Lipschitzstetige Funktion mit

$$|F(x) - F_M(x)| \leq \int_{x_0}^x |f(y) - f_M(y)| dy \leq \varepsilon \quad \text{für alle } x \in [c, d] .$$

Sprechweise

- Totalstetige Funktionen nennen wir auch **einmal ableitbare** Funktionen.
- Was eine **n -mal ableitbare** Funktion ist, definieren wir rekursiv: Eine stetig differenzierbare Funktion $f(\cdot)$ heißt n -mal ableitbar, wenn ihre Ableitung $(n-1)$ -mal ableitbar ist.

Für $0 \leq k \leq n-1$ bezeichnet

$$f^{(k)}(\cdot) \quad \text{oder} \quad D^{(k)}f(\cdot)$$

die k -te Ableitung.

Bemerke $f(\cdot)$ ist genau dann in (a, b) n -mal ableitbar, wenn $f^{(k)}(\cdot)$ in (a, b) $(n - k)$ -mal ableitbar ist.

Beachte Die n -te **Ableitung** einer n -mal ableitbaren Funktion existiert im allg. nicht als eine Funktion; sie ist vielmehr als eine Äquivalenzklasse von Funktionen zu verstehen, nämlich als der Integrand, wenn man die totalstetige Funktion $f^{(n-1)}(\cdot)$ als Stammfunktion schreibt

$$f^{(n-1)}(x) = f^{(n-1)}(x_0) + \int_{x_0}^x f^{(n)}(y) dy .$$

Beispiele

- 1) Ein Polynom ist beliebig oft differenzierbar. Für ein Polynom vom Grade $\leq n$ ist die n -te Ableitung eine Konstante

$$\begin{aligned} f(x) &= a_0 + a_1x + a_2x^2 + \dots + a_nx^n \\ f^{(1)}(x) &= a_1 + 2a_2x + 3a_3x^2 + \dots + na_nx^{n-1} \\ f^{(2)}(x) &= 2a_2 + 3 \cdot 2a_3x + 4 \cdot 3a_4x^2 + \dots + n(n-1)a_nx^{n-2} \\ &= [2]_2a_2 + [3]_2a_3x + [4]_2a_4x^2 + \dots + [n]_2a_nx^{n-2} \\ f^{(3)}(x) &= [3]_3a_3 + [4]_3a_4x + [5]_3a_5x^2 + \dots + [n]_3a_nx^{n-3} \\ &\dots \\ f^{(k)}(x) &= [k]_ka_k + [k+1]_ka_{k+1}x + [k+2]_ka_{k+2}x^2 + \dots + [n]_ka_nx^{n-k} . \end{aligned}$$

Dabei haben wir die Notation benützt

$$[\alpha]_k = \alpha(\alpha-1)(\alpha-2) \cdot \dots \cdot (\alpha-k+1) \quad \text{für } k = 1, 2, \dots, \quad \alpha \in \mathbb{R} .$$

- 2) $f(x) = \sqrt{|x|}$ für $x \in (-\infty, +\infty)$
 $f(\cdot)$ ist einmal ableitbar, aber nicht stetig differenzierbar in $(-\infty, +\infty)$
- 3) $f(x) = x^2 \sin \frac{1}{x}$ für $x \neq 0$, $f(0) = 0$
 $f(\cdot)$ ist einmal ableitbar, aber nicht stetig differenzierbar in $(-\infty, +\infty)$

$$f(x) = \int_0^x [2y \sin \frac{1}{y} - \cos \frac{1}{y}] dy \quad \text{für alle } x \in \mathbb{R} .$$

Der Integrand ist im Nullpunkt unstetig.

- 4) (Das Lemma von Morse im eindimensionalen Fall)

Sei $k(\cdot)$ in einer Umgebung des Nullpunkts zweimal stetig ableitbar mit

$$k(0) = 0, \quad k'(0) = 0, \quad k''(0) > 0 .$$

Setze

$$f(x) = \begin{cases} \sqrt{2k(x)} & \text{für } x \geq 0 \\ -\sqrt{2k(x)} & \text{für } x \leq 0 \end{cases}$$

Dann ist $f(\cdot)$ in einer Umgebung des Nullpunkts stetig differenzierbar.

Beweis In allen x mit $k(x) > 0$ haben wir

$$f'(x) = \pm \frac{\sqrt{2}}{2} \frac{k'(x)}{\sqrt{k(x)}} = \left(\frac{2k(x)}{x^2} \right)^{-1/2} \cdot \frac{1}{x} \int_0^x k''(y) dy .$$

Für $x \rightarrow 0$ konvergiert der erste Faktor gegen $(k''(0))^{-1/2}$ und der zweite gegen $k''(0)$; also gilt

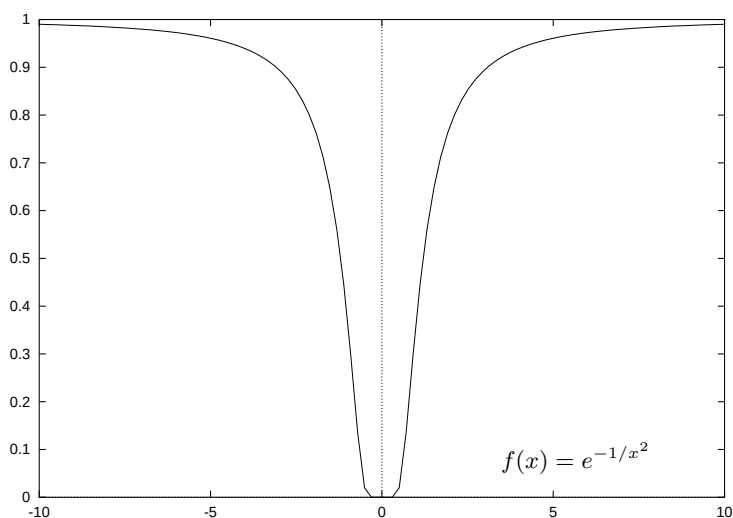
$$\lim_{x \rightarrow 0} f'(x) = \sqrt{k''(0)} .$$

Andererseits ist $f(\cdot)$ im Nullpunkt differenzierbar mit $f'(0) = \sqrt{k''(0)}$.

$$\frac{1}{x} [f(x) - x\sqrt{k''(0)}] = \frac{1}{x} \int_0^x [f'(y) - \sqrt{k''(0)}] dy \longrightarrow 0 \quad \text{für } x \rightarrow 0 .$$

5) $f(x) = \exp\left(-\frac{1}{x^2}\right)$ für $x \neq 0$, $f(0) = 0$

$f(\cdot)$ ist beliebig oft differenzierbar in $(-\infty, +\infty)$.



$$f'(x) = \exp\left(-\frac{1}{x^2}\right) \cdot \frac{2}{x^3}$$

Man sieht leicht, daß $x^{-n} \cdot e^{-1/x^2}$ für alle n nach 0 strebt, wenn $x \rightarrow 0$.

Exkurs (Zerlegung der Einsfunktion)

Gelegentlich möchte man eine unendlich oft differenzierbare Funktion f auf $\mathbb{R}$ darstellen als eine Summe von unendlich oft differenzierbaren Funktionen, die außerhalb von einem kurzen Intervall verschwinden. Die eben behandelte Funktion erweist sich als ein natürliches Instrument. Setze

$$c(x) := \begin{cases} \exp\left(-\frac{1}{x^2}\right) & \text{für } x > 0 \\ 0 & \text{für } x \leq 0 \end{cases}$$

$$C(x) := \text{const} \int_{-\infty}^x c(y) \cdot c(1-y) dy$$

Der Integrand verschwindet außerhalb des Einheitsintervalls; wir wählen die Normierungskonstante so, daß $\int_{-\infty}^x C(x) dx = 1$.

Setze nun

$$H(x) = C(x+1) - C(x).$$

Diese nichtnegative Funktion verschwindet für $|x| \geq 1$. Für alle x gilt

$$\sum_{j=-\infty}^{+\infty} H(x-j) = 1,$$

wobei nur zwei Summanden einen positiven Beitrag liefern.

Für jede unendlich oft differenzierbare Funktion $f(x)$ summieren sich die unendlich oft differenzierbaren Funktionen $f(x) \cdot H(x-j)$ zu $f(x)$ auf.

Satz Taylorsche Formel mit Restglied

Wenn f im Intervall (a, b) n -mal ableitbar ist, dann gilt für alle $x_0, x \in (a, b)$

$$\begin{aligned} f(x) &= f(x_0) + f'(x_0)(x-x_0) + \frac{1}{2!}f^{(2)}(x_0)(x-x_0)^2 \\ &\quad + \dots + \frac{1}{(n-1)!}f^{(n-1)}(x_0)(x-x_0)^{n-1} \\ &\quad + \int_{x_0}^x \frac{1}{(n-1)!}(x-y)^{n-1}f^{(n)}(y) dy \end{aligned}$$

Der letzte Term heißt das Restglied; das Polynom $(n-1)$ -ten Grades heißt das *Taylor-Polynom*

$$f(x) = T_{n-1}(x_0, x) + R_{n-1}(x_0, x).$$

Beweis

- 1) Für $n = 1$ besagt die Taylorsche Formel: Es existiert eine Funktion $f^{(1)}(\cdot)$ (genauer: eine Äquivalenzklasse von Funktionen), so daß

$$f(x) = f(x_0) + \int_{x_0}^x f'(y) dy .$$

Dies ist gerade die Bedingung der einmaligen Ableitbarkeit (= „Totalstetigkeit“).

- 2) Für $n > 1$ ergibt sich die Taylorsche Formel durch Rekursion mittels partieller Integration.

$$\begin{aligned} R_{n-1}(x_0, x) &= \int_{x_0}^x \frac{1}{(n-1)!} (x-y)^{n-1} f^{(n)}(y) dy \\ &= \left[\frac{1}{(n-1)!} (x-y)^{n-1} f^{(n-1)}(y) \right]_{x_0}^x + \int_{x_0}^x \frac{1}{(n-2)!} (x-y)^{n-2} f^{(n-1)}(y) dy \\ &= -\frac{1}{(n-1)!} (x-x_0)^{n-1} f^{(n-1)}(x_0) + R_{n-2}(x_0, x) . \end{aligned}$$

■

Spezialfall (Lagrangesche Form des Restglieds)

Wenn $f^{(n-1)}(\cdot)$ stetig differenzierbar ist, dann existiert für jedes x ein ξ zwischen x_0 und x , so daß

$$\begin{aligned} f(x) &= f(x_0) + f'(x_0)(x-x_0) \\ &\quad + \dots + \frac{1}{(n-1)!} f^{(n-1)}(x_0)(x-x_0)^{n-1} + \frac{1}{n!} f^{(n)}(\xi)(x-x_0)^n . \end{aligned}$$

Beweis Wir beschränken uns auf den Fall $x > x_0$. Betrachte für $y \in [x_0, x]$ die Funktion

$$p(y) = \frac{n!}{(x-x_0)^n} \cdot \frac{1}{(n-1)!} (x-y)^{n-1} \geq 0 \quad \text{für } x_0 \leq y \leq x .$$

Es gilt $\int_{x_0}^x p(y) dy = 1$.

$$\begin{aligned} \inf\{f^{(n)}(y) : y \in [x_0, x]\} &\leq \int_{x_0}^x p(y) \cdot f^{(n)}(y) dy \\ &\leq \sup\{f^{(n)}(y) : y \in [x_0, x]\} . \end{aligned}$$

Da $f^{(n)}(\cdot)$ stetig ist, gibt es nach dem Zwischenwertsatz ein $\xi \in [x_0, x]$ mit

$$\int_{x_0}^x p(y) \cdot f^{(n)}(y) dy = f^{(n)}(\xi) \cdot \int_{x_0}^x p(y) dy.$$

Zusatz Von einer komplexwertigen Funktion auf (a, b) sagt man, sie sei n -mal ableitbar, wenn sowohl der Realteil als auch der Imaginärteil n -mal ableitbar ist.

Beispiel Die Funktion

$$f(x) = e^{ix} = \cos x + i \sin x$$

ist in $(-\infty, +\infty)$ beliebig oft ableitbar. Es gilt

$$f^{(k)}(x) = i^k \cdot f(x) \quad \text{für alle } k$$

$$\left| e^{ix} - 1 - \frac{ix}{1!} - \frac{(ix)^2}{2!} - \dots - \frac{(ix)^{n-1}}{(n-1)!} \right| \leq \frac{|x|^n}{n!}.$$

Zu zeigen ist die Abschätzung des Restglieds

$$\left| \int_0^x \frac{1}{(n-1)!} (x-y)^{n-1} \cdot f^{(n)}(y) dy \right| \leq \frac{|x|^n}{n!}.$$

Wir haben

$$|f^{(n)}(y)| = |i^n \cdot e^{iy}| = 1 \quad \text{für alle } y$$

$$\left| \int_0^x \frac{1}{(n-1)!} (x-y)^{n-1} dy \right| = \frac{|x|^n}{n!}.$$

Satz (Taylorscher Satz ohne Restglied)

Wenn $f(\cdot)$ in (a, b) n -mal ableitbar ist und $f^{(n-1)}(\cdot)$ in x_0 stetig differenzierbar ist, dann ist $f(\cdot)$ polynomial approximierbar von der Ordnung n im Punkte x_0 . M.a.W.: Es existieren $a_0, a_1, \dots, a_n$ so, daß

$$f(x) = a_0 + a_1(x - x_0) + \dots + a_n(x - x_0)^n + o(|x - x_0|^n)$$

(Landaus Notation) .

Beweis Nach der Taylorsche Formel haben wir

$$\begin{aligned} f(x) &= f(x_0) + f'(x_0)(x - x_0) + \dots + \frac{1}{n!} f^{(n)}(x_0)(x - x_0)^n \\ &= \frac{1}{n!} (x - x_0)^n \int_{x_0}^x p(y) \cdot [f^{(n)}(y) - f^{(n)}(x_0)] dy \end{aligned}$$

mit der oben definierten Funktion

$$p(y) = p([x_0, x]; y) = \frac{n!}{(x - x_0)^n} \cdot \frac{1}{(n - 1)!} (x - y)^{n-1}.$$

Das Integral strebt nach 0 für $x \rightarrow x_0$.

Warnung Wenn es zu einer Funktion $f(\cdot)$ Zahlen $a_0, a_1, a_2, \dots$ gibt, so daß für jedes $n \in \mathbb{N}$

$$(x - x_0)^{-n} [f(x) - a_0 - a_1(x - x_0) - \dots - a_n(x - x_0)^n] \rightarrow 0 \quad \text{für } x \rightarrow x_0$$

dann bedeutet das noch lange nicht, daß es ein $x \neq x_0$ gibt, so daß

$$f(x) - \sum_{n=0}^N a_n (x - x_0)^n \xrightarrow{N \rightarrow \infty} 0.$$

Beispiel Die Funktion

$$f(x) = e^{-1/x^2}$$

ist beliebig oft differenzierbar im Nullpunkt. Alle Ableitungen verschwinden. Die Taylor-Polynome sind allesamt identisch 0. Die Folge der Taylor-Polynome konvergiert in keinem Punkt (außerhalb 0) gegen den Funktionswert $f(\cdot)$.

Anhang: Landaus Notation

Ein wichtiges Thema der klassischen Analysis ist das qualitative Studium von Kurven. Man interessiert sich manchmal für das globale Verhalten, manchmal aber auch nur für das ungefähre Verhalten bei Annäherung an ein x_0 (x_0 ist meistens ein Randpunkt des Definitionsbereichs oder auch $\pm\infty$). Die Techniken sind vielfältig; es kommt sehr darauf an, wie die Kurve gegeben ist. Wichtige Fälle sind

$$\begin{aligned} y &= f(x) && \text{ („explizite Gestalt“)} \\ F(x, y) &= \text{const} && \text{ („implizite Funktion“)} \\ y' &= F(x, y) && \text{ („Lösungskurven einer Differentialgleichung“)} \end{aligned}$$

Wenn man eine schwierige Funktion $f(x)$ in der Nähe von x_0 (auch $\pm\infty$) durch eine wohlbekannte einfachere Funktion $g(x)$ approximieren will, dann empfiehlt sich für die Beschreibung des Fehlers eine von E. LANDAU (1877–1938) eingeführte Notation: Der Approximationsfehler wird durch eine Funktion $h(x)$ (für $x \rightarrow x_0$) abgeschätzt.

Notation Seien $f(\cdot)$ und $g(\cdot)$ in $D \setminus \{x_0\}$ wohldefiniert und sei $|h(\cdot)|$ in $D \setminus \{x_0\}$ strikt positiv. Man schreibt:

$$\begin{aligned} \text{a)} \quad & f(x) - g(x) = o(h(x)) \quad \text{für } x \rightarrow x_0 \\ \text{oder} \quad & f(x) = g(x) + o(h(x)) \quad \text{für } x \rightarrow x_0 \end{aligned}$$

und liest: „ $f - g$ ist ein klein o von h für $x \rightarrow x_0$ “, wenn

$$\lim_{x \rightarrow x_0} \frac{f(x) - g(x)}{h(x)} = 0.$$

$$\text{b)} \quad f(x) - g(x) = O(h(x)) \quad \text{für } x \rightarrow x_0$$

„ $f - g$ ist ein Groß O von h für $x \rightarrow x_0$ “, wenn $\frac{f(x)-g(x)}{h(x)}$ beschränkt ist in einer Umgebung von x_0 .

Formal geschrieben:

$$\begin{aligned} \text{a)} \quad & f - g = o(h) \quad \text{für } x \rightarrow x_0 \iff \\ & \iff \forall \varepsilon > 0 \quad \exists \delta > 0 \quad \forall x \in D \setminus \{x_0\} \\ & \quad (d(x, x_0) < \delta \curvearrowright |f(x) - g(x)| \leq \varepsilon \cdot |h(x)|) \end{aligned}$$

$$\begin{aligned} \text{b)} \quad & f - g = O(h) \quad \text{für } x \rightarrow x_0 \iff \\ & \iff \exists \delta > 0 \quad \exists C > 0 \quad \forall x \in D \setminus \{x_0\} \\ & \quad (d(x, x_0) < \delta \curvearrowright |f(x) - g(x)| \leq C \cdot |h(x)|) \end{aligned}$$

Die Notation paßt für reellwertige (oder auch komplexwertige) Funktionen auf einem metrischen Raum $(D, d(\cdot, \cdot))$. Eine Variante ergibt sich, wenn D eine nach oben unbeschränkte Teilmenge von $\mathbb{R}$ ist. Man schreibt dann

$$\begin{aligned} \text{a')} \quad & f - g = o(h) \quad \text{für } x \rightarrow +\infty \iff \\ & \iff \forall \varepsilon > 0 \quad \exists M \quad \forall x \in D \\ & \quad (x \geq M \longrightarrow |f(x) - g(x)| \leq \varepsilon \cdot |h(x)|) \end{aligned}$$

$$\begin{aligned} \text{b')} \quad & f - g = O(h) \quad \text{für } x \rightarrow +\infty \iff \\ & \iff \exists M > 0 \quad \exists C > 0 \quad \forall x \\ & \quad (x \geq M \longrightarrow |f(x) - g(x)| \leq C \cdot |h(x)|) \end{aligned}$$

Beispiele

1) Sei $f(x) = \frac{1}{\sqrt{x}} e^{-x}$ für $x > 0$. Es gilt dann

$$\begin{aligned} f(x) &= O\left(\frac{1}{\sqrt{x}}\right) \quad \text{für } x \rightarrow 0 \\ f(x) &= o(e^{-x}) \quad \text{für } x \rightarrow +\infty \end{aligned}$$

2) Sei $f(x) = \ln x$ für $x > 0$. Es gilt dann für alle $\alpha > 0$

$$f(x) = o(x^\alpha) \quad \text{für} \quad x \rightarrow +\infty$$

$$f(x) = o(x^{-\alpha}) \quad \text{für} \quad x \rightarrow 0.$$

„Der Logarithmus wächst langsamer als jede Potenz.“

3) Für alle n gilt

$$x^n = o(e^x) \quad \text{für} \quad x \rightarrow \infty$$

„Die Exponentialfunktion wächst stärker als jede Potenz für $x \rightarrow \infty$.“

Bemerkung Die Aussagen in 2) und 3) sind in der Tat äquivalent.

Setze $y = \ln x$. $x \rightarrow \infty$ bedeutet dasselbe wie $y \rightarrow \infty$.

$\ln x = o(x^\alpha)$ für $x \rightarrow \infty$ bedeutet dasselbe wie $y = o(e^{\alpha y})$ für $y \rightarrow \infty$.

Mit der Notation von Landau läßt sich der Taylorsche Satz folgendermaßen formulieren:

Satz Wenn $f(\cdot)$ in einer Umgebung von x_0 n -mal ableitbar ist und $f^{(n-1)}(\cdot)$ in x_0 stetig differenzierbar ist, dann gilt

$$\begin{aligned} f(x) &= f(x_0) + f'(x_0)(x - x_0) + \dots \\ &\quad + \frac{1}{n!} f^{(n)}(x_0)(x - x_0)^n + o((x - x_0)^n) \quad \text{für} \quad x \rightarrow x_0. \end{aligned}$$

„In der Nähe von x_0 läßt sich $f(\cdot)$ bis auf einen Fehler $o(|x - x_0|^n)$ durch ein Polynom n -ten Grades approximieren.“

B.8 Totale Differenzierbarkeit

Wir betrachten hier stetige reellwertige Funktionen $F(\cdot)$ auf einer offenen Teilmenge D eines d -dimensionalen linearen Raums. (Der Buchstabe D suggerierte „domain“.)

Wir brauchen (vorübergehend) eine Norm auf dem Vektorraum V der Translationen des $\mathbb{R}$; diese liefert uns auch eine Metrik auf D

$$\rho(x, y) = \|y - x\| \quad x, y \in D.$$

In den konkreten Rechnungen wird V der Raum aller d -Spalten sein. Die Norm könnte dann eine der p -Normen ($p \geq 1$, auch $p = \infty$) sein; es wird aber nicht wichtig sein, welche Norm wir wählen.

V^* bezeichnet den Dualraum von V , also den Vektorraum der Linearformen auf V ; wenn V der Raum der d -Spalten ist, dann ist V^* der Raum der d -Zeilen.

Für $\vartheta \in V^*$ und $x - x_0 \in V$ schreiben wir

$$\vartheta(x - x_0) = (\vartheta_{(1)}, \dots, \vartheta_{(d)}) \begin{pmatrix} x^{(1)} - x_0^{(1)} \\ \vdots \\ x^{(d)} - x_0^{(d)} \end{pmatrix} = \sum_j \vartheta_{(j)} (x^{(j)} - x_0^{(j)}) .$$

(Die Notation mit unteren und oberen Indizes werden wir nicht durchhalten. Sie soll hier nur einmal die Physiker an die Unterscheidung von cogredienten und contragredienten Vektoren erinnern, die in der Tensoranalysis gepflegt wird.)

Wenn V mit der p -Norm ausgestattet ist, dann ist es vernünftig, V^* mit der q -Norm auszustatten, $\frac{1}{p} + \frac{1}{q} = 1$. Es kommt aber auf die in V^* gewählte Norm nicht an. Es hat allerdings schon manchen Anfänger verwirrt, wenn er an $p = 2 = q$ gedacht hat, weil ihm dann die Unterscheidung von V und V^* , die Unterscheidung zwischen Translationen und Linearformen, gefährdet werden könnte.

A) Richtungsableitbarkeit

Definition (Richtungsableitung)

- a) Sei $x_0 \in D$ und $\mathbf{t} \in V$. Wir sagen, die Funktion $F(\cdot)$ besitzt in x_0 eine **Richtungsableitung in der Richtung $\mathbf{t}$** , wenn es eine Zahl $\alpha(x_0)$ gibt, so daß

$$F(x_0 + t \cdot \mathbf{t}) = F(x_0) + \alpha(x_0) \cdot t + o(t) \quad \text{für } t \searrow 0$$

d.h.

$$\lim_{t \searrow 0} \frac{1}{t} [F(x_0 + t \cdot \mathbf{t}) - F(x_0)] = \alpha(x_0) .$$

- b) Wir sagen, $F(\cdot)$ sei **stetig differenzierbar in der Richtung $\mathbf{t}$** , wenn $\alpha(x_0)$ stetig von x_0 abhängt. Die Funktion $\alpha(\cdot) = \alpha_{\mathbf{t}}(\cdot)$ heißt die Richtungsableitung von $F(\cdot)$ in der Richtung $\mathbf{t}$.

(Bemerke: $\alpha_{\lambda \mathbf{t}}(\cdot) = \lambda \cdot \alpha_{\mathbf{t}}(\cdot)$ für alle $\lambda \geq 0$.)

Der Begriff der Richtungsableitung reduziert sich im eindimensionalen Fall im wesentlichen auf den Begriff der rechtsseitigen bzw. linksseitigen Ableitung

$$f'_+(x_0) := \lim_{x \searrow x_0} \frac{1}{x - x_0} (f(x) - f(x_0))$$

$$f'_-(x_0) := \lim_{x \nearrow x_0} \frac{1}{x - x_0} (f(x) - f(x_0))$$

Beispiele

- 1) Die Funktion $F(x) = |x|$ ist in allen Punkten rechtsseitig und linksseitig ableitbar.
 Für $x_0 > 0$ haben wir $f'_+(x_0) = f'_-(x_0) = 1$.
 Für $x_0 < 0$ haben wir $f'_+(x_0) = f'_-(x_0) = -1$.
 Außerdem $f'_+(0) = 1$, $f'_-(0) = -1$. Die Richtungsableitungen von $|x|$ sind also unstetig im Nullpunkt.

- 2) (Verallgemeinerung von 1))

$k(\cdot)$ sei eine konvexe Funktion. In den Übungen haben wir gezeigt, daß die rechtsseitige und die linksseitige Ableitung in allen Punkten existiert. Es handelt sich um isotone Funktionen, die in höchstens abzählbar vielen Punkten verschieden sind; die rechtsseitige Ableitung $f'_+(\cdot)$ ist rechtsseitig stetig, die linksseitige Ableitung ist linksseitig stetig.

- 3) Die Funktion $F(x) = x^2 \sin \frac{1}{x}$, $F(0) = 0$ ist in allen Punkten (auch in $x = 0$) rechtsseitig und linksseitig differenzierbar. Die rechtsseitige und die linksseitige Ableitung ist überall gleich. Die Ableitung ist aber im Nullpunkt unstetig

$$\begin{aligned} f'(0) &= \lim_{x \rightarrow 0} \frac{1}{x} (f(x) - f(0)) = \lim_{x \rightarrow 0} (x \sin \frac{1}{x}) = 0 \\ f'(x_0) &= 2x_0 \sin \frac{1}{x_0} - \cos \frac{1}{x_0} \quad \text{für } x_0 \neq 0. \end{aligned}$$

- 4) Sei $k(\cdot)$ zweimal stetig differenzierbar in einer Umgebung des Nullpunkts mit

$$k(0) = 0, \quad k'(0) = 0, \quad k''(0) > 0$$

$$k(x) = \frac{1}{2} x^2 k''(0) + o(x^2) \quad \text{für } x \rightarrow 0.$$

Setze

$$f(x) = \begin{cases} \sqrt{2k(x)} & \text{für } x \geq 0 \\ -\sqrt{2k(x)} & \text{für } x \leq 0 \end{cases}$$

Behauptung: $f(\cdot)$ ist stetig differenzierbar in einer Umgebung des Nullpunkts.

Beweis Betrachten wir zunächst die $x \neq 0$. Wir haben

$$f'(x) = \pm \frac{\sqrt{2}}{2} \frac{k'(x)}{\sqrt{k(x)}} = \left(\frac{2k(x)}{x^2} \right)^{-1/2} \frac{1}{x} \int_0^x k''(y) dy .$$

Diese stetige Funktion $f'(\cdot)$ konvergiert für $x \rightarrow 0$ gegen

$$[k''(0)]^{-1/2} \cdot k''(0) = \sqrt{k''(0)} .$$

Wir zeigen nun, daß $\sqrt{k''(0)}$ die Ableitung von $f'(\cdot)$ im Nullpunkt ist

$$\frac{1}{x} \left[f(x) - x \cdot \sqrt{k''(0)} \right] = \frac{1}{x} \int_0^x \left(f'(y) - \sqrt{k''(0)} \right) dy \longrightarrow 0 \quad \text{für } x \rightarrow 0 ,$$

d.h. $f(x) = x \cdot \sqrt{k''(0)} + o(|x|)$ für $x \rightarrow 0$.

B) Totale Differenzierbarkeit

Definition

- a) Die Funktion $F(\cdot)$ heißt **total differenzierbar** in x_0 , wenn es eine Linearform $\vartheta = \vartheta_{x_0}$ gibt, so daß

$$F(x) = F(x_0) + \vartheta(x - x_0) + o(\|x - x_0\|) \quad \text{für } x \rightarrow x_0 .$$

Wenn eine solche Linearform existiert, bezeichnen wir sie mit $F'(x_0)$ oder mit $\text{grad}F(x_0)$.

- b) Wenn $F'(x)$ in einer Umgebung von x_0 existiert und in x_0 stetig ist, sagen wir, $F(\cdot)$ sei stetig differenzierbar in x_0 .
- c) Wenn $F'(x)$ überall in D existiert und stetig ist, dann sagen wir, $F(\cdot)$ sei **stetig differenzierbar** in D .

Satz B.8.1 (*Produktregel*)

Wenn F und G in x_0 total differenzierbar sind, dann ist auch $F \cdot G$ in x_0 total differenzierbar mit

$$(F \cdot G)'(x_0) = F(x_0) \cdot G'(x_0) + G(x_0) \cdot F'(x_0) .$$

Beweis

$$\begin{aligned}
F(x) - F(x_0) &= F'(x_0)(x - x_0) + o(\|x - x_0\|) \quad \text{für } x \rightarrow x_0 \\
G(x) - G(x_0) &= G'(x_0)(x - x_0) + o(\|x - x_0\|) \\
(F \cdot G)(x) - (F \cdot G)(x_0) &= [F(x) - F(x_0)][G(x) - G(x_0)] \\
&\quad + F(x_0)[G(x) - G(x_0)] + G(x_0)[F(x) - F(x_0)] .
\end{aligned}$$

Der erste Term ist $O(\|x - x_0\|^2)$ also auch $o(\|x - x_0\|)$.

Neben der Produktregel $(F \cdot G)' = F \cdot G' + G \cdot F'$ ist die **Kettenregel** die bedeutendste Rechenregel für die totale Differentiation. Symbolisch geschrieben lautet sie

$$(H(G(\cdot)))' = H'(G(\cdot)) \cdot G'(\cdot) .$$

Wir beweisen sie im nächsten Abschnitt.

Hinweis Die Funktion f könnte man zweimal (total) differenzierbar im Punkt x_0 nennen, wenn es eine Linearform ϑ und eine symmetrische Bilinearform $B(\cdot, \cdot)$ auf V gibt, so daß

$$F(x) = F(x_0) + \vartheta(x - x_0) + \frac{1}{2} B(x - x_0, x - x_0) + o(\|x - x_0\|^2) .$$

Diejenigen, die in der linearen Algebra noch nichts von symmetrischen Bilinearformen gehört haben, mögen diese Definition hier übergehen. Man schreibt diese Bilinearform als eine symmetrische Matrix („Hesse-Matrix“) $F''(x_0)$ und hat dann

$$F(x) = F(x_0) + F'(x_0)(x - x_0) + \frac{1}{2}(x - x_0)^\top \cdot F''(x_0) \cdot (x - x_0) + o(\|x - x_0\|^2) .$$

Der Begriff der zweimaligen totalen Differenzierbarkeit in einem Punkt scheint unwichtig zu sein. Wichtig dagegen ist der Begriff der zweimaligen stetigen Differenzierbarkeit, wie wir sehen werden.

C) Die Jacobi-Matrix (Funktionalmatrix)

Ein m -tupel von (in x_0) total differenzierbaren Funktionen $G^{(1)}(\cdot), \dots, G^{(m)}(\cdot)$ fassen wir als eine Abbildung auf

$$G = \begin{pmatrix} G^{(1)} \\ \vdots \\ G^{(m)} \end{pmatrix} : \mathbb{R}^d \supseteq D \longrightarrow \mathbb{R}^m .$$

Die (als Zeilen geschriebenen) Gradienten fassen wir zu einer $m \times d$ -Matrix zusammen. Diese Matrix heißt die Jacobi-Matrix der Abbildung $G(\cdot)$ oder auch die Funktionalmatrix

$$J(x_0) = \begin{pmatrix} \text{grad } G^{(1)}(x_0) \\ \vdots \\ \text{grad } G^{(m)}(x_0) \end{pmatrix} = \left(\frac{\partial G^{(i)}}{\partial x_j} \right)_{i=1,\dots,m; j=1,\dots,d}.$$

(C.G. Jacobi, 1804–1851)

Definition

- a) Die Abbildung $G(\cdot)$ von einem Gebiet $D \subseteq \mathbb{R}^d$ in den $\mathbb{R}^m$ heißt **total differenzierbar in** $x_0 \in D$, wenn eine $m \times d$ -Matrix $J(x_0)$ existiert, so daß

$$G(x) = G(x_0) + J(x_0)(x - x_0) + o(\|x - x_0\|) \quad \text{für } x \rightarrow x_0.$$

(Diese Matrix ist notwendigerweise die Jacobi-Matrix; man notiert sie auch $G'(x_0)$.)

- b) Die Abbildung heißt **stetig differenzierbar in** x_0 , wenn $J(x)$ in einer Umgebung von x_0 existiert und $J(\cdot)$ im Punkte x_0 stetig ist.
- c) Die Abbildung heißt **stetig differenzierbar in** D , wenn sie in allen $x \in D$ stetig differenzierbar ist.

Satz B.8.2 (Kettenregel)

Betrachte $F(\cdot) = H(G(\cdot))$. Wenn $G(\cdot)$ in $x_0 \in \mathbb{R}^d$ total differenzierbar ist mit Werten im $\mathbb{R}^m$ und $H(\cdot)$ in $G(x_0)$ total differenzierbar ist (mit Werten im $\mathbb{R}^n$), dann ist die zusammengesetzte Abbildung in x_0 total differenzierbar. Es gilt

$$H(G(\cdot))'(x_0) = H'(G(x_0)) \cdot G'(x_0).$$

Beweis Mit $y_0 = G(x_0)$, $y \in G(x)$ gilt

$$\begin{aligned} H(G(x)) - H(G(x_0)) &= H'(G(x_0)) \cdot H'(x_0) \cdot (x - x_0) \\ &= [H(y) - H(y_0) - H'(y_0) \cdot (y - y_0)] + H'(y_0) [(y - y_0) - G'(x_0)(x - x_0)]. \end{aligned}$$

Der zweite Summand ist $o(\|x - x_0\|)$.

Der erste ist $o(\|y - y_0\|)$ d.h. $o(\|G(x) - G(x_0)\|)$.

Wegen der totalen Differenzierbarkeit von $G(\cdot)$ haben wir

$$\|G(x) - G(x_0)\| \leq \text{const} \cdot \|x - x_0\|.$$

Also ist auch der erste Summand $o(\|x - x_0\|)$.

q.e.d.

D) Stetige partielle Ableitungen

Notation (Partielle Ableitungen)

Sei D eine offene Menge im $\mathbb{R}^d$ (d -Spalten), $x_0 \in D$.

Wenn für den j -ten Einheitsvektor $\mathbf{e}_j$ der Limes

$$\lim_{t \rightarrow 0} \frac{F(x_0 + t\mathbf{e}_j) - F(x_0)}{t} = \alpha_j \quad \text{für } j = 1, \dots, d,$$

existiert, dann heißt α_j die j -te **partielle Ableitung** in x_0 . Man schreibt $\alpha_j = D_j F(x_0)$ oder auch $\alpha_j = \frac{\partial}{\partial x_j} F(x_0)$.

Satz B.8.3 Genau dann ist $F(\cdot)$ in x_0 total differenzierbar, wenn es eine d -Zeile $(\alpha_1, \dots, \alpha_d)$ gibt, so daß gleichmäßig in $\{\mathbf{t} : \|\mathbf{t}\| \leq 1\}$ gilt

$$\lim_{t \rightarrow 0} \left[\frac{F(x_0 + t \cdot \mathbf{t}) - F(x_0)}{t} - (\alpha_1, \dots, \alpha_d) \cdot \mathbf{t} \right] = 0.$$

Der Beweis ist trivial.

Satz B.8.4 (Kriterium für totale Differenzierbarkeit)

Wenn die partiellen Ableitungen von $F(\cdot)$ in einer Umgebung von x_0 existieren und in x_0 stetig sind, dann ist $F(\cdot)$ im Punkte x_0 total differenzierbar.

Beweis Betrachte $\mathbf{t}$ mit $\|\mathbf{t}\| \leq 1$, $\mathbf{t} = c_1\mathbf{e}_1 + c_2\mathbf{e}_2 + \dots + c_d\mathbf{e}_d$. Für kleines $|t|$ betrachten wir den Streckenzug durch die Punkte $x_0, x_1, \dots, x_d$, wobei

$$x_1 = x_0 + t c_1 \mathbf{e}_1, \quad x_2 = x_0 + t(c_1 \mathbf{e}_1 + c_2 \mathbf{e}_2), \dots, \quad x_d = x_0 + t \cdot \mathbf{t}$$

Für kleines t verläuft er ganz in einem Gebiet, in welchem sich die partiellen Ableitungen $F_1(\cdot), \dots, F_d(\cdot)$ um weniger als ε (beliebig klein vorgegbares $\varepsilon > 0$) von den Zahlen $F_1(x_0), \dots, F_d(x_0)$ unterscheiden. Andererseits gilt

$$\begin{aligned} & -F(x_0) + F(x_0 + t\mathbf{t}) \\ &= [-F(x_0) + F(x_1)] + [-F(x_1) + F(x_2)] + \dots + [-F(x_{d-1}) + F(x_d)] \\ &= c_1 \int_0^t F_1(x_0 + s c_1 \mathbf{e}_1) ds + c_2 \int_0^t F_2(x_1 + s c_2 \mathbf{e}_2) ds + \dots \\ &\quad \dots + c_d \int_0^t F_d(x_{d-1} + s c_d \mathbf{e}_d) ds \\ &= c_1 t (F_1(x_0) + \varepsilon_1) + c_2 t (F_2(x_0 + \varepsilon_2) + \dots + c_d t (F_d(x_0) + \varepsilon_d) \\ &\quad \text{mit } |\varepsilon_j| < \varepsilon \quad \text{für alle } j \\ &\left| \frac{F(x_0 + t\mathbf{t}) - F(x_0)}{t} - (c_1 F_1(x_0) + \dots + c_d F_d(x_0)) \right| \leq \varepsilon^* . \end{aligned}$$

Für alle $\mathbf{t}$ mit $\|\mathbf{t}\| \leq 1$ und $|t| \rightarrow 0$ gilt also

$$\frac{F(x_0 + t\mathbf{t}) - F(x_0)}{t} - \text{grad}F(x_0) \cdot \mathbf{t} \longrightarrow 0 \quad \text{gleichmäßig in } \|\mathbf{t}\| \leq 1.$$

Satz B.8.5 Wenn die partiellen Ableitungen in einer sternförmigen Umgebung von x_0 stetig sind, dann existiert dort eine stetige zeilenwertige Funktion

$$M(x) = (M_1(x), \dots, M_d(x)) \quad M(x_0) = \text{grad}F(x_0)$$

so daß

$$F(x) = F(x_0) + M(x) \cdot (x - x_0).$$

Beweis

$$\begin{aligned} F(x) - F(x_0) &= \int_0^1 \frac{d}{dt} F(x_0 + t(x - x_0)) dt \\ &= \sum_{j=1}^d \int (D_j F)(x_0 + t(x - x_0)) (x^{(j)} - x_0^{(j)}) dt \\ &= \sum_{j=1}^d M_j(x) (x^{(j)} - x_0^{(j)}) \\ &= M(x) \cdot (x - x_0) \end{aligned}$$

Satz B.8.6 Sei $F(\cdot)$ total differenzierbar in einer Umgebung von x_0 mit partiellen Ableitungen $F_j(x) = (D_j F)(x)$, welche selbst wieder partielle Ableitungen besitzen $(D_k F_j)(x) = (D_k D_j F)(x)$, die in x_0 stetig sind. Dann gilt für alle k, j

$$D_k D_j F(x_0) = D_j D_k F(x_0).$$

Beweis Betrachte F_j, F_k, F_{jk}, F_{kj} in dem Viereck mit den Ecken $x_0, x_j = x_0 + t\mathbf{n}_j$, $x_k = x_0 + t\mathbf{n}_k$, $x = x_0 + t(\mathbf{n}_j + \mathbf{n}_k)$. Wenn $|t|$ klein ist, dann sind diese Funktionen bis auf ein beliebig klein vorgegbares $\varepsilon > 0$ durch die Werte $F_j(x_0), F_k(x_0), F_{jk}(x_0), F_{kj}(x_0)$ zu approximieren. Andererseits haben wir

$$\begin{aligned} (F(x) - F(x_j)) - (F(x_k) - F(x_0)) &= (F(x) - F(x_k)) - (F(x_j) - F(x_0)) \\ (F(x) - F(x_j)) - (F(x_k) - F(x_0)) &= \int_0^t (F_k(x_j + s\mathbf{e}_k) - F_k(x_0 + s\mathbf{e}_k)) ds \\ &= \int_0^t \left(\int_0^t (D_j F_k)(x_0 + s\mathbf{e}_k + u\mathbf{e}_j) du \right) ds \\ &= t^2 [(D_j F_k)(x_0) + \varepsilon] \end{aligned}$$

Dieselbe Rechnung auf die rechte Seite angewandt ergibt

$$t^2 [(D_k F_j)(x_0) + \varepsilon''] \quad \text{mit } |\varepsilon'|, |\varepsilon''| < \varepsilon.$$

Dies ergibt $D_k F_j(x_0) = D_j F_k(x_0)$. Anders geschrieben

$$\frac{\partial^2 F}{\partial x_j \partial x_k}(x_0) = \frac{\partial^2 F}{\partial x_k \partial x_j}(x_0).$$

Korollar Wenn $F(\cdot)$ in $x_0 \in \mathbb{R}^d$ stetig differenzierbare partielle Ableitungen besitzt, dann existiert eine d -Zeile $F'(x_0)$ und eine symmetrische $d \times d$ -Matrix $F''(x_0)$, so daß

$$F(x) = F(x_0) + F'(x_0)(x - x_0) + \frac{1}{2}(x - x_0)^\top \cdot F''(x_0)(x - x_0) + o(\|x - x_0\|^2).$$

Die Matrix $F''(x_0)$ heißt die **Hesse-Matrix**; ihre Einträge sind die zweiten partiellen Ableitungen. (O. HESSE 1811–1874)

E) Komplexe Differenzierbarkeit

Definition

- a) Eine komplexwertige Funktion $f(\cdot)$, die in einer Umgebung von $z^* \in \mathbb{C}$ definiert ist, heißt in z^* komplex differenzierbar, wenn es eine komplexe Zahl γ gibt, so daß

$$f(z) - f(z^*) = \gamma(z - z^*) + o(|z - z^*|).$$

Man schreibt $\gamma = f'(z^*)$.

- b) Eine in einer offenen Teilmenge D von $\mathbb{C}$ definierte Funktion heißt komplex differenzierbar, wenn sie in jedem Punkt komplex differenzierbar ist.

Der folgende Satz, den E. GOURSAT (1858–1936) um 1900 bewiesen hat, hat sich als ein fundamentales Resultat der Funktionentheorie erwiesen.

Theorem B.8.7 Wenn $f(\cdot)$ in D komplex differenzierbar ist, dann existiert zu jedem z^* ein Kreis, in welchem $f(\cdot)$ durch eine konvergente Potenzreihe dargestellt wird.

Bemerke Insbesondere folgt also aus der komplexen Differenzierbarkeit in allen Punkten von G die stetige Differenzierbarkeit, die stetige Differenzierbarkeit der Ableitung

$$f'(z) := \lim_{w \rightarrow 0} \frac{f(z+w) - f(z)}{w} \quad \text{und so weiter.}$$

Eine in D komplex differenzierbare Funktion ist automatisch unendlich oft differenzierbar. Der Beweis wird in jeder Einführung in die komplexe Funktionentheorie geführt. Stichworte sind der Cauchysche Integralsatz und die Cauchysche Integralformel.

Sei $f(\cdot)$ in D komplex differenzierbar

$$\begin{aligned} z &= x + iy \quad \text{für } z \in D \\ f(z) &= u + iv = u(x, y) + iv(x, y) \end{aligned}$$

Für $z \rightarrow z^*$ gilt

$$\begin{aligned} f(z) - f(z^*) &= \\ &= [u(x, y) - u(x^*, y^*)] + i[v(x, y) - v(x^*, y^*)] \\ &= \left[\frac{\partial u}{\partial x} \bigg|_{(x^*, y^*)} \cdot (x - x^*) + \frac{\partial u}{\partial y} \bigg|_{(x^*, y^*)} (y - y^*) \right] \\ &\quad + i \left[\frac{\partial v}{\partial x} \bigg|_{(x^*, y^*)} \cdot (x - x^*) + \frac{\partial v}{\partial y} \bigg|_{(x^*, y^*)} (y - y^*) \right] + o(|z - z^*|) . \end{aligned}$$

Andererseits gilt mit $f'(z^*) = \alpha + i\beta$

$$\begin{aligned} f(z) - f(z^*) &= \\ &= f'(z^*)(z - z^*) + o(|z - z^*|) \\ &= [\alpha(x - x^*) - \beta(y - y^*)] + i[\alpha(y - y^*) + \beta(x - x^*)] + (z - z^*) + o(|z - z^*|) . \end{aligned}$$

Daraus ergibt sich

$$\begin{aligned} \frac{\partial u}{\partial x} \bigg|_{(x^*, y^*)} &= \alpha = \frac{\partial v}{\partial y} \bigg|_{(x^*, y^*)} \\ \frac{\partial u}{\partial y} \bigg|_{(x^*, y^*)} &= -\beta = -\frac{\partial v}{\partial x} \bigg|_{(x^*, y^*)} . \end{aligned}$$

Der Realteil $u(\cdot, \cdot)$ und der Imaginärteil $v(\cdot, \cdot)$ der komplex differenzierbaren Funktion $f(\cdot)$ erfüllen die „Cauchy–Riemannschen Differentialgleichungen“

$$\boxed{\frac{\partial u}{\partial x} = \frac{\partial v}{\partial y}, \quad \frac{\partial u}{\partial y} = -\frac{\partial v}{\partial x}}$$

Daraus ergibt sich, daß $u(\cdot, \cdot)$ und $v(\cdot, \cdot)$ „harmonische“ Funktionen sind, d.h. daß $u(\cdot, \cdot)$ und $v(\cdot, \cdot)$ Lösungen der sogenannten Laplaceschen Gleichung sind

$$\left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} \right) h = 0 , \quad \text{kurz geschrieben } \Delta h = 0 .$$

Die harmonischen (und noch mehr die superharmonischen) Funktionen spielen eine zentrale Rolle in der Potentialtheorie, einem wichtigen Kapitel der Reellen Analysis, das wir aber hier nicht behandeln können.

F) Bemerkungen zur Geschichte der Theorie der Differenzierbarkeit

Bei einer unendlich oft differenzierbaren reellen Funktion $f(\cdot)$ in einem Gebiet des $\mathbb{R}^d$ sagt das Verhalten in der Nähe eines Punkts x^* nichts aus über das Verhalten in irgendeinem anderen Punkt. Es gibt unendlich oft differenzierbare Funktionen $h(\cdot)$, die in einer vorgegebenen Umgebung U^* von x^* identisch $= 1$ sind und außerhalb einer vorgegebenen Umgebung $U^{**} \supseteq U^*$ identisch verschwinden (vgl. Übungen). Die Funktion $f \cdot h(\cdot)$ stimmt also in U^* mit $f(\cdot)$ überein, während sie außerhalb U^{**} verschwindet.

Bei komplex differenzierbaren Funktionen ist die Lage ganz anders. Das Verhalten in einer beliebig kleinen Umgebung bestimmt den globalen Verlauf. Diese Eigenschaft der komplex differenzierbaren Funktionen entsprach der bis ins 20. Jahrhundert hinein herrschenden Philosophie der Mathematik viel eher als der mangelnde Zusammenhang, den man bei den stetigen Funktionen im Sinne von BOLZANO, CAUCHY, DIRICHLET, RIEMANN u.a. feststellen mußte. Die meisten Mathematiker des 19. Jahrhunderts erwarteten von einer wertvollen Theorie der stetigen (oder glatten) Funktionen, daß sie ein „Prinzip der Kontinuität“ bewahrt. Funktionen sollten ja insbesondere dazu dienen, naturgesetzliche, vorausberechenbare Entwicklungen zu fassen. Ein Schlüsselbegriff mußte nach der vorherrschenden Meinung die „Kontinuität“ oder die „Stetigkeit“ sein.

Für die Geometrie hat PONCELET (1788–1867), ein Gründervater der projektiven Geometrie, das „Prinzip der Stetigkeit“ so ausgedrückt: „Wenn eine Figur aus einer anderen durch eine stetige Veränderung hervorgeht und ebenso allgemein ist wie die erste, dann kann eine für die erste Figur bewiesene Eigenschaft ohne erneute Untersuchung auf die andere übertragen werden.“ (zitiert nach Struik, S.191).

H. HANKEL (1839–1873) formulierte für die Algebra ein „Permanenzprinzip“ mit einer ähnlichen Zielrichtung. Hankel beschäftigte sich dann auch mit der Theorie der Funktionen und empfand es als außerordentlich enttäuschend, daß die Funktionentheorie in der Tradition von Cauchy offenbar kein solches Permanenzprinzip zuließ. Es war aber in der Mitte des 19. Jahrhunderts nicht zu verleugnen, daß die Funktionentheorie in der Tradition von Lagrange, die auf den Potenzreihenentwicklungen aufbaute, nicht auf die neue Theorie der Funktionen paßte, auf eine Theorie, die allein auf Axiome aufgebaut war, die also insbesondere die Stetigkeit (in der Form der Stetigkeitsdefinition von Bolzano und Cauchy) als ein Axiom behandelte. Die „mathematische Strenge“ erwies sich als notwendig, um die neue Auffassung von den stetigen Funktionen zu befreien von den tradierten intuitiven Vorstellungen von Stetigkeit. Solange die Intuition die neuen mathematischen Objekte nicht erfaßte, stand sie in Konkurrenz zur „Strenge“.

Für die komplexe Funktionentheorie verlief die Entwicklung ganz anders. Es entstanden da im 19. Jahrhundert rivalisierende Ansätze; man sprach vom Cauchy–Riemannschen An-

satz im Gegensatz zum Weierstraßschen Ansatz. Während Cauchy von der komplexen Differenzierbarkeit ausgehend über die „Cauchysche Integralformel“ zur Potenzreihenentwicklung gelangt, waren bei Weierstraß die (lokal) durch Potenzreihen gegebenen Funktionen der Ausgangspunkt. In der Theorie von Weierstraß werden die Eigenschaften der Funktionen dadurch erschlossen, daß die Eigenschaften der sie definierenden Potenzreihen untersucht wurden. „Eine ... Eigenart bei Weierstraß ist“, so schreibt 1875 der schwedische Mathematiker und Weierstraß-Schüler G. Mittag-Leffler, „daß er alle allgemeinen Definitionen und alle Beweise, die sich auf Funktionen insgesamt beziehen, vermeidet. Für ihn ist eine Funktion eine Potenzreihe und aus der Potenzreihe deduziert er alles“. (zitiert nach Richenhagen, S. 9).

Glücklicherweise hat sich die Weierstraßsche komplexe Funktionentheorie etwa um 1900 als inhaltlich äquivalent erweisen mit der komplexen Funktionentheorie von Cauchy und Riemann. Die endgültige Vereinigung der beiden im 19. Jahrhundert koexistierenden Theorien bewerkstelligte E. GOURSAT (1858–1936) mit dem Beweis, daß jede komplexwertige Funktion, die in jedem Punkt z_0 eines Gebiets komplex differenzierbar ist, holomorph ist: zu jedem z_0 gibt es eine Umgebung, in welcher die Funktion durch eine konvergente Potenzreihe dargestellt wird. Alle höheren Ableitungen existieren und die Ableitungen in einem einzigen Punkt z_0 bestimmen den globalen Verlauf der Funktion. Bei den analytischen Funktionen gibt es also einen direkten Zusammenhang zwischen dem lokalen Verhalten und dem globalen Verhalten. Man könnte diesen „Identitätssatz für holomorphe Funktionen“ als ein Exempel für das Kontinuitätsprinzip (im Sinne von Poncelet z.B.) ansehen. Der Identitätssatz kam jedenfalls den philosophischen Vorstellungen von einer mathematischen Erfassung der Weltgesetze viel mehr entgegen als die Theorie der reellen Funktionen im Sinne von Cauchy oder gar eine auf den Dirichletschen Funktionenbegriff aufbauende reelle Analysis. In der reellen Analysis ist es außerordentlich mühsam, einen Zusammenhang zwischen dem lokalen und dem globalen Verhalten der Funktionen herzustellen; die Theorie der Differentialgleichungen handelt davon. Vielen Mathematikern des 19. Jahrhunderts galt der Mangel an Glattheit bei den stetigen Funktionen im Sinne von Cauchy zunächst nur als ein bedauerlicher Schönheitsfehler. Immerhin haben viele geglaubt, daß jede stetige Funktion einer reellen Veränderlichen mit Ausnahme (evtl. auch einer unendlichen Anzahl) isoliert liegender Stellen überall differenzierbar ist. Ein renommierter Mathematiker wie JOSEPH BERTRAND (1822–1900) „bewies“ das auch in seinem Lehrbuch über Differential- und Integralrechnung (1864). Gegenbeispiele von Bolzano (1834) und Riemann (1854) waren von der Fachwelt nicht beachtet worden. Erst die Konstruktionen von Weierstraß schreckten die Mathematiker auf (siehe Edwards S. 330, Struik S. 184).

G) Didaktische Herausforderung

Die Determiniertheit der Funktionen durch ihr Verhalten in einer beliebig kleinen Umgebung eines einzigen Punktes ist, wie gesagt, nicht gegeben für die Funktionen, von welchen die reelle Analysis handelt. Weitverbreitete alltagsphilosophische Vorstellungen von der Mathematisierung unserer Welt bestimmenden Gesetze stehen diesem Befund entgegen. Mathemati-

sierung ist für viele fast identisch mit Vorausberechenbarkeit. In vielen Einführungen wird der Schock, den die Mathematiker Ende des 19. Jahrhunderts erlebten, dadurch verschleiert, daß in den Beispielen nur immer wieder von den speziellen Funktionen die Rede ist, welche die für die komplexe Funktionentheorie typischen Eigenschaften besitzen. (Das globale Verhalten ist durch das lokale Verhalten determiniert.) Die Stetigkeit (oder die stetige Differenzierbarkeit) erscheint in solchen „elementaranschaulichen Einführungen“ als eine unter vielen möglichen Eigenschaften, die man bei einer konkret vorgegebenen Funktion konstatieren und ausschalten kann. Sie tritt nicht als die theoretische Grundlage der weiteren Konstruktionen hervor.

Weierstraß war der erste, der die Unterschiede zwischen reeller und komplexer Analysis klar erkannt hat. Mit seinem kompromißlosen Festhalten an strenger Beweisführung hat er die allgegenwärtigen intuitiven Vorstellungen von Stetigkeit demontiert und die weiterführenden Begriffe neu sortiert. Vor demselben Problem steht jeder Student: die tragfähigen intuitiven Vorstellungen müssen von den irreführenden geschieden werden. Das didaktische Problem besteht darin, die „Beispiele“ so zu präsentieren, daß sie nicht als Beispiel für untypisches Verhalten rezipiert werden. Daß dies nicht leicht ist, lehrt die Geschichte. Bekanntlich stießen die Konstruktionen von Weierstraß auf Widerstand bei seinen Zeitgenossen. Diese hatten (ähnlich wie die heutigen tüchtigen Gymnasiasten) festverwurzelte Vorstellungen von Stetigkeit und die wollten sie nicht aufgeben. Die Reaktion auf das Weierstraßsche Beispiel einer stetigen nirgends differenzierbaren Funktion (1875 publiziert) reichte vom Erstaunen über Unglauben bis hin zu Abscheu. Viele Mathematiker weigerten sich, derartige Funktionen erst zu nehmen und nannten sie „pathologische“ Funktionen. Noch zwanzig Jahre später klagte HERMITE in einem Brief an TH.J. STIELTJES vom 20.5.1893: „Aber diese so eleganten Entwicklungen sind mit einem Fluch belegt. Die Analysis nimmt mit der einen Hand zurück, was die andere gibt. Ich wende mich mit Abscheu und Schrecken von dieser beklagenswerten Wunde der stetigen nirgendsdifferenzierbaren Funktionen ab.“ (zitiert nach Bölling 1994, Mitteilungen der DMV).

Wie sollte die reelle Analysis den heutigen Studenten leicht fallen? Was muß die Vorlesung tun, um intuitive Vorstellungen von den Gegenständen der reellen Analysis zu entwickeln, nach welchen das als das normale Verhalten erscheint, was sich aus den Deduktionen ergibt? Mit mathematischer Strenge ist allenfalls eine gewisse Disziplinierung zu erreichen, die für fruchtbares Mathematiktreiben nötige Intuition ist damit nicht aufzubauen. Man muß sich daher nach unserer Meinung bewußt von den „Beispielen“ aus der Zeit von Euler und den Bernoullis lösen; man muß den Blick auf die wesentlich abstrakteren Anwendungsintentionen der neuen Theorie lenken.

Hier steht die Didaktik der Analysis vor ungelösten Problemen. Daß die Theorie der Fourierreihen in den Kanon der Anfängervorlesung aufgenommen worden ist, könnte als ein Anfang angesehen werden. Auf der anderen Seite entspricht die Art und Weise, wie etwa die Integrationstheorie weithin behandelt wird, in keiner Weise diesem Anliegen. Vor lauter Bemühen, leicht verständlich zu sein, verpassen manche Einführungen das Ziel das Wichtige verständlich zu machen; die Intuition bleibt dem Geist des 19. Jahrhunderts verhaftet und der Zwang zum strengen Beweisen läuft der Intuition zuwider.

B.9 Von Cauchys Begriff der stetigen Funktion bis zum sog. allgemeinen Dirichletschen Funktionsbegriff; Zur Rezeption in der Lehre

- 1) Nach landläufiger Meinung müssen die zu untersuchenden Objekte zunächst einmal **mit irgendwelchen Mitteln vorgegeben** sein. Im 18. Jahrhundert wurde das auch nicht in Frage gestellt. Euler schreibt in seiner richtungsweisenden „Introductio ad analysin infinitorum“ (1748): „*Eine Funktion einer veränderlichen Zahlgröße ist ein analytischer Ausdruck, der auf irgendeine Weise aus der veränderlichen Zahlgröße und aus eigentlichen Zahlen oder aus konstanten Zahlengrößen zusammengesetzt ist.*“ Analytische Ausdrücke sind bei Euler Gebilde, die durch Anwendung der gängigen mathematischen Operation einschließlich der Grenzprozesse entstehen. Die Art, wie Euler mit Grenzprozessen hantierte, war aber mit Unsicherheiten behaftet.
- 2) Wir haben hier einige **Beispiele für Grenzprozesse mit Funktionen** gesehen. Wir haben durch Rekursion mit Hilfe der allereinfachsten Rechnungsarten Funktionen konstruiert, die monoton gegen die Wurzelfunktion $\sqrt{x}$ konvergieren. Die Funktion $|x|^\alpha$ kann man aber natürlich auch direkt als einen analytischen Ausdruck ansehen, so daß sich eine Konstruktion über Grenzprozesse erübrigt.

Bei unserer Konstruktion von e^x und $\ln x$ hatten wir schließlich monotone Konvergenz auf jedem beschränkten Bereich. Bei der Konstruktion von $e^z = e^x \cdot (\cos y + i \sin y)$ konnten wir nicht mit monotoner Konvergenz arbeiten; wir hatten hier gleichmäßige Konvergenz auf jedem Kompaktum $\{z : |z| \leq M\}$.

Hinweis Ein Mathematiker ist selten zufrieden, wenn er von einer Funktionenfolge nur weiß, daß sie punktweise konvergiert. In unserer Analysis I wird der Begriff der punktwisen Konvergenz nicht gebraucht. Zentrale Bedeutung hat der Begriff der (lokal) gleichmäßigen Konvergenz von Funktionenfolgen. In der Maß- und Integrationstheorie wird die „Konvergenz fast überall“ Bedeutung erlangen.

- 3) Man kann es sich zur Aufgabe machen, einzelne Funktionen, die mit irgendwelchen Mitteln gegeben sind, zu analysieren. Man kann z.B. fragen, ob eine gegebene Funktion stetig ist. Cauchy war unzufrieden mit dem **Stetigkeitsbegriff des 18. Jahrhunderts**. Cauchy schreibt (in den Gesammelten Abhandlungen VIII S. 145/6): „In den Werken von Euler und Lagrange heißt eine Funktion stetig oder unstetig, je nachdem die Funktionswerte zu den unterschiedlichen Werten der Variable durch dieselbe oder durch verschiedene Gleichungen geliefert werden ... Dieser Definition fehlt nicht nur die mathematische Präzision. Es kann auch passieren, daß die angegebenen algebraischen oder transzendenten Formeln, die eine Funktion repräsentieren, für manche Werte der Variablen dasselbe und für andere etwas verschiedenes liefern.“ (freie Übersetzung) Als

Beispiel gibt Cauchy an

$$\frac{2}{\pi} \int_0^{\infty} \frac{x^2}{t^2 + x^2} dt = \begin{cases} x & \text{wenn } x \geq 0 \\ -x & \text{wenn } x \leq 0. \end{cases}$$

In Eulers Theorie wäre, wie Cauchy ausführt, die linke Seite stetig zu nennen, die rechte aber unstetig. Diese Unbestimmtheit hat Cauchy durch seine Definition der Stetigkeit beseitigt.

Euler hatte unabhängig von allen analytischen Ausdrücken noch eine weitere Idee von einer stetigen Funktion. In den Institutiones Calculi Integralis III § 301 (1768–1770) versteht er unter einer Funktion eine Kurve, die man mit der freien Hand zeichnen kann („curva quaecunque libero manus ductu descripta“). Anleihen bei **geometrischen Vorstellungen** wurden von Bolzano, Cauchy und anderen in Frage gestellt.

- 4) Der **moderne Stetigkeitsbegriff** ist wohl zum ersten Male von B. Bolzano klar ausgesprochen worden. (Ob Cauchy ihn bei Bolzano abgeschrieben oder selbständig entwickelt hat, ist unter Historikern umstritten.)

„Nach einer richtigen Erklärung nämlich versteht man unter der Redensart, daß eine Funktion $f(x)$ für alle Werte von x , die innerhalb oder außerhalb gewisser Grenzen liegen, nach dem Gesetze der Stetigkeit sich ändere, nur so viel, daß, wenn x irgendein solcher Werth ist, der Unterschied $f(x + \omega) - f(x)$ kleiner als jede gegebene Größe gemacht werden könne, wenn man ω so klein, als man nur immer will, annehmen kann.“ (Ostwalds Klassiker Nr. 153, 1905 S. 8/9)

Bolzano möchte keine bloßen Gewißmachungen, er will Begründungen, d.h. Darstellungen jenes objektiven Grundes, den die zu beweisende Wahrheit hat. 1817 publizierte er seinen Aufsatz „Rein analytischer Beweis des Lehrsatzes, daß zwischen je zwei Werthen, die ein entgegengesetztes Resultat gewähren, wenigstens eine reelle Wurzel der Gleichung liege.“

Bolzano beweist den Zwischenwert recht ähnlich, wie es heute in den Lehrbüchern üblich ist. Bolzanos Beweis hat nur insofern eine gewisse Schwäche als er das Problem, das schließlich durch den Begriff des die Dedekindschen Schnitts geklärt wurde, etwas kursorisch behandelt. Bei Bolzano heißt es (S. 14):

„... woraus sich für jeden, der einen richtigen Begriff von Größe hat, ergibt, daß der Gedanke eines i , welches das größte derjenigen ist, von denen gesagt werden kann, daß alle unter ihm stehende die Eigenschaft besitzen, der Gedanke einer reellen, d.h. wirklichen Größe sei.“

Die meisten Lehrbuchautoren tun u.E. Bolzano insofern unrecht, als sie den Anschein erwecken, Bolzano hätte überhaupt nicht darüber nachgedacht, was unter einer **reellen Zahl** zu verstehen sei. Man sollte aber im angegebenen Zitat wohl eher einen Vorläufer von Dedekinds Axiom sehen.

- 5) **Zwischenwertsatz** Sei $f(\cdot)$ stetig auf $[a, b]$ mit $f(a) < 0 < f(b)$. Dann existiert mindestens ein x^* mit $f(x^*) = 0$.

Beweis Betrachte die Menge M_u aller derjenigen x für welche $f(y) < 0$ für alle $y < x$ gilt.

Dieses M_u bestimmt einen Dedekindschen Schnitt, also ein x^* . Wäre $f(x^*) < 0$, dann wäre wegen der Stetigkeit $f(x) < 0$ auch noch in einer genügend kleinen Umgebung, was ein Widerspruch zur Konstruktion von x^* ist. Ebenso führt die Annahme $f(x^*) > 0$ zu einem Widerspruch.

- 6) Bolzanos Idee vom Beweisen markiert den Unterschied zwischen der Analysis des 19. Jahrhunderts gegenüber der Analysis des 18. Jahrhunderts. Die alte Analysis stützt sich auf **Eigenschaften der Formeln**, durch welche eine Funktion vorgegeben ist, um Eigenschaften dieser Funktion herzuleiten. Demgegenüber bestimmt die neue Analysis die Eigenschaften der Funktionen aus der **impliziten Charakterisierung der Klasse von Funktionen**, die zur Debatte steht. Bei Cauchy (im Gegensatz zu Euler) sind die Funktionen keine vorgegebenen mathematischen Objekte, die es zu behandeln gilt; stetige Funktionen erscheinen vielmehr als diejenigen mathematischen Objekte, auf welche der entwickelte Kalkül paßt. Man hat die Herangehensweise von Bolzano, Cauchy u.a. eine **begriffliche Herangehensweise** genannt. Die stetige Funktion wird zu einem Begriff, d.h. zu einer gedanklichen Entität, die vielfältige Darstellungen und Darstellungsmöglichkeiten mit vielfältigen Problemen, Anwendungen und Bedeutungen verbindet.

Es ist unter Historikern umstritten, inwieweit sich Cauchy der Wendung zu einem amodalen Denken bewußt war, zu einem Denken also, das nicht mehr an feste mögliche Repräsentationen gebunden ist. Jedenfalls öffnete Cauchys „Cours d’analyse“ epistemologisch den Weg zum Begriff der „willkürlichen Funktion“. Diesen Weg beschritten aber zunächst nur wenige Mathematiker. DIRICHLET und RIEMANN, die Nachfolger von Gauß in Göttingen, waren die erfolgreichsten. L. KRONECKER (1823–1891) in Berlin war der einflußreichste Gegner der amodalen Herangehensweise.

- 7) Als der Schöpfer des modernen Funktionsbegriffs gilt P.L. DIRICHLET (1805–1859), der Mann, der 1855 Nachfolger von Gauß in Göttingen wurde. Felix Klein zitiert ihn so: „*Ist in einem Intervalle jedem einzelnen Werte x durch irgendwelche Mittel ein bestimmter Wert y zugeordnet, dann soll y eine Funktion von x heißen.*“ Bei Lobatschewskij (1793–1856) heißt es:

“Der allgemeine Begriff erfordert, daß eine Funktion von x eine Zahl genannt wird, die für jedes x gegeben ist und sich fortschreitend mit x ändert. Der Wert der Funktion kann gegeben sein entweder durch einen analytischen Ausdruck oder durch eine Bedingung, welche ein Mittel darbietet, alle Zahlen

zu prüfen und eine davon auszuwählen oder schließlich kann die Abhängigkeit bestehen aber unbekannt bleiben“

(zitiert nach A.P. Juschewitsch, The concept of function up to the middle of the 19th century, Arch. Hist. Exact Sci. 16, S. 77 (1976)).

- 8) Der allgemeine „Dirichletsche“ Begriff einer willkürlichen Funktion oder Zuordnung hat in seiner definitorischen Beschreibung während des ganzen 19. Jahrhunderts Unbehagen ausgelöst. So schreibt z.B. HANKEL 1870: *„Diese reine Nominaldefinition, der ich im folgenden den Namen Dirichlets beilegen werde, ... reicht nun aber für die Bedürfnisse der Analysis nicht aus, da Funktionen dieser Art allgemeine Eigenschaften nicht besitzen und damit alle Beziehungen von Funktionswerten für verschiedene Werte des Arguments in Wegfall kommen“*. Hankel sah auf der anderen Seite, daß man den Eulerschen Funktionsbegriff (wenn man die Theorie der Grenzprozesse ernsthaft betreibt) nicht dahin entwickeln kann, daß man einen organisch abgeschlossenen Kreis von Funktionen erhält, die gemeinsame Eigenschaften haben, wie etwa: Stetigkeit, Differenzierbarkeit, Integrierbarkeit, Bestimmtheit und Entwicklungsfähigkeit nach dem Taylorschen Lehrsatz — bei allen diesen Eigenschaften einzelne singuläre Punkte ausgenommen. Er leitete daraus die Befürchtung ab, daß eine allgemeine Theorie der Funktionen zu einer leeren Tautologie entarten würde; es würde schlechthin unmöglich sein, „aus gegebenen Relationen neue abzuleiten, indem man dabei die allgemeinen Eigenschaften des Begriffs, dem alle jene Objekte unterstehen, zur Anwendung bringt“. Als Ausweg empfahl Hankel das *„Prinzip der Permanenz formaler Gesetze“*, welches er bereits in seiner Theorie der komplexen Zahlen als das Prinzip jedes systematischen Fortschritts nachgewiesen hatte. (In den beiliegenden Auszügen (insbesondere S. 87,89) kommt Hankel ausführlich zu Wort). — Hankels Ansatz hat sich als unpassend für die moderne reelle Analysis erwiesen, hat aber bis heute starken Einfluß auf die Lehre der elementaren Analysis behalten. Cantors Mengenbegriff und Dirichlets Funktionsbegriff werden in der elementaren Analysis meistens nur deklamatorisch benützt; erst in den weiterführenden Veranstaltungen (Punktmengentopologie, Maßtheorie, Funktionalanalysis) machen die Dozenten ernst mit diesen abstrakteren Begriffen.
- 9) Durch die Dirichletsche Definition wird auf keinen anderweitig festgelegten Gegenstand verwiesen, was der dem Wiener Kreis zuzurechnende Topologe und Logiker K. Menger einmal durch die Bemerkung ausgedrückt hat, es handle sich um „die Definition der Definition von Funktion anstatt um die Definition von Funktion“. In der Tat gibt Dirichlets Funktionsbegriff lediglich eine Antwort auf die Frage, was man haben muß, um sagen zu können, daß man eine wohldefinierte Funktion vor sich hat.
- 10) Der heutige Student wird mit dem allgemeinen Begriff der Funktion als einer **beliebigen Zuordnung** überfallen. Die stetigen Funktionen oder auch die borelmeßbaren Funktionen werden ihm dann als Zuordnungen spezieller Art herausgehoben. Dieser Zugang stellt die historische Entwicklung auf den Kopf. Der Begriff der willkürlichen Funktion

hat sich erst langsam im 20. Jahrhundert im Zusammenhang mit Cantors Mengenlehre durchgesetzt. (Im Grundlagenstreit der 20er Jahre kam er nochmals sehr stark unter Druck.) Viele wichtige Mathematiker waren nicht bereit, ihre Wertschätzung der mathematische Formel hinter den abstrakten Begriff zurückzustellen. So schrieb z.B. L. KRONECKER 1884 in einem Brief an G. Cantor: *„Einen wahren wissenschaftlichen Wert erkenne ich — auf dem Gebiete der Mathematik nur in den konkreten mathematischen Wahrheiten, oder schärfer ausgedrückt, nur in mathematischen Formeln. Diese allein sind, wie die Geschichte der Mathematik zeigt, das Unvergängliche.“*

- 11) Die Anerkennung, die Cauchys Theorie der stetigen Funktionen fand, hatte nichts mit der (gar nicht explizit gemachten) amodalen Herangehensweise zu tun; sie hatte vielmehr überwiegend konventionelle Gründe. Der erkennbare Vorzug von Cauchys Theorie gegenüber den Auffassungen seiner Vorgänger bestand darin, daß sie auf vage geometrische Assoziationen verzichtete, Begriffe des unendlich Kleinen und unendlich Großen und den unsicheren Umgang mit divergenten Reihen vermied. Der Mathematiker und Philosoph A.N. WHITEHEAD (1861–1947) schreibt in seiner „Einführung in die Mathematik“: *„Wir haben gesagt, daß eine Funktion stetig sei, wenn sie bei allmählichen Änderungen des Arguments ihren Wert ebenfalls nur allmählich ändert und unstetig, wenn sie ihren Wert sprunghaft ändern kann. Dies ist die Definitionsweise, welche unseren mathematischen Vorfahren genügte, uns neuzeitliche Mathematiker aber nicht mehr befriedigen kann. ... Ein Unterschied zwischen der älteren und neueren Mathematik besteht darin, daß wir unbestimmte Ausdrücke wie ‚allmählich‘ aus den Erklärungen tilgen. ... Unsere Aufgabe ist es also, den Begriff der Stetigkeit zu definieren, ohne von den Ausdrücken ‚kleine‘ oder ‚allmähliche‘ Änderung des Funktionswerts Gebrauch zu machen.“* (In Auflage von 1958, S. 89/90).

Dies ist nun freilich eine recht bescheidene Aufgabe. Aber in der Tat versäumen es bis heute viele Einführungen in die Analysis, den Studenten den tiefergreifenden Bruch mit der Analysis des 18. Jahrhunderts vor Augen zu führen, den Cauchy, Dirichlet, Riemann und andere herbeigeführt haben.

Auszug aus dem Buch „Ostwalds Klassiker“, 153.

Untersuchungen über die unendlich oft oszillierenden und unstetigen Funktionen.

Von

Hermann Hankel †.

(Abdruck aus dem Gratulationsprogramm der Tübinger Universität
vom 6. März 1870¹⁾). (Mathematische Annalen. XX. 1882, S.63–112.)

Einleitende Bemerkungen über den Funktionsbegriff.

Nicht allein entschiedene Vorliebe für die Synthese und gänzliche Verleugnung allgemeiner Methoden charakterisiert die antike Mathematik gegenüber unserer neueren Wissenschaft; es gibt neben diesem mehr äußeren, formalen, noch einen tiefliegenden realen Gegensatz, welcher aus der verschiedenen Stellung entspringt, in welche sich beide zu der wissenschaftlichen Verwendung des Begriffes der *Veränderlichkeit* gesetzt haben.

Denn während die Alten den Begriff der Bewegung, des räumlichen Ausdruckes der Veränderlichkeit, aus Bedenken, die aus der philosophischen Schule der Eleaten auf sie übergegangen waren, in ihrem strengen Systeme niemals und auch in der Behandlung phoronomisch erzeugter Kurven nur vorübergehend verwenden, so datiert die neuere Mathematik von dem Augenblicke als DESCARTES von der rein algebraischen Behandlung der Gleichungen dazu fortschritt, die Größenveränderungen zu untersuchen, welche ein algebraischer Ausdruck erleidet, indem eine in ihm allgemein bezeichnete Größe eine stetige Folge von Werten durchläuft.

Die Abhängigkeit, in welcher hiernach die Wertreihen zweier veränderlicher Größen stehen, findet ihren anschaulichen Ausdruck in dem Verhältnis der Ordinate zur Abszisse einer krummen Linie, und bedurfte so lange keines besonderen Terminus, als man sich auf die Untersuchung spezieller Fälle beschränkte.

Die Entdeckung des Infinitesimalkalküls aber mit seinen umfassenden Methoden drängte zu einer allgemeinen Bezeichnung dieses Abhängigkeitsverhältnisses, das nach LEIBNITZens²⁾ Vorgange JOH. BERNOULLI³⁾ im Jahre 1718 als »Funktion« bezeichnete.

¹⁾Indem wir die letzte Arbeit von HERMANN HANKEL, welche bisher durch die Art ihrer Veröffentlichung nur schwer zugänglich war, erneut zum Abdruck bringen, rechnen wir auf den Dank des mathematischen Publikums; wird doch auf dieselbe bei fast allen neueren Untersuchungen über den Funktionsbegriff zurückgegangen! Dabei wird man es nur billigen, daß wir den Abdruck durchaus wörtlich gestalteten und also die Inkorrektheiten, welche die Arbeit im einzelnen enthalten mag, als nebensächlich betrachteten.

Die Red. d. math. Annalen.

²⁾Acta Erudit. 1692. April p. 170, und 1694. Juli p. 316

³⁾Opera omnia. t. II. p. 241. Vgl. auch LEIBNITII et JOH. BERNOULLII commerc. epist. t. I. p. 386, und JAC. BERNOULLII, Oper. omn. t. II. p.788.

Der Begriff der Funktion ist von da an der fundamentale Ausgangspunkt für die Analysis geworden, deren epochemachenden Fortschritte aufs engste mit der Ausbildung zusammenhängen, welche er allmählich erlitten hat.

EULER, der das wissenschaftliche Bewußtsein in der Mitte des vorigen Jahrhunderts am vollständigsten vertritt, definiert in wesentlicher Übereinstimmung mit seinen Vorgängern⁴⁾:

»Functio quantitatis variabilis est expressio analytica quomodocunque composita ex illa quantitate variabili et numeris seu quantitativis constantibus.«

Wir haben uns unter einem solchen »analytischen Ausdruck«, dessen Bedeutung EULER näher zu bestimmen nicht für nötig hält, nichts anderes zu denken, als eine nach dem Typus algebraischer Funktionen gebildete Größenabhängigkeit. Denn die transzendenten Funktionen dachte man sich durch Entwicklungen bestimmt, welche aus einer Anzahl von Elementaroperationen (Addition, Subtraktion, Multiplikation und Division) ebenso gebildet werden, als die algebraischen durch eine endliche Anzahl solcher, ohne daß man sich über die Zulässigkeit einer solchen grenzenlosen Fortsetzung der Operationen, Bedenken gemacht hätte. Zu jenen vier Spezies traten dann noch die Differentiation und Integration als die der Analysis eigenen Operationen hinzu.

Indem man sich immer an jenes Vorbild der algebraischen Funktionen anlehnte, so übertrug man alle allgemeinen Eigenschaften, die man an den algebraischen Funktionen bemerkte, ihre besondere Stetigkeit, ihre Singularitäten, ihre besondere Weise, unstetig oder unendlich zu werden, ohne weiteres auf alle Funktionen. Da man bei den algebraischen Funktionen die Reihenentwicklung nach Potenzen eines Inkrementes, den Differentialquotienten und das Integral direkt ableiten konnte, so hielt man sich nicht nur für berechtigt, die Existenz einer solchen Reihe, des Differentialquotienten und Integrales ganz allgemein für alle Funktionen anzunehmen, sondern man kam überhaupt nicht auf den Gedanken, daß hierin eine Behauptung, sei es nun als Axiom oder Theorem liegt, — so selbstverständlich erschien diese, von der geometrischen Anschauung der die Funktionen repräsentierenden Kurven scheinbar unterstützte Übertragung der Eigenschaften algebraischer Funktionen auf die Transzendenten; und Beispiele, in denen recht eigentlich analytische Funktionen⁵⁾ Singularitäten zeigten, die von denen der algebraischen Funktionen wesentlich verschieden waren, blieben gänzlich unbeachtet.

Wenn beliebig gezeichnete Kurven in irgend einer Beziehung von dem Verhalten algebraischer Kurven abwichen, so nannte man sie⁶⁾ »curvae discontinuae seu mixtae seu irregulares« im Gegensatz zu den in einer Gleichung enthaltenen »curvae continuae«, und war der Überzeugung, daß jene auf keine Weise durch eine analytische Gleichung darstellbar seien. Indem man jedoch die Abhängigkeit der Ordinate solcher Kurven von der Abszisse als willkürliche Funktion in dem Problem der Saitenschwingungen einführte, ging man bereits über die eigentliche Definition des Funktionsbegriffes hinaus, und D’ALEMBERT hatte recht, wenn er behauptete, dies sei »contre les règles d’analyse«⁷⁾.

⁴⁾ Introd. in an. inf. 1748. t. I. p. 4.

⁵⁾ Wie $\sin \frac{1}{x}$, $e^{\frac{1}{x}}$, $\frac{1}{1+e^{\frac{1}{x}}}$, $\int \frac{dx}{1+e^{\frac{1}{x}}}$ u.a. für $x = 0$. Zum Beweise des im Texte Gesagten, diese z.B. D’ALEMBERT Hist. de l’Acad. Berlin für 1747, p. 236.

⁶⁾ EULER, Introd. in an. inf. t. II, p. 6.

⁷⁾ Opusc. math. t. I. p. 32. Die ganze Polemik, die sich über das Problem de chordis vibrantibus zwischen D’ALEMBERT, EULER, DAN. BERNOULLI und LAGRANGE erhob und uns vollständigen Einblick in die Anschau-

In noch viel stärkerem Maße aber geschah dies, indem man sich gedrängt sah, die für ein gewisses Intervall der Veränderlichen wohl definierten Funktionen auf ein weiteres Gebiet auszudehnen, und Funktionen für alle reellen Werte des Argumentes zu bestimmen, wenn auch ihre ursprüngliche Definition auf ein bestimmtes Gebiet beschränkt war. Das erste Beispiel einer solchen Fortsetzung war die Bestimmung der Logarithmen negativer Zahlen. Damit aber hatte man den Boden gänzlich verlassen, auf den man sich durch die Definition der funktionellen Abhängigkeit gestellt hatte, und ein neues Prinzip geschaffen, das ich etwa folgendermaßen aussprechen zu können glaube:

Wenn für ein Gebiet der Veränderlichen eine Entwicklung gegeben ist, welche nicht unmittelbar über die Grenzen desselben fortgesetzt werden kann, weil sie außerhalb ihre Bedeutung verliert, wenn ferner für ein anderes Gebiet der Veränderlichen eine andere Entwicklung gegeben ist, welche in das erste Gebiet nicht fortgesetzt werden kann, und es keine für beide Gebiete gültige Entwicklung gibt, so sollen beide Entwicklungen als zu einer Funktion gehörig angesehen werden, wenn die Funktion in jedem der beiden Gebiete dieselben *Eigenschaften* zeigt.

Daß man hiermit einen wesentlich neuen Funktionsbegriff gesetzt hatte, konnte um so weniger bemerkt werden, als man niemals dieses Prinzip als solches erkannte, wenigstens nicht ausdrücklich formulierte, und so sah man sich denn auch nicht veranlaßt, genauer zu untersuchen, welche Eigenschaften und Relationen zwischen den Funktionswerten dazu erforderlich seien, um die Einheit der funktionellen Abhängigkeit in zwei verschiedenen analytischen Entwicklungen für verschiedene Größengebiete zu verbürgen.

Eine ganze Reihe von unzureichend oder fehlerhaft erwiesenen, nicht genügend determinierten, sowie von durchaus falschen Sätzen⁸⁾ legen genügendes Zeugnis davon ab, wie unsicher das Fundament war, auf das man die ganze Funktionentheorie gebaut hatte.

Diese ganze Auffassung des Funktionsbegriffes, die ich kurz als die EULERSche bezeichnen werde, erhielt den ersten schweren Stoß im Jahre 1807 durch FOURIERS bedeutende Entdeckung⁹⁾, daß es möglich ist, durch periodische Reihen nicht nur *analytische*, etwa daneben in Potenzreihen entwickelbare Funktionen (funct. continuæ), sondern auch ganz beliebige, keinem einfachen Gesetz genügende,

ungen gewährt, die man damals mit dem Funktionsbegriffe verband, ist in meisterhafter Weise von RIEMANN dargestellt worden (Üb. d. Darstellbarkeit einer Funktion d. eine trigon. Reihe. 1867. Abh. d. Gött. Ges. t. 13. Art. 1).

⁸⁾In dieser Beziehung erinnere ich nur an LANGRANGES Beweis des TAYLORSchen Lehrsatzes, wie er ihn noch 1813 (théor. d. fonct. 2 éd.) gab, und daß CAUCHY (Leç. sur le calc. diff. p. 105) 1829 »entdeckte«, wie dieser Satz zuweilen »en défaut« ist; ferner daran, daß GAUSS (Theor. d. res. funct. dem. tertia, art. 4) erst 1816 in bestimmter Weise auf die Unzulässigkeit einer Integration über unendliche Werte der Funktion hinwies¹⁾; daß erst 1826 ABEL (Crelle, Journ. t. I. p. 311) eine richtige Bestimmung der Potenz und der binomischen Reihe lieferte, nachdem POISSON auf verschiedene Paradoxien in dieser Beziehung aufmerksam gemacht, und sich Männer, wie POINSOT (Rech. sur l'analyse d. sections angul. 1825) und viele andere vergeblich mit deren Auflösung beschäftigt hatten; ich erinnere ferner an die vielfachen ungenügenden und fehlerhaften Versuche, die allgemeine Begriffsbestimmung der numerischen Fakultäten für gebrochene und negative Werte der Veränderlichen zu geben, die erst WEIERSTRASS (Crelle, Journ. t. 51 p. 1) 1856 durch deren richtige Definition ersetzte usw. usw.

⁹⁾Bull. d. scienc. p. l. soc. philomatique. t. I. p. 112.

oder verschiedene Gesetze in ihren verschiedenen Teilen folgende Funktionen (funct. discontinuae), die ich *illegitime* nennen werde, darzustellen.

Damit aber war, wie man, freilich zögernd anerkannte, der ganze ältere Begriff unmöglich geworden: Wenn jene functiones discontinuae auch durch analytische Ausdrücke darstellbar waren, so konnten die Eigenschaften algebraischer Funktionen nicht mehr als typisch auf alle Transzendenten übertragen werden; und wenn eine und dieselbe Reihe für verschiedene aneinanderstoßende Gebiete der Veränderlichen verschiedene analytische Gesetze repräsentieren konnte, so fiel das oben formulierte Prinzip der Fortsetzung in sich selbst zusammen.

Es blieb nur übrig, das Axiom, daß jene Eigenschaften der algebraischen Funktionen, die sich auf ihre Stetigkeit, ihre Entwickelbarkeit in Potenzreihen usw. beziehen, allen analytischen Funktionen ebenfalls zukommen, und daher auch die Forderung, eine Funktion solle analytisch darstellbar sein, als eine bedeutungslose fallen zu lassen; und indem man so den Knoten zerhieb, folgende Erklärung zu geben:

Eine Funktion heißt y von x , wenn jedem Werte der veränderlichen Größe x innerhalb eines gewissen Intervalles ein bestimmter Wert von y entspricht; gleichviel, ob y in dem ganzen Intervalle nach demselben Gesetze von x abhängt oder nicht; ob die Abhängigkeit durch mathematische Operationen ausgedrückt werden kann oder nicht.

Diese reine Nominaldefinition, der ich im folgenden den Namen DIRICHLETS beilegen werde, weil sie seinen Arbeiten über die FOURIERSchen Reihen¹⁰⁾, welche die Unhaltbarkeit jenes älteren Begriffes zweifellos dargetan haben, zugrunde liegt, reicht nun aber für die Bedürfnisse der Analysis nicht aus, da Funktionen dieser Art allgemeine Eigenschaften nicht besitzen, und damit alle Beziehungen von Funktionswerten für verschiedene Werte des Argumentes in Wegfall kommen.

Es ist so eine empfindliche Lücke in den analytischen Fundamentalbegriffen entstanden, die, obgleich sie überall mit Stillschweigen übergangen wird, doch nicht minder vorhanden ist; wie ein Blick selbst auf die besseren Lehrbücher der Analysis lehrt. Das eine definiert die Funktionen wesentlich im EULERSchen Sinne, das zweite verlangt, y solle sich »gesetzmäßig« mit x ändern, ohne daß eine Erklärung dieses dunklen Begriffes gegeben ist, das dritte definiert sie in der Weise DIRICHLETS, das vierte gar nicht; alle aber leiten aus ihrem Begriffe Folgerungen ab, die nicht in ihm enthalten sind.

Es war nötig, über den DIRICHLETSchen Begriff hinauszugehen, wie bereits CAUCHY dies seit dem Jahre 1815 zu tun begann². Aber erst in neuester Zeit (1851) ist durch RIEMANN¹¹⁾ mit echt philosophischem Geiste ein fester Grund neu gelegt worden, indem er, von dem DIRICHLETSchen Begriffe ausgehend, den der (monogenen) Funktion einer *komplexen Variablen* begründete, und so jener leeren Definition wieder einen Inhalt gab, der sich dem des älteren Begriffes annäherte.

Es war dem Gründer der Funktionentheorie komplexer Veränderlicher leider nicht vergönnt, sein System allseitig und nach dem einheitlichen Plane aufzubauen, dessen Großartigkeit wir aus den schönen uns bekannten Bruchstücken noch erschließen können.

Die Grundpfeiler des neuen Systems werden hier und da noch nicht für sicher und fest genug gehalten, um von ihnen aus das ganze Gebäude der Analysis aufzuführen, und namentlich das von RIEMANN

¹⁰⁾ Repert. d. Physik, her. v. DOVE t. I. 1837. p. 152. (S. Klass. 116, 3–34.)

¹¹⁾ Grundlagen f. e. allgem. Theor. d. Funkt. Göttingen. Ostwalds Klassiker. 153

in edler Bescheidenheit nach DIRICHLET benannte Prinzip ist neuerdings mehrfach angegriffen und verteidigt worden.

Die Bedenken sind hergenommen von gewissen a priori denkbaren Ausnahmefällen, in denen die Funktionen Unstetigkeiten zeigen, die nicht ohne weiteres ausgeschlossen werden können und doch die Schlüsse, die man auf alle der Definition genügenden Funktionen anzuwenden hoffen durfte, alterieren³.

Ich glaube nun, daß der einzige Weg, um über diese Unstetigkeiten ins Klare zu kommen und damit den Entscheid über die Natur der Funktionen vorzubereiten, der sei, sich von allen Vorstellungen, wie sie auch dem modernsten Mathematiker noch aus dem EULERSchen Funktionsbegriffe anhaften, loszusagen und zunächst einmal die Mannigfaltigkeit der in dem *reinen* DIRICHLETSchen Funktionsbegriffe enthaltenen, möglichen Größenbeziehungen zweier Veränderlichen auseinanderzulegen, dabei aber besondere Aufmerksamkeit den bisher wenig oder garnicht beachteten, den *illegitimen* Funktionen zu schenken.

In der folgenden Abhandlung habe ich den Versuch gemacht, diese Paradoxien der Funktionen zu behandeln, indem ich mich zunächst of *reelle* Variabele und *reelle endliche* Werte der Funktionen einer Veränderlichen beschränke.

Nachdem in § 1 der Sinn, den ich in dieser Abhandlung mit dem Worte »Funktion« verbinden werde, festgelegt, in § 2 die verschiedenen denkbaren Arten der Stetigkeit und Unstetigkeit von Funktionen in Punkten erörtert sind, so gehe ich in § 3 zu der Betrachtung stetiger Funktionen im allgemeinen über und zeige, daß außer den gewöhnlichen analytischen Funktionen mit endlichen, in endlicher Zahl auf einem gewissen Intervalle liegenden Oszillationen, auch Funktionen mit unendlich vielen Oszillationen unendlich kleiner Amplitude denkbar sind. Ausschließlich solche Funktionen dieser Klasse, welche nur in der Nachbarschaft einzelner Punkte jene unendlich vielen Oszillationen zeigen, sind bisher, soviel mir bekannt, durch analytische Ausdrücke dargestellt worden. Mit Hilfe eines Prinzipes, auf welches ich durch ein Beispiel RIEMANNS¹²⁾ aufmerksam wurde⁴, und welches ich das der *Kondensation von Singularitäten* nenne (§ 4), ist es mir gelungen, in § 5 analytische, unbedingt konvergente Reihen aufzustellen, die in ganzen Intervallen durchaus oszillieren.

Während nun alle diese Funktionen stets der Integration unterzogen werden können, so bietet die Frage nach der *Existenz* eines *Differentialquotienten* eigentümliche Schwierigkeiten dar. Es ist diese bisher nur selten einer Erörterung unterzogen worden¹³⁾, weil man die Existenz einer Tangente an jedem Punkte einer Kurve als aus unmittelbarer geometrischer Anschauung gewiß, und für eine selbstverständliche Folge der »lex continuitatis« ansah, die man wie eine *Naturnotwendigkeit* im Gebiete der Mathematik respektierte. Wenn aber auch dies dunkle »Gesetz der Stetigkeit« in der Tat alle Bewegungen in der Natur beherrschen sollte, so darf es doch das Gebiet der reinen Mathematik auf keine Weise beschränken; und jene unmittelbare intuitive Gewißheit hat man selbst in durchaus geometrischen Gebieten als so trügerisch befunden, daß sie auf den Rang eines wissenschaftlichen Beweises nicht

¹²⁾ Üb. d. Darst. e. Funkt. art. 6

¹³⁾ Der, soviel mir bekannt, einzige Versuch AMPÈRES, die Existenz eines solchen a priori für alle Funktionen zu erweisen (Journ. de l'éc. polyt. cah. XIII, 1806. p. 148) ist ganz verunglückt⁵.

mehr Anspruch machen darf. Nach dem Vorgange von GAUSS¹⁴⁾, DIRICHLET, JACOBI¹⁵⁾ u. a. ist denn auch die Überzeugungen, daß die Existenz eines Differentialquotienten stetiger Funktionen *keine notwendige* Folge der Stetigkeit sei, sondern eine besondere Voraussetzung involviere, unter den neueren Mathematikern eine ziemlich allgemeine geworden, obgleich noch gar mancher an stetige Funktionen ohne Differentialquotienten »nicht zu glauben« erklärt. Ich darf hoffen, in § 5 diesen Unglauben als ein Vorurteil dargetan zu haben, indem dort völlig bestimmte Funktionen der angegebenen Art in *analytischem* Ausdrücke durch unbedingt konvergente Reihen dargestellt werden⁶.

In § 6 und 7 habe ich die *linear* unstetigen Funktionen untersucht, d. h. solche, welche auf einem endlichen Intervalle eine unendliche Anzahl von Lösungen der Kontinuität zeigen, und diese nach ihrem inneren Charakter in zwei wesentlich verschiedenen Klassen, die *punktiert*–, und die in Linien *total* unstetigen geteilt, die sich auch (§ 8) insofern wesentlich verschieden verhalten, als die ersteren immer, die letzteren niemals integrabel sind⁷. In § 9 sind beide Klassen unstetiger Funktionen durch analytische Ausdrücke dargestellt.

In den Schlußbetrachtungen habe ich endlich die erhaltenen Resultate zu einer Kritik des Funktionsbegriffes verwendet und zu zeigen versucht, daß diese mit Notwendigkeit auf die RIEMANNsche Definition führt, an welche sich in der III. Note einige Bemerkungen über die lineare Unstetigkeit von Funktionen komplexer Veränderlicher anschließen.

Soviel zur vorläufigen Orientierung über diesen Beitrag zur Metaphysik unserer Wissenschaft, den ich im folgenden dem Urteile der Mathematiker unterbreite. Ich verdanke die Anregung zu diesen Studien wesentlich RIEMANNs Schriften, insbesondere seiner glänzenden Arbeit über die trigonometrischen Reihen, nach deren Erscheinen es keiner Entschuldigung mehr bedarf, sich mit diesen Fragen zu beschäftigen, welche, wie ihr Verfasser in Übereinstimmung mit DIRICHLET bemerkt, »mit den Prinzipien der Infinitesimalrechnung in der engsten Verbindung stehen und dazu dienen können, diese zu größerer Klarheit und Bestimmtheit zu bringen«. Sollte mein Versuch, auf einem so schlüpfrigen, bisher nur selten betretenen Boden festen Fuß zu fassen, teilweise mißlungen sein, und es den Resultaten hier und da an dem wünschenswerten Abschluß fehlen, so glaube ich, auf einige wohlwollende Nachsicht rechnen zu dürfen. Denn es war bei der gelegentlichen Veröffentlichung dieser Gedanken meine einzige Absicht, andere Gelehrte zu veranlassen, ihr Interesse auch diesen fundamentalen Problemen der Wissenschaft zuzuwenden, welche die neuere Funktionenlehre nicht mehr abzuweisen gestattet.

§ 1.

Begriff der Funktion.

Um alle Unbestimmtheit zu beseitigen, erkläre ich hier, daß im folgenden, mit Ausnahme weniger Stellen, an denen dies besonders bemerkt werden wird, nur von *reellen* Werten der Veränderlichen und von *reellen* Funktionen die Rede sein wird, und ein *Unendlichwerden* der Funktion überall ausgeschlossen sein soll. Demgemäß definiere ich:

¹⁴⁾ Allg. Lehrs. in Bez. auf d. im verk. Verh. etc. art. 16

¹⁵⁾ Mündlicher Überlieferung nach soll JACOBI in seinen Vorlesungen zuweilen bemerkt haben, »man könne sich stetige Kurven mit unendlich vielen Spitzen denken.«

Eine *Funktion* von x nennt man $f(x)$, wenn zu jedem Werte von x innerhalb eines gewissen Intervalles ein *einzig* bestimmter Wert von $f(x)$ gehört.

Es ist dabei gleichgültig, woher und wie man $f(x)$ bestimme, ob durch analytische Größenoperationen oder auf andere Weise. Nur muß der Wert von $f(x)$ überall eindeutig und bestimmt sein. So wäre $f(x) = \sin \frac{1}{x}$ in einem $x = 0$ umgebenden Intervalle keine überall bestimmte Funktion, solange sie nur durch ihren analytischen Ausdruck definiert wird. Denn dieser wird in $x = 0$ völlig unbestimmbar und unbestimmt, während er in jedem $x = 0$ noch so nahen Punkte allerdings bestimmt ist. Soll sie obigem Begriffe einer völlig bestimmten Funktion genügen, so muß der Wert $f(0)$ noch besonders, etwa $= 0$, oder $= 1$ oder anders festgestellt sein.

Ebensowenig ist $f(x) = e^{\frac{1}{x}}$ in $x = 0$ bestimmt, da es für positive unendlich zu Null abnehmende x unendlich zu-, für negative x unendlich abnimmt. Wohl aber ist

$$f(x) = \frac{\sin x}{1} - \frac{\sin 2x}{2} + \frac{\sin 3x}{3} - \dots$$

trotz ihrer Sprünge eine überall bestimmte Funktion.

§ 2.

Begriff der Stetigkeit.

Eine Funktion $f(x)$, welche für $x = a$ einen bestimmten endlichen Wert $f(a)$ hat, nennt man in *diesem Punkte stetig*⁸, wenn ε stets so klein angenommen werden kann, daß die Differenz:

$$f(a + \delta) - f(a)$$

für alle δ , die numerisch $< \varepsilon$ sind, numerisch kleiner, als eine beliebig kleine Größe σ ist.

Unstetig heißt sie im Punkte $x = a$, wenn ε *nicht* so klein genommen werden kann, daß

$$f(a + \delta) - f(a)$$

für alle δ , die numerisch $< \varepsilon$ sind, numerisch unter jeder beliebig kleinen Größe σ liege.

...

Ein anderes, weniger einfaches Beispiel einer solchen total unstetigen Funktion ist in der II. Note im Anhang behandelt worden.

§ 10.

Schlußbetrachtungen zur Feststellung des Funktionsbegriffes.

Blicken wir nun auf das durchwanderte Gebiet zurück, und fragen wir, welchen Nutzen unsere Untersuchungen für die Aufklärung des Funktionsbegriffes gebracht haben.

Es war von vornherein gewiß, daß wenn wir den DIRICHLETSchen Begriff der Funktion zugrunde legen, es allgemeine Eigenschaften aller Funktionen, die unter denselben fallen, nicht geben kann.

Aber man ist geneigt, zu vermuten, daß wenn der ältere EULERSche Begriff zugrunde gelegt, und unter einer Funktion nur eine in Rechnungsoperationen gesetzmäßig dargestellte Größenbeziehung verstanden wird, die so definierten Funktionen einen engeren, in sich organisch abgeschlossenen Kreis bilden möchten, dem eine Anzahl gemeinsamer Eigenschaften allerdings zukäme, insbesondere: die Stetigkeit, die Differenzierbarkeit, die Integrabilität, Bestimmtheit und Entwicklungsfähigkeit nach dem TAYLORSchen Lehrsatz — bei allen diesen Eigenschaften einzelne singuläre Punkte angenommen. Diese Vermutung hat sich *nicht* bestätigt. Wir haben gezeigt, daß es durch konvergente Reihen von Größenoperationen möglich ist, Funktionen darzustellen, welche in unendlich vielen Punkten eines endlichen Liniestückes unstetig oder unendlich in der verschiedensten Weise oder unbestimmt¹⁶⁾ werden, welche nirgends integriert und in Potenzreihen entwickelt werden können, und welche selbst, wo sie stetig sind, einen bestimmten Differentialquotienten nicht besitzen. Kurz: Jene beiden Begriffe der Funktionen, so verschieden sie auf den ersten Blick zu sein scheinen, haben sich insofern wesentlich verschieden nicht erwiesen, als wir Funktionen aller möglichen Arten¹⁷⁾ durch analytische Formeln darzustellen imstande waren²¹⁾.

Beide Begriffe sind somit unzureichend, um die Grundlage einer allgemeinen Theorie der Funktionen bilden zu können. Denn die Aufgabe aller Mathematik ist die, aus gegebenen Relationen, die sich auf ein Gebiet von Objekten beziehen, neue abzuleiten, indem man dabei die allgemeinen Eigenschaften des Begriffes, dem alle jene Objekte unterstehen, zur Anwendung bringt. Sind aber solche allgemeinen Eigenschaften nicht vorhanden, so bleibt jede Transformation der gegebenen Relationen eine leere Tautologie.

So, wenn es sich z. B. darum handelte, eine Funktion $f(x)$ aus der Funktionalgleichung:

$$f(x+y) = f(x) + f(y)$$

abzuleiten, so ergibt sich sofort für alle rationalen μ

$$f(\mu) = \mu f(1);$$

ebenso können alle Funktionswerte für Argumente $\mu\xi$, welche rational von einer und derselben irrationalen Größe ξ abhängen, durch den Funktionswert $f(\xi)$ dargestellt werden, indem $f(\mu\xi) =$

¹⁶⁾Die Funktionen, welche in einzelnen Punkten oder gar Linien *unbestimmt* werden, haben wir aus unseren obigen Betrachtungen ausgeschlossen, und es mag daher hier bemerkt werden, daß solche in analytischer Darstellung mittels des Prinzips der Kondensation

$$f(x) = \sum_{n=1}^{\infty} \frac{\varphi(\sin nx\pi)}{n^s}$$

leicht gegeben werden können, wenn $\varphi(y)$ eine für $y = 0$ nicht unbestimmt scheinende, sondern *wirklich* und *notwendig unbestimmte* Größe ist. Setzt man z. B.

$$\varphi(y) = \sin \frac{1}{y}, \quad \text{oder} \quad \varphi(y) = \frac{1}{1 + e^{\frac{1}{y}}},$$

so ist $f(x)$ für alle irrationalen x bestimmt und endlich, für rationale x aber durchaus unbestimmt.

¹⁷⁾Dies ist cum grano salis zu verstehen; denn es ist noch nicht ausgemacht, ob sich einzelne *denkbare* Fälle nicht doch der Analysis entziehen.

$\mu f(\xi)$. Auch stehen die Funktionswerte für rationale und irrationale Werte des Argumentes, und letztere untereinander vielfach im Zusammenhang, durch Gleichungen, wie z.B.

$$f(2 + \sqrt{3}) + f(2 - \sqrt{3}) = f(4)$$

usw. Trotzdem ist, soviel ich sehe, es nicht möglich, direkt anzugeben, wie $f(\xi)$ von seinem Argumente abhängt; denn die verschiedenen Wertreihen für die wesentlich verschiedenen Irrationalitäten, z.B. $f(\mu\sqrt{2})$, $f(\mu\sqrt{5})$, $f(\mu\sqrt[3]{2})$ usw. stehen außer Zusammenhang miteinander; und, wie es auch sei, es liegt jedenfalls nicht in der Absicht eines Analytikers, der auf jene Funktionalgleichung bei seinen Untersuchungen stößt, in diese intrikaten Verhältnisse der Irrationalitäten einzugehen. Man wird vielmehr die Bestimmung $f(\mu) = \mu f(1)$ sofort auch für irrationale μ in Anspruch nehmen, d. h.:

Macht man die Voraussetzung, die Funktion solle überall stetig sein (oder die engere, $f(y)$, solle sich in unendlicher Nähe von $y = 0$ stetig der Null nähern, von der dann die allgemeine Stetigkeit die Folge ist), so ist das Problem mit einem Schlage gelöst, indem die Funktionswerte für irrationale x sich zwischen die $f(x) = x f(1)$ für rationale x stetig einordnen.

Oder eine anderes Beispiel: Es sei eine Funktion für alle Werte von x zwischen 0 und 1 durch die unendliche Reihe

$$f(x) = 1 + \frac{x}{1} + \frac{x^2}{2!} + \frac{x^3}{3!} + \dots$$

gegeben; welche Werte hat $f(x)$ für andere x ? Diese Frage ist schlechterdings nicht zu beantworten, solange wir den DIRICHLETSchen Funktionsbegriff zugrunde legen, weil eben gar kein Zusammenhang zwischen den benachbarten Funktionswerten besteht, mit anderen Worten: weil eine Funktion bisher nicht ein *Gesetz* repräsentiert, sondern ihre Werte ganz willkürlich aufeinanderfolgen.

Ein *Gesetz* ist überall da vorhanden, wo aus einer Reihe von Merkmalen auf das Vorhandensein anderer, nicht gegebener oder beobachteter geschlossen wird. Die Funktion unter unserem bisherigen Begriffe, der rein nominell und ohne realen Inhalt ist, sind daher gesetzlos, *illegitim* zu nennen; sie können aber zu *legitimen* werden, indem sie sich einem Gesetze unterordnen.

Welches soll nun das Gesetz sein, dem diese Funktionen unterworfen werden? Es hält nicht schwer, sich zu überzeugen, daß, soll anders der Funktionsbegriff angemessen sein und den Bedürfnissen des Analytikers entsprechen, wir uns des allgemeinen hodegetischen Prinzipes bedienen müssen, das ich bereits in meiner Theorie der komplexen Zahlen als das jedes systematischen Fortschrittes für jenes Gebiet nachgewiesen und das *Prinzip der Permanenz formaler Gesetze* genannt habe¹⁸⁾.

Ich bin dort von den ganzen Zahlen und den auf sie bezüglichen elementaren Rechnungsoperationen ausgegangen und zeigte, wie man von diesen aus das Zahlengebiet allmählich erweitert, indem man zur Auflösung von Aufgaben, die in dem bekannten Gebiete unrealisierbar sind, neue Zeichen, Zahlen einführt, welche denselben formalen Operationsgesetzen, als die vier Spezies sie enthalten, unterliegen. Nachdem so das Gebiet bis auf das der gemeinen komplexen Zahlen erweitert war, wurde nachgewiesen, daß jede Aufgabe, welche nur eine endliche Anzahl jener Operationen zu ihrer Aufstellung voraussetzt, in dem Gebiete der gemeinen komplexen Zahlen ihre Lösung findet, und so ein organischer Abschluß in diesem Zahlengebiete gefunden wird. Gleichwohl waren noch andere komplexe Zahlensysteme *möglich*,

¹⁸⁾ Vorlesungen über die komplexen Zahlen und ihre Funktionen. Leipzig 1867. Teil I, S. 10–13.

welche teilweise andern formalen Gesetzen folgten, und es war nötig, dies festzustellen, weil nur dadurch die Eigentümlichkeit der gemeinen komplexen Zahlen in helles Licht gesetzt wurde.

In ähnlicher Weise nun, als für das System der Zahlen (oder Größen) die ganzen Zahlen und Operationsgesetze typisch waren und weiterhin für permanent erklärt wurden, so müssen nun für das System der Funktionen die algebraischen Ausdrücke, als die aus den vier Spezies direkt abgeleiteten, die permanenten Gesetze an die Hand geben. Denn wie dort die ganzen Zahlen gewissermaßen das Gerüst bilden, auf das mehr oder weniger direkt alle andern Zahlen gestützt werden, so sind die algebraischen Funktionen, als die, deren Werte *eigentlich* berechnet werden können, die typischen Formen für Funktionen im allgemeinen, die, insofern sie eben Größenbeziehungen darstellen und sich daher *berechnen* lassen sollen, ebenfalls durch die vier Elementaroperationen ausgedrückt werden müssen; nur werden die transzendenten Funktionen durch eine endliche Reihe von Operationen nicht erschöpft, sondern bedürfen eines unendlichen Prozesses.

Die so nach dem Typus algebraischer ganzer und gebrochener Ausdrücke gebildeten transzendenten Funktionen entsprechen nun in gewisser Weise den irrationalen und gemeinen komplexen Zahlen, die nach dem Typus der ganzen und rational gebrochenen Zahlen gebildet sind; und ebenso, wie es unendlich viele Klassen z. B. von irrationalen Größen gibt, so auch unendlich viele Klassen transzendenter Funktionen.

Wie ferner höhere komplexe Zahlen gedacht und konstruiert werden können, die sich den formalen Gesetzen der Arithmetik nicht unterordnen¹⁹⁾, so auch *illegitime* Funktionen, welche sich nicht an den Typus algebraischer Formen anlehnen. Ganz ebenso aber, wie die Betrachtung dieser höheren Zahlensysteme ein systematisches Bedürfnis war, um sich einerseits erfahrungsmäßig von der Möglichkeit derselben, andererseits davon zu überzeugen, daß der Begriff der gemeinen komplexen Zahlen den allgemeinen Begriff der Zahl keineswegs erschöpfe, so ist auch die tatsächliche Nachweisung der Existenz von Funktionen der illegitimsten Arten, wie ich sie in dieser Abhandlung zu geben versucht habe, notwendig gewesen, um unzweifelhaft darzutun, daß die Legitimität der Funktionen uns keineswegs von einer mysteriösen, eisernen Notwendigkeit diktiert wird, die in der »Natur der Sache« liegt, wie man häufig hören kann, sondern daß sie eine konventionelle, aber weise und adäquate Beschränkung ist, die wir uns auferlegen, — und über deren Grenzen wir eben deshalb noch nicht im Reinen sind.

Wenn es sich nämlich nun darum handelt, das bestimmte System der Bedingungen aufzustellen, welches über die Legitimität der Funktionen entscheidet, so werden wir dasselbe zunächst so wählen, daß es mit den Voraussetzungen, die man zuweilen bewußt, meistens unbewußt, bei analytischen Untersuchungen zu machen pflegt, übereinstimmt und nennen

legitime Funktionen solche, welche für alle reellen Werte des Argumentes, mit Ausnahme einzelner, auf jedem endlichen Intervall in niemals unendlicher Zahl vorhandener Punkte, bestimmt, endlich und stetig sind, und deren sämtliche Differentialquotienten dieselben Eigenschaften besitzen.

Solche Funktionen sind aber nach dem TAYLORSchen Satze immer entwickelbar²²⁾, d. h. es läßt sich stets ein Intervall für x angeben, in dem die Entwicklung von $f(a+x)$ nach aufsteigenden Potenzen

¹⁹⁾ Ebenda S. 99

von x konvergiert, solange a nicht mit einem der Punkte zusammenfällt, in dem die Funktion oder einer ihrer Differentialquotienten unbestimmt, unendlich oder unstetig wird. Die Werte einer Funktion sind daher zwischen zwei Unstetigkeitspunkten immer in der Weise miteinander verbunden, daß, wenn nur die Funktionswerte für das kleinste endliche Intervall bekannt sind, sie für die ganze Strecke mit Notwendigkeit bestimmt sind. Sind in einer solchen Strecke für zwei endliche Intervalle zwei Reihen von Werten gegeben, so kann es niemals zweifelhaft sein, ob sie zu derselben Funktion gehören oder nicht.

Wie steht es aber mit der Fortsetzung einer solchen Funktion über jene Unstetigkeitspunkte? Die TAYLORSchen Reihen reichen über diese nicht hinaus, und es bleibt gänzlich unbestimmt, was wir unter einer zwischen zwei Unstetigkeitspunkten völlig bestimmten Funktion außerhalb dieser Strecke verstehen sollen. Jene Punkte trennen, wie Barrieren, die verschiedenen Strecken, und es kann nach der gegebenen Definition nicht entschieden werden, ob zwei, durch einen solchen Punkt getrennte Wertreihen zu einer und derselben Funktion gehören oder nicht.

In welcher Weise man durch Einführung eines neuen Prinzips diese Hindernisse zu überwinden versucht hat, ist schon oben S. 47 gezeigt; eine unzweideutige und annehmbare Form ist aber diesem Prinzipie weder jemals gegeben worden, noch dürfte sie sich überhaupt finden lassen.

Die Wissenschaft mußte andere Mittel suchen, um die Funktionen über jene Unstetigkeitspunkte hinaus fortzusetzen; und man entdeckte, daß man, anstatt diese Punkte gewissermaßen zu *überspringen*, sie vielmehr *umgehen* könne, indem man die Variablen nicht nur reelle, sondern auch *komplexe* Werte durchlaufen läßt.

Ist hierdurch dargetan, daß man bei der Beschränkung der Variabilität auf reelle Werte des Argumentes zu einer die Bedürfnisse der Analyse befriedigenden Definition der Funktion nicht gelangen kann, so ist damit das Ziel erreicht, das wir uns in dieser Abhandlung gesteckt haben.

Sind wir aber einmal im Gange, so wird es erlaubt sein, noch einige Schritte weiter zu tun, um eine Perspektive in das neue Gebiet zu gewinnen.

Eine Funktion w zweier Variablen x, y wird dann als die einer Veränderlichen $(x+yi)$ angesehen werden können, wenn sie sich in eine Form bringen läßt, daß x und y in ihr nur in der Verbindung $(x+yi)$ erscheinen. Diese vorläufige Definition hat nur dann eine Bedeutung, wenn die Funktion nach dem Typus algebraischer Funktionen gebaut ist, d. h. aus ihren Veränderlichen durch die vier arithmetischen Grundoperationen abgeleitet. Sie hat dann, wenn ihre Differentialquotienten überhaupt bestimmt sind, die Eigenschaft, daß

$$i \frac{\partial w}{\partial x} = \frac{\partial w}{\partial y}.$$

Während nun jene Definition ihre Bedeutung verliert, wenn w nicht durch arithmetische Operationen bestimmt wird, so behält diese Gleichung ihre Bedeutung in jedem Falle; und wir definieren nun eine Funktion w zweier Veränderlichen x, y als eine *monogene* Funktion der komplexen Veränderlichen $(x+yi)$, indem wir letztere Eigenschaft für permanent erklären.

Damit ist aber zugleich gesetzt, daß jene Differentialquotienten $\frac{\partial w}{\partial x}$, $\frac{\partial w}{\partial y}$ überhaupt existieren, einen bestimmten und endlichen Wert haben; denn hätten die Quotienten der Inkremente von w und x, y innerhalb eines endlichen Stückes der Ebene, in welcher die komplexe Veränderliche

$(x + yi)$ dargestellt wird, keine bestimmten und endlichen Grenzwerte, so ließe sich in demselben gar nicht konstatieren, ob die Funktionen als die *einer* komplexen Variablen angesehen werden kann. So ergibt sich die Voraussetzung, daß eine legitime Funktion im allgemeinen bestimmte und endliche Differentialquotienten hat, als eine fundamentale, und wir definieren nun:

Eine monogene Funktion von $(x + yi)$ heißt eine veränderliche komplexe Größe w , wenn sie sich mit x und y so ändert, daß im allgemeinen, d. h. abgesehen von einzelnen Punkten und Linien, die Differentialquotienten $\frac{\partial w}{\partial x}$, $\frac{\partial w}{\partial y}$ überall bestimmt und endlich sind und der Bedingung

$$i \frac{\partial w}{\partial x} = \frac{\partial w}{\partial y}$$

genügen.

Dies ist nun das allgemeine einfache Gesetz, dem alle legitimen, d. h. monogenen Funktionen untergeordnet werden, und das wir als die RIEMANNSche Definition²⁰⁾ zu bezeichnen haben. Da nach bekannten Sätzen mit der Endlichkeit, Stetigkeit und Bestimmtheit der Funktion selbst und ihrer ersten Differentialquotienten dieselben Eigenschaften auch für alle andern Differentialquotienten gesetzt sind, so kann eine monogene Funktion immer nach dem TAYLORSchen Lehrsatz entwickelt werden, und die obige Bestimmung der im reellen Gebiete stetigen Funktion erscheint jetzt unserer neuen Definition untergeordnet.

Nun erst gewinnt das Problem, eine für eine endliche Strecke oder innerhalb eines endlichen Flächenstückes gegebene Funktion fortzusetzen, d.h. für alle Punkte der ganzen Ebene zu bestimmen, eine feste Bedeutung, wenn man eben die Funktion obiger Bedingung unterwirft. Nun stehen die Funktionswerte für die verschiedenen Argumente in unlöslicher Verbindung miteinander, und es wird nicht zweifelhaft sein, ob zwei in verschiedenen Teilen der Ebene gegebene Wertreihen zu derselben Funktion gehören, — wenn sich nicht jetzt, wo wir jene, die Fortsetzung im reellen Gebiete störenden, und alle anderen Unstetigkeitspunkte, leicht umgehen können, neuen Barrieren erheben in Gestalt von *Linien*, in denen die Funktion aufhört, bestimmt, endlich und stetig zu sein. Obgleich wir nämlich Unstetigkeiten, die über ganze Flächenstücke ausgedehnt sind, durch obige Definition der Monogenität ausgeschlossen haben, so konnten wir nicht umhin, Unstetigkeiten in Punkten und Linien zuzulassen, wenn sie sich mit Notwendigkeit aus dem Verlaufe der Funktion in ihren stetigen Teilen ergeben. Daß solche Linien, in denen die Funktion punktiert oder gar total unstetig ist, trotz der Monogenität in der Tat vorkommen, ist im Anhang (Note III.) gezeigt, und zwar scheidet in den angeführten Beispielen diese Unstetigkeitslinie als Kreis den inneren Raum, in dem die Funktion zunächst festgestellt gedacht wird, in der Weise von dem äußeren ab, daß beide durch keinen Flächenstreifen von endlicher Breite, in dem die Funktion die zur legitimen Fortsetzung notwendigen Eigenschaften hätte, miteinander zusammenhängen.

Die Funktion ist dann auf die innere Kreisfläche beschränkt, sie verliert für Punkte außerhalb des Kreises ihre Bedeutung, d. h. unser Funktionsbegriff gibt keinerlei Aufschluß darüber, was wir uns unter Funktionswerten im äußeren Raume zu denken haben.

²⁰⁾ Grundl. f. e. allg. Th., S. 2.

Es ist nun aber in der Note III. das merkwürdige Faktum hervorgehoben, daß es trotzdem analytische Entwicklungen solcher Funktionen geben kann, welche auch außerhalb ihre Bedeutung nicht verlieren, sondern innerhalb und außerhalb des Kreises in gleicher Weise gebildet, und mit Ausnahme jener Kreisperipherie überall konvergent und monogen sind.

Insofern man nun solchen Entwicklungen, welche zu beiden Seiten der geschlossenen Unstetigkeitslinien dieselben bleiben, die Kraft nicht absprechen kann, die innerhalb gegebene Funktion auch außerhalb des Kreises zu repräsentieren, so erhebt sich die Frage:

In welchem Sinne sind diese beiden, durch eine übersteigliche Schranke getrennten Wertreihen als zu *einer* Funktion gehörig anzusehen²³⁾?

Note I (zu S. 54)

Begriff der Grenze

Gegenüber den häufig mangelhaften, teils zu viel, teils zu wenig enthaltenden Definitionen des in unsere Untersuchungen tief eingreifenden Begriffes der »Grenze« glaube ich hier seine strenge Bestimmung in aller Kürze geben zu sollen, indem ich in bezug auf die Begründung des Folgenden auf eine andere kleine Arbeit von mir verweise²¹⁾. Der Begriff Funktion ist hier im Sinne von § 1 zu nehmen:

Wenn es für eine Funktion $f(x)$ eine bestimmte endliche Größe A gibt, und ein Intervall von dem festen Punkte $x = a$ aus bis $x = a + \varepsilon$ für jedes beliebig kleine σ bestimmt werden kann, so daß die Differenz $A - f(a + \delta)$ jene Größe σ numerisch nicht übersteigt, während δ *alle* innerhalb jedes Intervalles 0 bis ε liegenden Werte, $\delta = 0$ *ausgeschlossen*, durchläuft, so heißt A die *Grenze*, welcher sich $f(a + \delta)$ mit unendlich abnehmendem δ unendlich annähert.

Wird hierin ε positiv vorausgesetzt, so habe ich, indem ich $f(x)$ als Ordinate zu der Abszisse x konstruiert denke, A die Grenze genannt, welcher sich $f(x)$ nähert, wenn man von rechts her zu dem Punkte $x = a$ gelangt. Die Grenze, der sich $f(a - \delta)$ nähert, nenne ich die Grenze auf der linken Seite.

Wenn es dagegen *keine* bestimmte endliche Größe A gibt, so daß die Differenz $A - f(a + \delta)$ für *alle* von Null verschiedenen $\delta < \varepsilon$ *nicht immer* numerisch kleiner als eine beliebig kleine Größe σ bleibt, wie klein man auch ε bestimmen möge, so sagen wir, es nähere sich $f(a + \delta)$ für unendlich abnehmende δ *keiner* Grenze.

In vielen, und zwar in den meisten uns oben interessierenden Fällen ist es nicht möglich, jene Grenzwerte, denen sich Funktionen annähern, ihrer wirklichen Größe nach zu bestimmen; vielmehr wird man sich begnügen müssen, die *Existenz* einer Grenze überhaupt festzustellen, und es dient dazu folgender aus dem Begriffe der Grenze abgeleiteter Satz:

Wenn sich ein Intervall ε so bestimmen läßt, daß für alle in diesem enthaltenen, von Null verschiedenen δ die Differenz

$$f(a + \varepsilon) - f(a + \delta)$$

²¹⁾Encyklopädie von ERSCH u. GRUBER. Art. Grenze.

numerisch kleiner als ein beliebig kleine Größe σ ist, so nähert sich $f(a + \delta)$ mit unendlich abnehmenden δ einer endlichen bestimmten Grenze²⁴.

Man bemerkt übrigens, daß nach der in § 2 festgesetzten Terminologie der Begriff auch so gefaßt werden könnte

Einer Grenze nähert sich $f(a + \delta)$ für unendlich abnehmende δ , wenn $f(x)$ in unmittelbarer Nähe von $x = a$ (auf der rechten Seite) stetig ist.

Note II (zu S. 86)

**Beispiele von Funktionen, welche in ganzen Linien
total unstetig sind.**

Ein in mancher Hinsicht interessantes Beispiel einer total unstetigen Funktion gibt die Reihe

$$f(x) = \sum_{m=1}^{\infty} \frac{w^m}{(m \sin mx\pi)^2},$$

worin w einen positiven echten Bruch bedeuten soll²⁵.

Schluß : Didaktische Konsequenzen

Die Ideen von Cauchy haben mit einiger Verspätung auch die Lehre der Analysis an den deutschen Universitäten und Gymnasien beeinflusst. J.A. GRUNERT (1797–1872) hat in Halle bei Pfaff und in Göttingen bei Gauß studiert. Er war dann Gymnasiallehrer, bis er 1834 als Professor an die Universität Greifswald berufen wurde. In dem von ihm gegründeten „Archiv der Mathematik und Physik“ sollte eine Verbindung zwischen Schul- und Universitätsmathematik hergestellt werden. Diese Zeitschrift versuchte die Auffassungen von Cauchy zu propagieren, strebte dabei aber Kompromisse zwischen dem komplizierten Vorgehen von Cauchy und der heuristischen Leichtigkeit der älteren Verfahren an. Grunert lobt z.B. in seiner Besprechung das „vollständige und reichhaltige Lehrbuch der sogenannten Analysis des Endlichen, der Differentialrechnung und der Integralrechnung“ von Salomon (1844), weil es die beiden Methoden vereinigt, um auf das Publikum Rücksicht zu nehmen. Er erkennt an, daß „z.B. die Methode der unbestimmten Koeffizienten bei allen ihren Unvollkommenheiten dem Praktiker, dem es oft weniger auf die Methode an sich, als vielmehr auf die Auffindung eines Resultats ankommt, in vielen Fällen vortreffliche Dienste leisten kann.“ (Zitiert nach Jahnke: „Motive und Probleme der Arithmetisierung der Mathematik in der ersten Hälfte des 19. Jahrhunderts“. Archive for History of Exact Sciences, Vol. 37, Nr. 2 (1987) S. 166). Wir zitieren weiter Jahnke:

„Eines der ersten deutschen Lehrbücher, das sich ganz bewußt auf den Standpunkt CAUCHYS stellte, war die „Algebraische Analysis“ von OSKAR SCHLÖMILCH aus dem Jahre 1845. Dieses Buch versuchte, bei aller Radikalität des Eintretens für CAUCHYS Ansatz, die Idee eines Kompromisses zwischen Praktikabilität und Strenge im Sinne CAUCHYS zu realisieren. In dem historisch außerordentlich interessanten Vorwort bemerkte der Autor zunächst, daß in den 20 Jahren seit dem Erscheinen der Werke CAUCHYS eine neue Epoche in der Geschichte der Wissenschaft begonnen habe, deren Eigentümlichkeit in der *Kritik der Methode* bestehe. Zu den grundsätzlichen Anforderungen an eine wissenschaftliche Behandlungsweise der Mathematik rechne er: heuristischen Gedankengang, Strenge und architektonisches Gefüge. Vergleiche man unter diesen Gesichtspunkten die frühere Behandlungsweise der Analysis mit der jetzigen, durch CAUCHY angeregten, so sehe man, daß die ältere Methode die Forderung nach heuristischem Gedankengang immer, die anderen aber sehr wenig befriedige. Diese heuristische Vorgehensweise habe zwar für den Praktiker einigen Wert, aber für ein System der Wissenschaft sei sie unbrauchbar, da sie nur isolierte Resultate liefere ohne die „Schnur, an der man die gefundenen Perlen aufreihen könnte, auch nur im Entferntesten ahnen zu lassen.“ Die Methode des Erfinders sei nicht immer die systematisch beste.

Dagegen bescheinigte er CAUCHYS „Cours d’analyse“ die „größte Strenge bei vieler Kürze in der Entwicklung“. Allerdings leide bei CAUCHY die „Schönheit des architektonischen Baus“ und das „Leben der Erfindung“ fehle gänzlich. „Dies mag auch vielleicht bei Manchem den Glauben hervorgerufen haben, daß Cauchy bloß

Seiltänzerkunststücke mache, und zeigen wolle, daß er Fertigkeit genug besitze, um Alles auf den Kopf stellen zu könne.“ Daraus leitete er für sein eigenes Werk die Zielstellung ab: „Zwischen diesen beiden Extremen habe ich einen Mittelweg einzuschlagen gesucht, welcher das Interesse des heuristischen Gedankenganges mit der Strenge des französischen Analytikers vereinen und dem Gange ein besseres architektonisches Gefüge verleihen soll, als man bisher an diesem Theile der Mathematik bemerkt hat.“

CANTOR schrieb über Schlömilchs Buch aus dem Jahre 1845: „Sie war allerdings in Anlehnung an CAUCHYS Analyse algébrique von 1821 entstanden, aber in ihrer Ausarbeitung ein durchaus selbständiges eigenartiges Werk geworden, eigenartig auch durch die erfrischende Grobheit der in der Vorrede sowohl als in Fußnoten sich kundgebenden Angriffe gegen das, was wir als Sorglosigkeit der Entwicklung bezeichnet haben. In den späteren Auflagen hat SCHLÖMILCH die polemischen Abschweifungen allmählich gestrichen. Sie waren nicht mehr notwendig, aber daß sie entbehrlich wurden, ist unzweifelhaft eine Folge der durch die erste Auflage vollzogenen Aufrüttelung der deutschen Mathematiker.“ (Cantor, 1901, S. 262)

Schlömilch war übrigens später als Beamter im sächsischen Ministerium (bis 1885) ein wichtiger Förderer der Entwicklung des mathematischen Seminars an der Universität Leipzig, wo man sich unter Felix Klein sehr um die Verbesserung des mathematischen Unterrichts bemühte. Jahnke schreibt weiter: „Die Rezeption und Propagierung des Cauchyschen Ansatzes durch SCHLÖMILCH ist in vielerlei Hinsicht charakteristisch für die Situation, auf die dieser Ansatz in Deutschland traf. So enthüllten die vorstehend referierten Überlegungen nicht nur die Notwendigkeit des Kompromisses mit den Erfordernissen der Praxis, sondern SCHLÖMILCH führte mit dem Begriff des „architektonischen Gefüges“ zusätzlich einen Terminus ein, der in dem philosophischen und bildungstheoretischen Kontext eine Rolle spielte. In diesem Kontext wurden Theorienideale diskutiert und entwickelt, die in der ersten Hälfte des 19. Jahrhunderts in Deutschland ein wirksames Motiv auch für die Entwicklung der Mathematik darstellten.“

Jahnke stellt fest, daß die Haltung von Schlömilch dogmatische Züge annahm: „Was für Cauchy, Abel, Dirichlet und andere Analytiker zweifellos eine bestimmte Gesamtsicht derjenigen Begriffe war, mit denen man angemessen die „Transzendente Analysis“ behandeln konnte, reduzierte sich in der Propagierung durch GRUNERT und SCHLÖMILCH auf die Durchsetzung bestimmter Normen mathematischer Strenge.“ (S. 169)

Das von Cauchy eingeleitete **begriffliche Denken** fand nur sehr zögernd Eingang in die Lehrbücher; und bis heute ist dort wenig von der **Öffnung des mathematischen Denkens** zu spüren, welches damit bewirkt worden ist. (Die „begriffliche“ Herangehensweise (an die Gegenstände der Mathematik) wird hier als Gegensatz verstanden zu einer „formalen“, im Sinne einer an speziellen Darstellungen orientierten Herangehensweise.)

Dirichlet war der erste, der die neue Denkungsart klar ausgesprochen und in konkreten Fällen verwirklicht hat. Bei ihm zeigte sich, daß überall dort die größten Durchbrüche und die überraschendsten Erfolge erzielt wurden, wo es gelang, sich von überkommenen Prozeduren

zu lösen und sozusagen das „Wesen“ des in Frage stehenden Problems zu studieren. (Jahnke, S. 173). Das Ideal des begrifflichen diskursiven Denkens hatte es sehr schwer, gegen die Programme der Elementarisierungen einen anderen, weniger auf Disziplinierung bedachten Typus der Lehre von Mathematik zu initiieren.

Dirichlet scheint dem Humboldtschen Ideal von „Einheit von Forschung und Lehre“ besonders gut entsprochen zu haben. H. Minkowski schreibt in seinem Nachruf: Peter Gustav Lejeune Dirichlet und seine Bedeutung für die heutige Mathematik. (Jahresbericht der DMV 14 (1905) S. 149–163) (zitiert nach Jahnke):

„Dirichlets durch wunderbare Klarheit und Einfachheit berühmter Beweis, daß jedem Minimum potentieller Energie Stabilität des Gleichgewichts entspricht, ist eine zweite Tat solcher Art, daß er sich von verkehrten Hilfsmitteln freimachte, indem er statt der analytischen Regeln für die Bestimmung der Minima einer Funktion nur den ursprünglichen Begriff des Minimums heranzog. Immer wieder sind doch bedeutende Fortschritte in der Mathematik auf der Stelle errungen, so wie man nur wahrnimmt, daß Umstände, die stets als zusammengehörig betrachtet wurden, nichts miteinander zu tun haben.“

Dirichlets Lehrtätigkeit beschrieb Minkowski so: *„Sein Vortrag war dadurch so eindringlich, weil es schien, als wenn er im Begriff stünde, eben erst den ganzen Bau zu schaffen; es war in hohem Grade fesselnd, ihm bei dieser Arbeit zu folgen. Er entwickelte den Stoff in vollster Natürlichkeit. Kein Kunstgriff trat auf als deus ex machina, tragisch geschürzte Knoten zu unerwartet günstiger Lösung zu führen.“* Minkowski spricht von dem anderen Dirichletschen Prinzip „mit einem Minimum an blinder Rechnung, einem Maximum an sehenden Gedanken die Probleme zu zwingen“, von dem die Neuzeit in der Geschichte der Mathematik datiere.

Die Forderung nach mathematischer Strenge wirkt sich auch im heutigen Analysisunterricht noch sehr oft als eine Fessel der begrifflichen Herangehensweise aus. Es ist aber zu beachten, daß das Bemühen um Strenge nicht nur disziplinierende sondern auch antidogmatische Effekte haben kann. Die Studenten können den Dozenten korrigieren ohne seine überlegene Erfahrung in Frage zu stellen. Max Dehn meinte in seiner Frankfurter Universitätsrede 1928 („Über die geistige Eigenart des Mathematikers“): *„Zweifelloos wird jeder, der sich mit Mathematik beschäftigt, hierdurch stark beeinflusst; sein Charakter wird in eigenartiger Weise geformt. Die Mathematik ist der einzige Lehrstoff, der ganz undogmatisch vorgetragen werden kann. Deswegen spielt der mathematische Unterricht eine hervorragende Rolle an allen höheren Schulen seit dem Altertum. In richtiger Weise vorgetragen, wird die Mathematik den Schüler befähigen, soweit es ihm nach seiner Anlage möglich ist, klar und selbständig zu denken. Für seine mathematischen Arbeiten kann er wirklich die volle Verantwortung übernehmen.“*

Die mathematische Strenge kann offenbar vielerlei Rollen spielen. Bei Weierstraß, dem großen Advokaten der Strenge, ging es vor allem darum, die Vorurteile über Funktionen, welche die Mathematiker seiner Zeit mit sich herumtrugen, zu sortieren. Durch seine Arbeiten wurden zum einen die Kluft zwischen reeller und komplexer Analysis deutlich; und zum

anderen wurde der Boden geschaffen für die (um 1900) geglückte Vereinigung der beiden Linien der komplexen Analysis. (Die auf die Cauchy–Riemannschen Differentialgleichungen und das Dirichlet–Prinzip gegründete Funktionentheorie erwies sich als inhaltlich gleichwertig mit der auf die konvergenten Potenzreihen gegründeten Theorie.) Weierstraß selbst hat den Funktionsbegriff stark an seine analytische Darstellung gebunden gesehen und den relationalen (Dirichletschen) Begriff der „willkürlichen“ Funktion als inhaltsleer abgelehnt. (Jahnke, S. 163 u. 173). (Richenhagen). Dennoch hat er (vgl. Purkert S. 72) die Arbeiten Cantors wohlwollend verfolgt. Er hat Cantor sogar Anregungen gegeben. Weierstraß trat aber nicht öffentlich für die revolutionären Ideen ein.

Insofern kann sich ein Mathematikdozent unserer Tage, der sich im Analysis–Unterricht lieber an Formeln hält als an Begriffe, nicht nur auf die Reaktionäre um Kronecker berufen, sondern im gewissen Sinne auch auf Weierstraß. Es sollte ihm aber klar sein, daß die Weierstraßsche Strenge nichts zu tun hat mit dem Einüben begriffsloser „streng mathematischer“ Argumentation, zu welchem die Analysis vielerorts mißbraucht wird.

Plädoyer für eine Modernisierung der Lehre

Nach meiner Meinung erreichen die heute gängigen ersten Einführungen in die Analysis im Betreff des Funktionsbegriffs nicht das gebotene Abstraktionsniveau. Dem „Dirichletschen Funktionsbegriff“ wird ähnlich wie dem Cantorsche Mengenbegriff nur ein Lippenbekenntnis geschenkt. In dem, was entwickelt wird, orientiert man sich weiterhin recht unentschlossen am oben beschriebenen Permanenzprinzip von H. HANKEL (1870).

Im Hinblick auf die Integrationstheorie kann man feststellen: Der Begriff der borelmeßbaren Funktion, den Hankel noch nicht kannte und der ja auch den Ideen Hankels nicht entgegenkommt, gilt den Autoren der elementaren Lehrbücher als allzu abstrakt für den Begriffshorizont des Anfängers. In der Konsequenz hält man sich mit dem Riemann–Integral auf, anstatt die natürliche Integrationstheorie von Lebesgue (zunächst einmal ohne Beweise) vorzustellen; die Integrationstheorie schwankt in unerfreulicher Weise zwischen dem begrifflichen Niveau der Stammfunktionenbildung und dem Exhaustionsprinzip; sie stößt nicht vor auf das begriffliche Niveau der linearen monotonstetigen Funktionale.

Was die Differentiation betrifft, plagt man sich mit Absonderlichkeiten des Verhaltens einer Funktion in einzelnen Punkten herum (partielle Ableitbarkeit versus totale Differenzierbarkeit z.B.), anstatt entschlossen die natürlichere Theorie der stetigen Differenzierbarkeit auf einer Mannigfaltigkeit (mit diversen lokalen Koordinatensystemen) anzupeilen.

In der Fortsetzung unserer Einführung in die Analysis wollen wir uns mit metrischen Räumen im Sinne der Punktmengentopologie, mit der Lebesgueschen (oder Daniellschen) Integrationstheorie und mit der Theorie der Differentialformen auf einer (Riemannschen) Mannigfaltigkeit befassen.

Kapitel C

ANALYSIS I

C.1 Zahlen aus der Sicht der Mengenlehre

1.1 Aufbau des Zahlensystems

Die natürlichen Zahlen sind allgemein bekannt:

$$1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, \dots$$

oder auch

$$1, 10, 11, 100, 101, 110, 111, 1000, 1001, \dots$$

Es kommt nicht auf die Zahlsymbole an, sondern darauf, daß die Aufzählung mit der Zahl „Eins“ beginnt, daß es zu jeder natürlichen Zahl einen wohlbestimmten Nachfolger gibt und daß man durch das Fortschreiten von der Eins aus jede natürliche Zahl erreicht. Man kann das mit Formeln ausdrücken (die Nachfolgerrelation wird üblicherweise mit „ $'$ “ bezeichnet) und man kann sich dann um eine axiomatische Charakterisierung der strukturierten Menge $(\mathbb{N},')$ bemühen, wo $'$ die Nachfolgerrelation bedeutet. Man kann auf dieser Grundlage („Definition durch Induktion“ usw.) die Operationen der Addition und der Multiplikation natürlicher Zahlen definieren. Und man kann dann natürlich für diese Operationen allerlei beweisen.

Dieser axiomatische Zugang zum System der natürlichen Zahlen $(\mathbb{N}, 1, +, \cdot)$ ist im Detail durchgeführt in dem Büchlein

E. LANDAU: Grundlagen der Analysis, 3. Auflage, New York 1960.

Wenn man zur Menge $\mathbb{N}$ der natürlichen Zahlen die Zahl 0 und die negativen Zahlen hinzufügt, gelangt man zur Menge $\mathbb{Z}$ der ganzen Zahlen. Die Erweiterung der Addition und der Multiplikation auf $\mathbb{Z}$ ist wohlbekannt, in $\mathbb{Z}$ kann man unbeschränkt subtrahieren. Zu jedem $a, b \in \mathbb{Z}$ existiert genau ein x mit $a + x = b$. Man schreibt $x = b - a$.

Auf der anderen Seite kann man sich um die Division sorgen. Nicht für jedes Paar natürliche Zahlen a, b gibt es eine natürliche Zahl x mit $a \cdot x = b$. Unbeschränktes Dividieren wird aber möglich, wenn man den Zahlbegriff erweitert. Die Menge der (echt) positiven rationalen Zahlen konstruiert man als die Menge der Äquivalenzklassen von Paaren natürlichen Zahlen, wobei

$$(a_1, b_1) \sim (a_2, b_2) \iff a_1 b_2 = a_2 b_1 .$$

Die Äquivalenzklasse von (a, b) bezeichnet man mit $\frac{a}{b}$.

$$\frac{a_1}{b_1} = \frac{a_2}{b_2} \iff a_1 b_2 = a_2 b_1 .$$

Diese Äquivalenzklassen kann man addieren und multiplizieren

$$\frac{a}{b} \cdot \frac{c}{d} = \frac{a \cdot c}{b \cdot d}, \quad \frac{a}{b} + \frac{c}{d} = \frac{ad + bc}{bd} \quad (\text{„Bruchrechnen“})$$

Wenn man die Null und die formalen Ausdrücke $-\frac{a}{b}$ hinzunimmt, erhält man die Menge $\mathbb{Q}$ der rationalen Zahlen. Wenn man die Operationen der Addition und Multiplikation in der naheliegenden Weise erweitert, erhält man den Körper $\mathbb{Q}$ der rationalen Zahlen

$$(\mathbb{Q}, =, 0, 1, +, \cdot) .$$

Wenn man $\mathbb{Q}$ in der bekannten naheliegenden Weise ordnet, erhält man den linear angeordneten Körper

$$(\mathbb{Q}, \leq, 0, 1, +, \cdot) .$$

Dies ist die bekannte Konstruktion des archimedisch angeordneten Körpers der rationalen Zahlen aus den natürlichen Zahlen. Nach allgemeiner Meinung sind die natürlichen Zahlen unserer Anschauung oder unserem Denken unmittelbar vorgegeben. KRONECKER (1823 – 1891), der einflußreiche Berliner Mathematiker, hat auch das Umgekehrte vertreten. Er soll gesagt haben: „*Die natürlichen Zahlen hat der liebe Gott geschaffen, alles andere ist Menschenwerk.*“ Hier trennt sich nun die Mathematik des 20. Jahrhunderts von der des 19. Jahrhunderts. Die allgemeine Mengenlehre wurde die umfassende Basis aller mathematischer Argumentation. Der Bruch begann bei den Irrationalzahlen. Die Haltung Kroneckers und seiner Anhänger zu den Irrationalzahlen war für GEORG CANTOR, den Schöpfer der Mengenlehre ein geradezu existentielles Problem. (Siehe in PURKERT u. ILGAUDS: Georg Cantor 1845–1918).

In der Tat liegt in der „Konstruktion“ der reellen Zahlen durch die Dedekindschen Schnitte eine radikale erkenntnistheoretische Wende – und die hat dann etwas später tatsächlich eine tiefgreifende Krise der mathematischen Logik herbeigeführt. Wir zitieren nochmals aus K. Reidemeisters Artikel „*Eine Begründung der Infinitesimalrechnung*“:

„Alle reellen Zahlen sind nach Dedekind soviel wie ‚alle Schnitte rationaler Zahlen‘, also fast soviel wie ‚alle Eigenschaften rationaler Zahlen‘ oder wie ‚alle Mengen aus rationalen Zahlen‘ – und der Begriff ‚alle Mengen‘ ist bei unvorsichtigem Gebrauch widerspruchsvoll. Um den wesentlichen Bezug auf alle reellen Zahlen kommt man aber z.B. bei der Definition der kleinsten oberen Grenze nicht herum.“

Prinzip Zu jeder nichtleeren nach oben beschränkten Menge A von reellen Zahlen existiert die kleinste obere Schranke.

Was es bedeutet, daß a die kleinste obere Schranke von A ist, ist anschaulich klar. Man kann es aber auch formal schreiben in einer Notation, die wir demnächst entwickeln werden

$$a = \sup A \iff (\forall x \in A : x \leq a) \wedge (\forall c : (\forall x \in A : x \leq c) \rightarrow (a \leq c)) .$$

Der Beweis des Prinzips ergibt sich sofort aus dem Dedekindschen Vollständigkeitsaxiom. Das, was den Logikern Probleme bereitet, ist der unbekümmert benützte Begriff einer „beliebigen“ nach oben beschränkten Zahlenmenge. Es ist nicht die Rede davon, wie eine solche Menge „konkret“ vorgegeben sein könnte, daß sie durch irgendeine Vorschrift gegeben sein sollte. Das nichtkonstruktive Element in dem Prinzip ist in der Tat der springende Punkt. Mit den Problemen, die daraus entstehen, haben sich viele große Mathematiker des 20. Jahrhunderts auseinandergesetzt. Siehe z.B. H.N. JAHNKE: „*Hilbert, Weyl und die Philosophie der Mathematik*“, Mathematische Semesterberichte Bd. 38, S. 159–179 (1991).

1.2 Cantors Mengendefinition

Die „Beiträge zur Begründung der Mengenlehre“ von Georg Cantor aus dem Jahre 1895 beginnen mit den Worten:

„Unter einer Menge M verstehen wir jede Zusammenfassung von bestimmten, wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens (welche die „Elemente“ von M genannt werden) zu einem Ganzen.“

Diese „Definition“ des Begriffs Menge ist der Ausgangspunkt der „naiven“ Mengenlehre, die das Bild der Mathematik des 20. Jahrhunderts bestimmt. Die Objekte der mathematischen Betrachtung werden nicht irgendwie konstruiert, sie werden als wohlunterschiedene Objekte gegeben angenommen. Die Gleichheit ist die grundlegende Relation der Mengenlehre. Was eine Gleichheitsrelation ist, kann man axiomatisch fassen.

Definition Eine Gleichheits- oder Äquivalenzrelation ist eine zweistellige Relation mit den Eigenschaften

- (i) $x = x$ für jedes x (Reflexivität)
- (ii) Wenn $x = z$ und $y = z$, dann $x = y$ (Komparativität)

(Insbesondere gilt: Wenn $x = y$ dann $y = x$ (Symmetrie).)

Ein Prototyp von wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens sind die ganzen Zahlen. Eine Äquivalenzrelation in der Menge $(\mathbb{Z}, =)$ ist neben der wirklichen Gleichheit z.B. auch die „Gleichheit modulo n “ für eine vorgegebene natürliche Zahl n . (Für $x, y \in \mathbb{Z}$ definiert man $x \sim y \iff (y - x)$ ist ein ganzzahliges Vielfaches von n .)

Die allgemeine Mengenlehre hat es mit abstrakten Gleichheiten zu tun. Zum Gleichheitsbegriff schreibt der Philosoph und Mathematiker E. CASSIRER (1920):

„Gegenüber der empirischen Lehre, die die Gleichheit bestimmter Vorstellungsinhalte als eine selbstverständliche psychologische Tatsache hinnimmt und für die Erklärung des Prozesses der Begriffsbildung verwendet, ist mit Recht darauf verwiesen worden, daß von Gleichheit irgendwelcher Elemente nur dann mit Sinn geredet werden kann, wenn bereits eine Hinsicht festgestellt ist, in welcher die Elemente als gleich oder ungleich bezeichnet werden sollen. Diese Identität der Hinsicht, des Gesichtspunkts, unter welchen die Vergleichenng stattfindet, ist jedoch ein Eigenartiges und Neues gegenüber den verglichenen Inhalten selbst.“

Die allgemeine Mengenlehre kann nicht Gegenstand einer Einführung in die Analysis sein. Wir gebrauchen die Symbolik der Mengenlehre als eine Art Kurzschrift, und wir benutzen die Vorstellungsweisen der Mengenlehre in der üblichen naiven Weise.

1.3 Spezielle Mengen von Zahlen

Wenn man bei den oben diskutierten Zahlbereichen von allen Relationen und Verknüpfungen (mit Ausnahme der Gleichheitsrelation!) absieht, erhält man Mengen mit allgemein akzeptierten Namen:

$$\mathbb{N} \subseteq \begin{matrix} \mathbb{Z} \\ \mathbb{Q}_+ \end{matrix} \subseteq \mathbb{Q} \subseteq \mathbb{R} \subseteq \mathbb{C} \quad (\text{Menge der komplexen Zahlen}) \quad .$$

Die Notation ist an einer Stelle nicht ganz einheitlich: $\mathbb{Q}_+$ bezeichnet häufig nicht die Menge der strikt positiven rationalen Zahlen, sondern die Menge der nichtnegativen rationalen Zahlen (die 0 ist also eingeschlossen). $\mathbb{R}_+$ bezeichnet üblicherweise die Menge der nichtnegativen reellen Zahlen. $\mathbb{Z}_+ = \mathbb{N} \cup \{0\}$. Wenn man zu $\mathbb{R}$ die Elemente $+\infty$ und $-\infty$ hinzunimmt, erhält man $\overline{\mathbb{R}} = \mathbb{R} \cup \{-\infty, +\infty\}$. Zu $\mathbb{C}$ nimmt man gelegentlich den unendlichfernen Punkt ∞ dazu und schreibt $\overline{\mathbb{C}} = \mathbb{C} \cup \{\infty\}$.

Die Ordnung auf $\mathbb{R}$ wird zu einer Ordnung auf $\overline{\mathbb{R}}$ fortgesetzt und mit $+\infty$ als größtem und $-\infty$ als kleinstem Element. Die Elemente von $\overline{\mathbb{R}}_+ = \mathbb{R}_+ \cup \{+\infty\}$ kann man addieren und mit nichtnegativen Zahlen multiplizieren, wobei die Konvention $0 \cdot \infty = 0$ üblich ist. Das Dedekindsche Prinzip kann man auch so formulieren: *Jede Teilmenge von $\overline{\mathbb{R}}$ besitzt ein Supremum in $\overline{\mathbb{R}}$* ; dieses Supremum ist $+\infty$, wenn die Menge den Punkt $+\infty$ enthält oder nach oben unbeschränkt ist. Für die leere Teilmenge von $\overline{\mathbb{R}}$ setzt man das Supremum gleich $-\infty$.

Wichtige Teilmengen von $\mathbb{R}$ bzw. $\overline{\mathbb{R}}$ sind die „Intervalle“

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

wobei für a auch der Wert $-\infty$ und für b der Wert $+\infty$ zugelassen ist.

1.4 Cartesische Produkte

Wenn E_1 und E_2 Mengen sind, bezeichnet $E_1 \times E_2$, das *cartesische Produkt*, die Menge der Paare

$$E_1 \times E_2 = \{(x_1, x_2) : x_1 \in E_1, x_2 \in E_2\}.$$

Entsprechend sind die cartesischen Produkte

$$E_1 \times E_2 \times E_3, \quad E_1 \times E_2 \times E_3 \times E_4, \quad \dots$$

definiert.

Wenn die E_i alle dieselbe Menge E sind, schreibt man $E \times E = E^2$, $E \times E \times E = E^3$, ... für die cartesischen Produkte.

So bezeichnet z.B. $[0, 1]^n$ den abgeschlossenen Einheitswürfel in n Dimensionen, also die Menge aller n -Tupel reeller Zahlen $(x_1, \dots, x_n)$ mit $0 \leq x_i \leq 1$ für $i = 1, 2, \dots, n$.

Mit $\{0, 1\}^n$ bezeichnet man die Menge der Null-Eins-Folgen der Länge n ; es handelt sich um eine Menge mit 2^n Punkten.

Relationen Eine zweistellige Relation über der Menge E kann man mit einer Teilmenge von $E \times E$ identifizieren: Für jedes Paar $(x, y) \in E \times E$ steht fest, ob es in Relation steht oder nicht.

Ein wichtiger Typ von zweistelligen Relationen stellen die partiellen Ordnungen dar; ein Spezialfall davon sind die Äquivalenzrelationen.

Eine dreistellige Relation über E kann man als eine Teilmenge von $E \times E \times E$ auffassen: Für jedes Tripel steht fest, ob es in Relation steht oder nicht. Ein Beispiel ist z.B. die Addition in $\mathbb{R}$: $(a, b, c) \in \mathbb{R}^3$ steht in Relation genau dann, wenn $a + b = c$.

1.5 Abbildungen und ihre Graphen

Seien E und F Mengen. Wenn jedem $x \in E$ ein eindeutig bestimmtes Element $\varphi(x) \in F$ zugeordnet ist, dann nennt man diese Zuordnung φ eine **Abbildung** von E nach F und man definiert den **Graph der Abbildung** $\varphi(\cdot)$ als die Menge der Paare

$$\{(x, y) : y = \varphi(x)\} \subseteq E \times F.$$

Nicht jede Teilmenge von $E \times F$ ist der Graph einer Abbildung. Eine Teilmenge Γ von $E \times F$ ist genau dann der Graph einer Abbildung $\varphi : E \rightarrow F$, wenn zu jedem $x \in E$ genau ein $y \in F$ existiert mit $(x, y) \in \Gamma$.

Beispiel Die Menge $\Gamma = \{(x, y) : x \in \mathbb{R}, y \in \mathbb{R}, x^2 + y^2 = 1\} \subseteq \mathbb{R}^2$ ist nicht der Graph einer Abbildung von $\mathbb{R}$ nach $\mathbb{R}$. Jedoch ist die Menge

$$\{(x, y) : x \in [0, 1], y \in \mathbb{R}_+, x^2 + y^2 = 1\}$$

der Graph einer Abbildung von $[0, 1]$ nach $\mathbb{R}_+$. Es handelt sich bei dieser Abbildung um die in $[0, 1]$ definierte Funktion

$$[0, 1] \ni x \mapsto f(x) = \sqrt{1 - x^2} \in \mathbb{R}_+.$$

Notation

a) Wenn φ eine Abbildung von E nach F ist, symbolisiert man das folgendermaßen

$$\varphi : E \longrightarrow F$$

$$\varphi : E \ni x \mapsto \varphi(x) \in F$$

Beachte, daß der Pfeil $\longrightarrow$ stets vom Pfeil $\mapsto$ graphisch unterschieden wird; der erste steht zwischen Mengen, der zweite zwischen Elementen dieser Mengen.

b) Abbildungen kann man hintereinanderschalten. Seien $\Omega', \Omega'', \Omega'''$ Mengen und

$$\varphi : \Omega' \longrightarrow \Omega'', \quad \psi : \Omega'' \longrightarrow \Omega'''$$

Abbildungen. Dann bezeichnet man mit $\psi \circ \varphi$ die zusammengesetzte Abbildung

$$\psi \circ \varphi : \Omega' \ni \omega' \mapsto \psi(\varphi(\omega')) \in \Omega'''.$$

Man notiert auch

$$\begin{array}{ccccc} \Omega' & \xrightarrow{\varphi} & \Omega'' & \xrightarrow{\psi} & \Omega''' \\ & \searrow & & \nearrow & \\ & & \psi \circ \varphi & & \end{array}$$

oder

$$\begin{array}{ccc} \Omega' & \xrightarrow{\varphi} & \Omega'' \\ & \searrow \chi & \downarrow \psi \\ & & \Omega''' \end{array} \quad \chi = \psi \circ \varphi$$

Sprechweisen Eine Abbildung $\varphi : E \rightarrow F$ heißt

- a) injektiv, wenn verschiedene Elemente von E verschiedene Bilder in F haben.
- b) surjektiv, wenn es zu jedem $y \in F$ mindestens ein $x \in E$ gibt, so daß $y = \varphi(x)$
- c) bijektiv, wenn sie sowohl injektiv als auch surjektiv ist.

Bemerkung Zu einer bijektiven Abbildung $\varphi : E \rightarrow F$ gibt es eine wohlbestimmte Umkehrabbildung, d.h. eine Abbildung $\psi : F \rightarrow E$ mit

$$\begin{aligned}\psi(\varphi(x)) &= x \text{ für alle } x \in E \\ \varphi(\psi(y)) &= y \text{ für alle } y \in F.\end{aligned}$$

Man notiert das auch so

$$\psi \circ \varphi = \text{id}_E, \quad \varphi \circ \psi = \text{id}_F.$$

id_E bezeichnet die Identitätsabbildung von E , entsprechend id_F .

Beispiele

- 1) Wenn man die Exponentialfunktion $\exp(\cdot)$ als eine Abbildung von $\mathbb{R}$ nach $(0, \infty)$ auffaßt

$$\mathbb{R} \ni x \longmapsto e^x \in (0, \infty),$$

hat man eine bijektive Abbildung.

Ihre Umkehrabbildung ist der natürliche Logarithmus

$$(0, \infty) \ni y \longmapsto \ln y \in \mathbb{R}$$

$$\ln(e^x) = x \text{ für alle } x \in \mathbb{R}, \quad e^{\ln y} = y \text{ für alle } y \in (0, \infty).$$

- 2) Wenn man die Sinusfunktion $\sin(\cdot)$ auf das Intervall $[-\frac{\pi}{2}, +\frac{\pi}{2}]$ einschränkt, dann erhält man eine bijektive Abbildung auf $[-1, +1]$

$$[-\frac{\pi}{2}, +\frac{\pi}{2}] \ni x \longmapsto \sin x \in [-1, +1].$$

Ihre Umkehrabbildung ist die Arcussinusfunktion

$$[-1, +1] \ni y \longmapsto \arcsin y \in [-\frac{\pi}{2}, +\frac{\pi}{2}]$$

$$\arcsin(\sin x) = x \text{ für alle } |x| \leq \frac{\pi}{2}, \quad \sin(\arcsin y) = y \text{ für alle } |y| \leq 1.$$

(Die Graphen dieser bijektiven Abbildungen sind in A1 geometrisch dargestellt.)

Vergewissern Sie sich der Tatsachen: Wenn man zwei injektive Abbildungen hintereinanderschaltet, erhält man eine injektive Abbildung. Wenn man zwei surjektive Abbildungen hintereinanderschaltet, erhält man eine surjektive Abbildung.

1.6 Abzählbare und überabzählbare Mengen

Eine surjektive Abbildung

$$\varphi : \mathbb{N} \rightarrow M$$

heißt auch eine Aufzählung der Elemente von M .

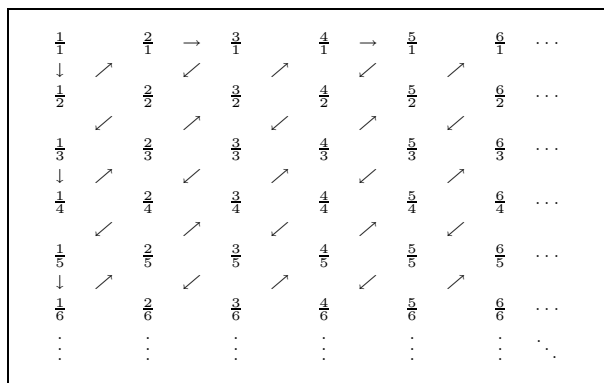
Beispiele

1) Eine Aufzählung von $\mathbb{Z}$ ist z.B. die folgende

$$0, 1, -1, +2, -2, +3, -3, \dots$$

$$\text{d.h. } \varphi(1) = 0, \varphi(2) = 1, \varphi(3) = -1, \varphi(4) = +2, \dots$$

2) Eine Aufzählung von $\mathbb{Q}_+$ ist im folgenden Diagramm beschrieben.



Die Pfeile sind hier keine Abbildungspfeile; sie zeigen nur an, wie man die Menge der Zahlenpaare durchlaufen soll. Wenn man die eingekreisten Paare $\frac{m}{n}$ überspringt, erhält man eine bijektive Abbildung, für die m und n nicht teilerfremd sind

$$\psi : \mathbb{N} \longrightarrow \mathbb{Q}_+ .$$

Definition Eine Menge M heißt abzählbare Menge, wenn es eine Aufzählung der Elemente von M , d.h. eine surjektive Abbildung $\varphi : \mathbb{N} \rightarrow M$ gibt.

Bemerke Jede endliche Menge ist abzählbar. Wenn M abzählbar unendlich ist, dann gibt es eine bijektive Abbildung $\psi : \mathbb{N} \rightarrow M$.

Satz Sind $M_1, M_2, M_3, \dots$ abzählbare Mengen, dann ist auch $M = \bigcup^\infty M_i$ abzählbar.

Beweis Aus irgendwelchen Aufzählungen $\varphi^{(1)}$ für M_1 , $\varphi^{(2)}$ für $M_2, \dots$ kann man wie oben in Beispiel 2) eine Aufzählung von M gewinnen.

Satz Die Menge Ω aller Null-Eins-Folgen ist nicht abzählbar.

Beweis (durch Cantors Diagonalverfahren)

Angenommen, wir hätten eine Abzählung : $\omega^{(1)}, \omega^{(2)}, \dots$

$$\omega^{(1)} = \delta_1^{(1)}, \delta_2^{(1)}, \delta_3^{(1)}, \dots$$

$$\omega^{(2)} = \delta_1^{(2)}, \delta_2^{(2)}, \delta_3^{(2)}, \dots$$

$$\omega^{(3)} = \delta_1^{(3)}, \delta_2^{(3)}, \delta_3^{(3)}, \dots$$

Betrachte nun die Null-Eins-Folge

$$\omega^* = \delta_1^* \delta_2^* \delta_3^* \dots \quad \text{mit} \quad \delta_k^* = 1 - \delta_k^{(k)} \quad \text{für} \quad k = 1, 2, \dots$$

Diese Null-Eins-Folge ω^* ist verschieden von allen $\omega^{(n)}$; denn ω^* differiert zumindest in der n -ten Position von $\omega^{(n)}$. Die Annahme, daß jede Null-Eins-Folge in der Liste $\omega^{(1)}, \omega^{(2)}, \dots$ vorkommt, ist damit widerlegt.

Korollar Die Menge aller reellen Zahlen zwischen 0 und 1 ist überabzählbar.

Beweis Jedes reelle $x \in (0, 1]$ besitzt eine dyadische Entwicklung

$$x = \delta_1 \cdot \frac{1}{2} + \delta_2 \cdot \frac{1}{4} + \delta_3 \cdot \frac{1}{8} + \dots = \sum \delta_i \cdot 2^{-i}$$

mit $\delta_i \in \{0, 1\}$. Man erhält eindeutig bestimmte $\delta_i(x)$, wenn man fordert, daß unendlich viele $\delta_i(x)$ ungleich 1 sind. $x \mapsto (\delta_1(x), \delta_2(x), \dots)$ ist dann bijektiv. Die Menge aller Null-Eins-Folgen kann man aber nicht abzählen, wie wir gesehen haben. Die abzählbar vielen, die nur endlich viele Einsen enthalten, spielen keine Rolle.

Hinweis Cantor nannte die Mächtigkeit von Mengen Cardinalzahlen. Die natürlichen Zahlen sind die endlichen Cardinalzahlen. Cantor hatte in der Tat eine Erweiterung des Zahlbegriffs im Auge. Genauer siehe : H.N. JAHNKE: *Mathematikgeschichte für Lehrer — Gründe und Beispiele*, Mathem. Semesterber. (1996) 43:21–46. Wörtlich heißt es auf S. 31–37:

„ 4 Die Kulturelle Dimension der Mathematik — G. Cantors Einführung und Begründung der transfiniten Ordinalzahlen

Als zweites Beispiel sei die Einführung und Begründung der transfiniten Ordinalzahlen durch Georg Cantor behandelt. Cantor führt die transfiniten Ordinalzahlen als symbolische Größen ein, die gewisse Eigenschaften haben sollen. Für ihn ist dieses Vorgehen nicht so selbstverständlich wie für uns heute, und so führt er eine ganze Serie von Argumenten an, die nur auf dem Hintergrund des kulturellen und philosophischen Klimas seiner Zeit verständlich sind. Diese sollen im folgenden näher betrachtet werden. Cantor wird sich dabei als typischer Vertreter einer Übergangszeit vom 19. Jahrhundert in die mathematische Moderne erweisen, und die Kenntnis seiner Denkweise führt uns in einzigartiger Weise an die Entstehungsbedingungen des modernen mathematischen Denkens.

Um dies zu verstehen, muß man sich zunächst vergegenwärtigen, daß es in der cantorschen Mengenlehre zwei verschiedene Komponenten gibt. Die eine ist die mengentheoretische Beschreibung von Punktaggregaten des linearen Kontinuums oder höherdimensionaler Räume. Diese lag dem Bewußtsein der Zeitgenossen Cantors nahe und war durch informelle Sprechweisen lange vorbereitet. Cantor war jedoch der erste, der auf diesem Feld ein systematisches Studium eröffnete und zu bemerkenswerten und überraschenden Resultaten kam. Seine entsprechenden Ergebnisse wurden von seinen Kollegen rezipiert und benutzt. Ganz anders aber steht es mit der eigenständigen Theorie der transfiniten Kardinal- und Ordinalzahlen, die Hilbert später als „eine der kühnsten Schöpfungen des menschlichen Geistes“ bezeichnet hat. Dieser Teil seiner Theorie stieß am Anfang auf weitgehendes Unverständnis.

Vergegenwärtigen wir uns zunächst die ersten Schritte in der Herausbildung der Mengenlehre. Zu Cantors Biographie und zur Herausbildung seiner Ideen vergleiche man Dauben [1979] und Purkert und Ilgauds [1987]. 1872 publizierte Cantor eine Arbeit über die Eindeutigkeit der Darstellung einer Funktion durch eine Fourierreihe. Zur genauen Formulierung dieses Satzes benötigte er einen neuen Begriff, nämlich den einer „abgeleiteten Menge“, um gewisse Teilmengen der reellen Zahlen zu charakterisieren. Er definierte die Ableitung P' einer Teilmenge P der reellen Zahlen als die Menge aller Häufungspunkte von P . In derselben Arbeit gab er zudem eine Begründung des Systems der reellen Zahlen, die allein schon hingereicht hätte, um ihm einen Platz in der Geschichte der Mathematik zu sichern. Es folgte dann 1874 der Beweis, daß die algebraischen Zahlen abzählbar, die reellen Zahlen hingegen nicht abzählbar sind. Nach langen Bemühungen kam er schließlich 1877 zu dem für ihn selbst überraschenden Ergebnis, daß die Menge aller Punkte des Einheitsquadrats gleichmächtig ist zur Menge aller Punkte des Einheitsintervalls.

Obwohl einige dieser Ergebnisse sensationell waren, bewegten sie sich doch im Bereich der akzeptierten mengentheoretischen Beschreibung des Kontinuums. Ab 1879 begann Cantor dann mit der Publikation einer Serie von Artikeln unter dem Titel „Über lineare unendlich Punktmannichfaltigkeiten“ (Teil 1 bis 6), in der er den entscheidenden Schritt zur Theorie der transfiniten Ordinalzahlen vollzog. Dieser erfolgte in der 5. Abhandlung, die 1883 erschien (Cantor [1833]). Cantor leitete sie mit den Worten ein, daß seine Untersuchungen an einen Punkt gelangt seien, „wo ihre Fortführung von einer Erweiterung des realen ganzen Zahlbegriffs über die bisherigen Grenzen hinaus abhängig wird, und zwar fällt diese Erweiterung in eine Richtung, in welcher sie bisher von niemandem gesucht worden ist.“ (l.c., 73) Er ist sich wohl bewußt, daß sein Unternehmen in einem „gewissen Gegensatz zu weitverbreiteten Anschauungen über das mathematische Unendliche und zu häufig vertretenen Ansichten über das Wesen der Zahlgrößen“ (l.c.) steht. Daher widmet er den Rest dieser Arbeit historischen und philosophischen Erörterungen über das Unendliche und über seine neuen transfiniten Ordinalzahlen. In ihrem ganzen Charakter fällt diese Arbeit aus dem Rahmen des mathematisch Üblichen heraus, und es war der erstaunliche Weitblick eines Felix Klein nötig, damit sie überhaupt in einer mathematischen Zeitschrift erscheinen konnte.

Um Cantors epistemologisches Problem und die Kühnheit seines Schrittes zu verstehen, müssen wir etwas zurückgreifen. Cantor hatte die transfiniten Ordinalzahlen bereits

in der 2. Abhandlung (1880) eingeführt, und zwar als sogenannte „*Operationssymbole*“ (Cantor [1880]). Ist P eine Punktmenge und sind $P', P'', \dots, P^{(n)}, \dots$ ihre Ableitungen mit endlicher Ordnung, so erhält man die erste Ableitung mit unendlicher Ordnung und damit die erste unendliche Ordinalzahl ω gemäß $P^{(\omega)} = \bigcap_{n=1}^{\infty} P^{(n)}$. Auf $P^{(n)}$ folgen dann $P^{(\omega+1)}, P^{(\omega+2)}$ usw. Die ω -te Ableitung von $P^{(\omega)}$ ist dann durch $P^{(\omega \cdot 2)} = \bigcap_{n=1}^{\infty} P^{(\omega+n)}$ definiert. Diesen Prozeß kann man fortsetzen über alle Polynome in ω , bis zu Bildungen wie $P^{\omega^\omega}, P^{\omega^{\omega^\omega}}$ usw. Die für unseren Zusammenhang wichtige Tatsache ist nun, daß für alle diese Mengen tatsächlich *anschaulich-geometrische Beispiele* angegeben werden können, wie Cantor zeigt. Damit haben aber auch die transfiniten Operationssymbole $\omega, \omega^\omega, \omega^{\omega^\omega}$ einen wohlbestimmten anschaulichen Sinn. Um sich die Situation zu verdeutlichen, kann man vielleicht eine Analogie mit der Bruchrechnung machen. Brüche wie $\frac{3}{4}$ oder $\frac{1}{7}$ können als Operationssymbole verstanden werden, die aus einem Ganzen (z.B. einem Kuchen) Teile erzeugen. In diesem Sinne haben Brüche eine wohlbestimmte anschauliche Bedeutung. Etwas ganz anderes ist es aber, sie von dieser anschaulichen Bedeutung zu lösen und als autonome symbolische Objekte aufzufassen, für die man einen abstrakten Kalkül, die „Bruchrechnung“, aufbaut.

Genau diesen Schritt tut Cantor nun in der erwähnten 5. Abhandlung. Die transfiniten Operationssymbole (Cantor spricht von „*bestimmt definierten Unendlichkeitssymbolen*“) werden aus ihrem inhaltlich-anschaulichen Kontext gelöst und als „*unendliche reale, ganze Zahlen*“ aufgefaßt (Cantor [1883], 74). Wie wir gesehen haben, war sich Cantor der Kühnheit dieses Schrittes wohl bewußt, und im folgenden möchte ich einige der Argumente diskutieren, die er zu seiner Rechtfertigung anführt. Es ist nicht die Absicht, die mathematisch-philosophischen Auffassungen Cantors vollständig zu rekonstruieren. Vielmehr sollen seine Argumente klassifiziert und auf ihren Hintergrund im kulturell-geistigen Klima des 19. Jahrhunderts bezogen werden. Insgesamt kann man sieben Argumenttypen unterscheiden.

1 *Das wissenschaftstheoretische Argument: Mathematik als formale Wissenschaft*

Die entsprechenden Ausführungen Cantors gehören zu den berühmtesten und meist zitierten philosophischen Aussagen Cantors (l.c., 90/91). Er unterscheidet zwischen zwei Bedeutungen der Realität der ganzen (endlichen oder unendlichen) Zahlen, nämlich ihrer *intrasubjektiven* oder *immanenten* Realität, und ihrer *transsubjektiven* oder *transienten* Realität. Erstere besteht in den Beziehungen der Zahlen zu anderen „*Bestandteilen unseres Denkens*“, letztere beinhaltet, daß die Zahlen auch als Ausdruck oder Abbild von Vorgängen in der dem Intellekt gegenüberstehenden Außenwelt zu betrachten sind. Für die immanente Realität eines Begriffes ist es lediglich erforderlich, daß er in wohl definierten und widerspruchsfreien Beziehungen zu den anderen vorher gebildeten Begriffen steht. Die transiente Realität ist nur dann gegeben, wenn sich eine Beziehung des Begriffes auf Objekte der körperlichen und/oder geistigen Natur aufweisen läßt. Nach Cantor hat sich die reine Mathematik nur um die immanente Realität eines Begriffes zu kümmern und unterliegt keiner metaphysischen Kontrolle. In diesem Sinne ist sie „*freie Mathematik*“.

Dagegen muß die „*angewandte Mathematik*“ (analytische Mechanik und mathematische Physik) die Beziehung auf äußere Objekte berücksichtigen und ist insofern in ihren Grundlagen und Zielen metaphysisch.

Die strategische Funktion dieser Argumente ist klar. Cantor wollte sich angesichts der bereits erfolgten und noch zu erwartenden Kritik von Fachgenossen einen Freiraum verschaffen. Wären Gauß, Abel Jacobi, Dirichlet, Weierstraß, Hermite und Riemann einer metaphysischen Kontrolle unterworfen gewesen, so argumentiert er, niemals hätte sich die neuere (komplexe) Funktionentheorie, die ja zunächst keine Anwendungen hatte, entwickeln können. Cantor untermauerte diesen Gedankengang durch ein mathematikspezifisches Argument.

2 Ideale Elemente

Die transfiniten Ordinalzahlen und das Aktual-Unendliche seien genauso legitime und zugleich für den Alltagsverstand befremdliche Schöpfungen wie es einst die komplexen Zahlen gewesen seien (l.c., 74). Wir können feststellen, daß bereits Leibniz und andere ihren Gebrauch infinitesimal kleiner Größen durch denselben Hinweis gerechtfertigt hatten. Im 19. Jahrhundert war die Methode, neue anschaulich nicht interpretierbare Begriffe rein symbolisch einzuführen, von verschiedenen Mathematikern zu einer systematischen Technik ausgearbeitet worden, z.B. unendlich ferne Punkte, Geraden und Ebenen in der projektiven Geometrie, unendlich ferner Punkt in der komplexen Ebene der Funktionentheorie, die Einführung der idealen algebraischen Zahlen durch E.E. Kummer. Cantor wies in der vorliegenden Arbeit auf alle diese Beispiele hin, und in der Tat ist es für den heutigen Mathematiker, der mit diesen Techniken groß wird, schwer zu verstehen, daß Cantor einer derart starken Kritik ausgesetzt war. Dreißig Jahre später benutzte David Hilbert genau dasselbe Argument, um seine formale Begründung der Mathematik zu legitimieren. Hilbert betrachtete nur die endlichen kombinatorischen Objekte als anschaulich interpretierbar, der ganze Rest der Mathematik verwandelte sich für ihn in eine Universum „idealer Elemente“. Cantors Proklamierung der Freiheit der reinen Mathematik war daher, philosophisch gesehen, ein Schritt in die mathematische Moderne, der dennoch seine erkennbaren Wurzeln im kulturellen und philosophischen Klima des 19. Jahrhunderts hatte.

Diese Verwurzelung im 19. Jahrhundert zeigt sich auch daran, daß Cantor eben kein Formalist war und sich auch in Bezug auf die reine Mathematik mit einer formalen Sichtweise nicht zufrieden gab. Für ihn gabe es keinen Zweifel, daß sich die immanente und die transiente Realität stets zusammenfinden, und den Grund dafür sah er „*in der Einheit des Alls, zu welchem wir selbst mitgehören.*“ (l.c., 91) Unser Denken und die äußere Realität sind letztlich durch eine innere Harmonie verknüpft, wenn auch diese nicht jederzeit und in jeder Hinsicht offensichtlich ist. Um die inhaltliche Bedeutung der transfiniten Ordinalzahlen herauszuarbeiten, brachte Cantor daher eine ganze Reihe weiterer Argumente vor.

3 *Der Appell an die innere Anschauung*

An verschiedenen Stellen seines Textes appelliert Cantor an die „*innere Anschauung*“ als Beweis für die Existenz der transfiniten Ordinalzahlen. An einer Stelle sagt er, daß der Begriff der Anzahl in unserer inneren Anschauung eine unmittelbare gegenständliche Realität besitze und daß dies die Realität der Anzahlen auch in dem Fall beweise, wo sie bestimmt unendliche seien (l.c., 76). „Anzahl“ ist für Cantor das Ergebnis des Zählens, und der Kern seines Argumentes ist, daß auch die transfiniten Ordinalzahlen Ergebnis eines gewissen Zählprozesses sind. Wenig später stellt er fest, daß die Gesetze diese Anzahlen aus der „*unmittelbaren innere Anschauung*“ (l.c., 78) mit „*apodictischer Gewißheit*“ erschlossen werden können. Er erklärt sogar den Satz, daß jede Menge wohlgeordnet werden könne, zu einem allgemeingültigen Denkgesetz (l.c.). In seiner vollen Tragweite kann dieser Appell an die innere Anschauung nur im kulturellen Klima des 19. Jahrhunderts verstanden werden. Der Begriff selbst tritt wohl erstmalig bei Kant auf, wo er das psychologische Korrelat der reinen Anschauung ist. Er bezeichnet unsere Fähigkeit, unabhängig von jeder äußeren empirischen Anschauung räumliche und zeitliche Formen als das Material von Geometrie und Arithmetik innerlich zu erfassen. In der Folge wurde dieser Begriff dann ausgeweitet und mit dem der „*produktiven Einbildungskraft*“ zusammengebracht und bezeichnete dann die Fähigkeit, sich auch Objekte, die raum– zeitliche Anschauung übersteigen, innerlich vorzustellen. Innere Anschauung meinte also die Fähigkeit, sich Welten vorzustellen, die jenseits unserer empirischen Alltagswelt liegen. Der Begriff hatte auch eine wichtige pädagogische Funktion, insofern sich alle pädagogischen Autoren von Herbart bis Diesterweg einig waren, daß es im Unterricht der Mathematik nicht um die äußere Anschauung gehen können, sondern daß letztlich die innere Anschauung entwickelt werden müsse (Jahnke [1990], 34–62). Im ganzen 19. Jahrhundert spielte der Begriff in der Selbstdarstellung der Mathematik eine große Rolle, weil ideale Elemente, wie etwa die unendlich fernen Punkte der projektiven Geometrie, zwar nicht empirisch erfaßt werden können, aber doch der inneren Anschauung zugänglich erschienen, und in genau derselben Funktion benutzt ihn Cantor in der vorliegenden Arbeit. Von der Existenz und der Legitimationskraft einer solchen inneren Anschauung war er überzeugt und er konnte auch darauf vertrauen, daß das Argument für seine Leser Überzeugungskraft hatte.

4 *Das ästhetische Argument*

Die Schönheit einer Theorie war ein vielbenutzter Topos in der wissenschaftlichen Literatur des 19. Jahrhunderts, und so sprach auch Cantor davon, daß es ihm ein „wahrer Genuß“ sei zu sehen, wie der Zahlbegriff sich in zwei Begriffe (Mächtigkeit und Ordinalzahl) spalte, wenn man zum Unendlichen aufsteige (l.c.).

5 *Die Legitimität abstrakter Spekulation*

Kronecker, der erbitterte Gegner Cantors, vertrat bekanntlich die Auffassung, daß nur die endlichen ganzen Zahlen wirklich existieren und daß alle anderen analytischen Be-

griffe letztlich als Beziehungen zwischen endlichen ganzen Zahlen aufzufassen seien. Diese Konzeption, später als früher Intuitionismus bezeichnet, war gegen die meisten Analytiker seiner Zeit gerichtet, in besonderem Maße aber gegen Cantor. An einer sehr bezeichnenden Stelle unterstellt Cantor Kronecker nun, daß dahinter das Prinzip liege, „*den Flug der mathematischen Speculations- und Konzeptionslust in die wahren Grenzen zu weisen, wo sie keine Gefahr läuft, in den Abgrund des ‚Transcendenten‘ zu geraten.*“ (l.c., 82) Damit sind wir mitten in einer der wichtigsten und sehr deutschen Auseinandersetzungen in der Wissenschaft des 19. Jahrhunderts. Waren am Anfang des Jahrhunderts, etwa in der romantischen Naturforschung, Philosophie und Wissenschaft eng verzahnt, so trat später eine anti-philosophische Wende ein, die dazu führte, daß stärker positivistische Auffassungen herrschend wurden. Cantor sieht sich selber offenbar noch 1883 auf Seiten des spekulativen Denkens und im Gegensatz zur positivistischen Richtung. Er glaubt ein starkes Argument für sich zu haben, wenn er Kronecker vorwirft, mit seinen Auffassungen nur dem „Nützlichkeitsdenken“ anzuhängen. Schon C.G.J. Jacobi hatte gegen Fourier darauf insistiert, daß er die Wissenschaft „zur Ehre des menschlichen Geistes betreibe“. In diese Tradition reiht sich Cantor ein. In der Tat, trotz allem Positivismus, war es im deutschen akademischen Klima nicht opportun als Anhänger des „Nützlichkeitsdenkens“ zu gelten.

6 Organische Naturauffassung

Cantor sagt es nicht direkt, aber er legt dem Leser im Zusammenhang einer Bemerkung zur Philosophie von Leibniz und Spinoza die Idee nahe, daß mit den transfiniten Zahlen ein erster Schritt zu einer Mathematik getan werde, die für eine organische Naturerklärung geeignet wäre, während die bisherige mathematische Analyse der mechanischen Naturerklärung zur Seite stehe. Erneut konstatieren wir ein gedankliches Motiv der romantischen Naturforschung (l.c., 86).

7 Mathematische Verallgemeinerung und Differenzierung von Merkmalen

Gegen diejenigen Mathematiker und Philosophen, die in der Existenz des aktual-Unendlichen einen Widerspruch zu anderen Begriffen und Prinzipien der Mathematik sahen, argumentierte Cantor, daß es im Wesen der mathematischen Verallgemeinerung liege, daß die verallgemeinerten Objekte nicht mehr notwendig diejenigen Eigenschaften haben, wie die, von denen man ausgegangen ist. Für endliche Mengen fallen Mächtigkeit und Anzahl in dem Sinne zusammen, daß Mengen derselben Mächtigkeit auch dieselbe Anzahl haben. Dies stimmt nicht länger für unendliche Mengen. Mengen derselben Mächtigkeit können durchaus unterschiedliche Anzahlen haben. Der Begriff der Anzahl ist nämlich für unendlichen Mengen abhängig von der Ordnung, für endliche hingegen nicht. Im Endlichen gilt für Anzahlen das kommutative Gesetz der Addition, im Unendlichen hingegen nicht, denn $1 + \omega = \omega \neq \omega + 1$. Im Endlichen ist eine Zahl entweder gerade oder ungerade, in Unendlichen kann sie beides sein wegen $\omega \cdot 2 = 1 + \omega \cdot 2$. Es verhält sich hier wie bei

den komplexen Zahlen, wo man auch nicht mehr zwischen positiven und negativen Zahlen unterscheiden kann (l.c., 75, 77, 78). Lehrerstudenten sollten in der Lage sein, die Analogien zu sehen zu Problemen, die Lernende haben, wenn sie vom Rechnen mit natürlichen Zahlen zu gebrochenen oder negativen Zahlen übergehen. So gilt etwa nicht mehr das „Gesetz“, daß Multiplizieren eine Zahl vergrößert.

Versucht man eine Zusammenfassung dieser Argumente, so wird man die nicht als ein in sich geschlossenes Gedankengebäude empfinden können. Dieselbe innere Anschauung etwa, die Cantor seine transfiniten Ordinalzahlen „sehen“ ließ, diente ihm als Argument, um die unendlich kleinen Größen für mathematisch illegitim zu halten. Dennoch gibt es, wie wir gesehen haben, einen vielleicht überraschenden Bezugspunkt, der in der deutschen Kulturgeschichte des 19. Jahrhunderts liegt. Viele Argumente Cantors haben bereits 80 Jahre vorher in der Romantikistigen Außenwelt zu haben scheinen, und zum zweiten vertreten viele ihrer Professoren Auffassungen, die denen Cantors zumindest nahekommen. Im Studium gibt es allerdings im allgemeinen keine Gelegenheit, solche Auffassungen explizit zu besprechen, und auf die sachlichen und historischen Gründe für ihre Entstehung einzugehen.

Die Frage nach der Wirklichkeit und Existenz der mathematischen Begriffe ist das Thema dieses Textstücks, und es zeigt, wie sehr Cantor um seine Position gerungen hat. So vehement er die formale Auffassung vertreten hat, daß Widerspruchsfreiheit (immanente Realität) das einzige Kriterium für die Akzeptanz eines Begriffs in der reinen Mathematik sein sollte, so wenig konnte er sich damit völlig zufrieden geben. Für Studenten der Mathematik ist der Text daher sehr lesenswert. Sie werden hier mit der Tatsache konfrontiert, daß das Selbstverständnis der reinen Mathematik, mit dem sie es im Studium theoretisch und praktisch immer wieder zu tun haben, durchaus nicht immer so existiert hat, sondern daß es historisch unter ganz bestimmten Bedingungen entstanden ist. Wichtig und aufklärend ist es auch zu sehen, daß diese Position von Cantor unter großen persönlichen Opfern erkämpft worden ist, und daß er es sich dabei nicht leicht gemacht hat. Ihm gingen bestimmte Aussagen nicht so einfach über die Lippen wie vielen heutigen Mathematikern.

Hinter Cantors Ausführungen steht das *Spannungsverhältnis zwischen anschaulicher Interpretation und theoretischem Begriff*, ein Problem, das jeden Mathematiker betrifft. Genau so betrifft es die Lernenden und ihre Lehrer. Aber es gibt keine überzeitlichen Antworten auf diese Frage. Gerade deswegen ist die historische Sichtweise unerläßlich. Cantors Ausführungen enthalten eine faszinierende Mischung von Gesichtspunkten, die sich unmittelbar mit Erfahrungen der heutigen Lernenden verknüpfen lassen und anderen, die erst aus dem zeitgeschichtlichen Kontext erschlossen werden können. Selbst der romantische Grundzug in Cantors Denken ist unserer heutigen Problemlage nicht so fremd, wie es auf den ersten Blick erscheinen mag. Das große Thema der Romantik, sich in Welten hineinzudenken, die jenseits der Alltagswelt liegen, hat im Computerzeitalter mit ihren künstlichen Welten überraschende Aktualität gewonnen. So führt uns die historische Sicht nicht nur in die Vergangenheit zurück, sondern macht uns fähig, unsere heutigen Probleme differenzierter zu sehen.

Auch unter rein mathematischem Gesichtspunkt stellt die Artikelserie „Über unendliche lineare Punktmannichfaltigkeiten“, die Cantor zwischen 1879 und 1883 publizierte,

eine einzigartige Gemengelage von reeller Analysis (speziell Fourier-Analyse), Maß- und Integrationstheorie, mengentheoretischer Topologie und Mengenlehre dar. Die innere Einheit der Analysis kann wohl an keinem Dokument der Mathematikgeschichte so eindringlich aufgezeigt werden, wie in diesen Arbeiten Cantors.“

J.N. JAHNKE: *Mathematikgeschichte für Lehrer — Gründe und Beispiele*, Mathem. Semesterberichte (1996), S. 21–46.

Fazit Man kann die reellen Zahlen nicht in derselben Weise aufweisen wie die natürlichen, rationalen oder auch die algebraischen Zahlen. Bei den Dedekindschen Schnitten handelt es sich um eine neuartige Bildung mathematischer Objekte. Das Prinzip der kleinsten oberen Grenze sprengt den Rahmen der traditionellen mathematischen Konstruktionen und ist deswegen von vielen bedeutenden Mathematikern abgelehnt worden. Kronecker (1823–1891) hat nicht nur Cantor sondern auch seinen Berliner Kollegen Weierstraß (1815–1897) wegen dieser Schlußweisen heftig angegriffen. Die erkenntnistheoretischen Probleme sind durch Cantors Mengenlehre verdeutlicht und zugespitzt aber nicht bereinigt worden. Der entscheidende Wandel vollzog sich bei der Auffassung von der Existenz mathematischer Objekte. In einem Brief an G. Frege, der in seinen Büchern (1893–1904) die Mathematik in voller Ausführlichkeit mit reiner Logik zu begründen versucht, schreibt D. HILBERT 1899: „*Sie schreiben: Aus der Tatsache der Axiome folgt, daß sie einander nicht widersprechen. Es hat mich sehr interessiert, gerade diesen Satz bei Ihnen zu lesen, da ich nämlich solange ich über solche Dinge denke, schreibe und vortrage, immer wieder gerade umgekehrt sage: Wenn sich die willkürlich gesetzten Axiome nicht widersprechen, mit sämtlichen Folgen, so sind sie wahr, so existieren die durch die Axiome definierten Dinge.*“

Die Grundlagenkrise der Mathematik, von der man in den zwanziger Jahren viel sprach, gilt heutzutage als überwunden. Die heutigen Mathematiker haben sich damit abgefunden, daß die Analysis auf erkenntnistheoretisch schlecht fundiertem Boden blüht. Hilbert hat die, die sich dadurch (zeitweise) beirren haben lassen, immer wieder aufgerüttelt. Er meinte, die Erkenntnistheorie müsse sich eben der Mathematik anpassen. „*Fruchtbaren Begriffsbildungen und Schlußweisen wollen wir, wo immer nur die geringste Aussicht sich bietet, sorgfältig nachspüren und sie pflegen, stützen und gebrauchsfähig machen. Aus dem Paradies, das Cantor uns geschaffen hat, soll uns niemand vertreiben können.*“ (D. Hilbert: Über das Unendliche. Math. Ann. 95, (1925), 161–190).

P. BERNAYS, der Mitarbeiter Hilberts, hat 1930 in einer seiner philosophischen Abhandlungen erklärt: „*Es erscheint als eine unberechtigte Zumutung von seiten der Philosophie an die Mathematik, daß sie ihre einfachere und leistungsfähigere Methode zugunsten einer be-*

schwerlicheren und an Systematik zurückstehenden Methode aufgeben solle, ohne durch eine innere Nötigung dazu veranlaßt zu sein.“

Besonders eindrucksvoll hat H. WEYL (1885–1947) die Grundlagenkrise durchlitten. Seine Konklusionen beziehen auch Fragen der Anwendung in der Physik und allgemeine Vorstellungen von der geschichtlichen Entwicklung mit ein. Die theoretischen Begriffe sind für ihn „getragen von dem an uns sich vollziehenden Lebensprozeß des Geists und werden niemals als ein totes ‚endgültiges‘ Resultat abgesetzt“ (zitiert nach Jahnke, Math. Semesterberichte, Band 38 (1991), S. 175). Die Mathematik kann nicht mehr als ein Reich von ewigen Wahrheiten verstanden werden. Insbesondere hat die heutige Analysis mit den Motiven von Leibniz, dem Erfinder der Infinitesimalrechnung, nicht mehr viel zu tun. Der Aufsatz von HARDY GRANT: Leibniz – Beyond the calculus. Mathem. Semesterberichte Bd 39, 3 – 11, beginnt mit den Worten: *„Für Leibniz war die Mathematik eine Sammlung von notwendigen und ewigen Wahrheiten. Darüberhinaus glaubte er, daß bei der Erschaffung der Welt Gott durch quasi-mathematische Überlegungen geleitet war. . . .“*

Es wäre ein Erfolg, wenn eine Einführung in die Analysis den Bogen sichtbar machen könnte. Mit Jahnke möchte ich die These vertreten: Wissenschaftliche Theorien sind komplizierte selbstbezügliche Systeme und die Mathematik bringt in ständig wachsendem Maße ein spezifisches Element des Formalen, des Stimmigen und der Durchsichtigkeit in die Wirklichkeitsdeutung.

C.2 Aussagen

Die Mathematiker bedienen sich gerne, aber nicht systematisch der Notation der formellen Aussagenlogik. Die Rechenobjekte der Aussagenlogik sind Aussagen, die wahr oder falsch sein können, tertium non datur. Daß im „tertium non datur“ ein logisches Problem liegt, wird dem Anfänger gern verschwiegen. Die meisten Mathematiker glauben jedenfalls an das

Rechnen mit Aussagen

- 1) Jede Aussage von der betrachteten Art kann man negieren; mit A ist auch $\neg A$ („non A“) eine Aussage und es gilt

$$(i) \quad \neg(\neg A) = A$$

- 2) Beliebige Aussagen A und B kann man mit „Und“ und „Oder“ verknüpfen (das „Oder“ ist das „nichtausschließende Oder“; das Verknüpfungssymbol $\vee$ erinnert an das lateinische „vel“ im Gegensatz zum „ausschließenden Oder“, dem lateinischen „aut“). Es gilt

$$(ii) \quad A \vee B = B \vee A, \quad (A \vee B) \vee C = A \vee (B \vee C) = A \vee B \vee C.$$

Entsprechend gilt für die „Und“-Verknüpfung

$$(iii) \quad A \wedge B = B \wedge A, \quad (A \wedge B) \wedge C = A \wedge (B \wedge C) = A \wedge B \wedge C.$$

Ferner gelten die Distributivgesetze

$$(iv) \quad A \wedge (B \vee C) = (A \wedge B) \vee (A \wedge C)$$

$$A \vee (B \wedge C) = (A \vee B) \wedge (A \vee C)$$

und die sog. de Morganschen Regeln

$$(v) \quad \neg(A \wedge B) = (\neg A) \vee (\neg B)$$

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

- 3) Schließlich hat man als ausgezeichnete Aussagen die „immer falsche Aussage“ $\tilde{0}$ und die „immer richtige Aussagen“ $\tilde{1}$. Wir bemerken

$$(vi) \quad \tilde{0} \vee A = A, \quad \tilde{0} \wedge A = \tilde{0} \quad \text{für alle } A.$$

$$(vii) \quad A \wedge (\neg A) = \tilde{0}, \quad A \vee (\neg A) = \tilde{1} \quad \text{für alle } A.$$

Hinweis Eine strukturierte Menge

$$(\mathfrak{A}, =, \tilde{0}, \tilde{1}, \neg, \vee, \wedge)$$

mit den Eigenschaften (i) ... (vii) heißt eine Boolesche Algebra. (Wir haben uns nicht bemüht, möglichst viele überflüssige Axiome, d.h. solche, die aus anderen folgen, wegzulassen.) Die professionelle Aussagenlogik ist aber nicht genau dasselbe wie das Rechnen in einer Booleschen Algebra. Ein großes Problem liegt dort im Begriff der Gleichheit von Aussagen. Die Theorie der Booleschen Algebren schafft es durch die Axiome der Gleichheit aus der Welt.

Notation Manche Überlegungen werden suggestiver, wenn man $A \rightarrow B$ für $(\neg A) \vee B$ schreibt und „Wenn A , dann B “ liest. Die Verknüpfung $\rightarrow$ heißt Subjunktion. Der Pfeil ist die offizielle Bezeichnung in der Logik. Man beachte aber, daß der Pfeil in der Mathematik auch als Abbildungspfeil und als Konvergenzpfeil ausgiebig strapaziert wird.

Das Beweisen mathematischer Aussagen; Implikationen

Die Mathematiker sind ständig damit beschäftigt, von gewissen Aussagen zu beweisen, daß sie wahr sind. Die Behauptungen, die es zu beweisen gilt, sind meistens sogenannte Implikationen „ $A \Rightarrow B$ “; man will nachweisen, daß B aus A folgt, d.h. daß $A \rightarrow B$ wahr ist. Wenn A falsch ist, dann ist nichts zu beweisen; wenn aber A wahr ist, dann ist nachzuweisen, daß auch B wahr ist. Das Beweisen geschieht mit formalen Schlußregeln auf der Grundlage der Axiome.

Beim *direkten Beweis* von „ $A \Rightarrow B$ “ nimmt man an, daß A wahr ist und folgert daraus B . Beim *indirekten Beweis* nimmt man an, daß B falsch ist und folgert, daß auch A falsch ist. Beim *Widerspruchsbeweis* nimmt man an, daß A und $\neg B$ wahr sind; wenn es gelungen ist, aus dieser Annahme eine offensichtlich falsche Aussage, einen „Widerspruch“, abzuleiten, dann hat man „ $A \Rightarrow B$ “ durch die *reductio ad absurdum* bewiesen. Daß $A \wedge (\neg B)$ falsch ist, bedeutet in der Tat, daß $(\neg A) \vee B$ wahr ist.

Didaktische Anmerkung Widerspruchsbeweise sind unbeliebt; denn im Laufe der Deduktionen aus der falschen Annahme, daß A und $\neg B$ wahr sind, lernt man nichts nützliches über die wahren Sachverhalte.

Äquivalente Aussagen Mit „ $A \Leftrightarrow B$ “ (gelesen: A und B sind äquivalent) bezeichnet man die Behauptung, daß A genau dann wahr ist, wenn B wahr ist. „ $A \Leftrightarrow B$ “ behauptet

also die *Äquivalenz* der Aussagen A und B . Häufig entstehen Äquivalenzen durch Definitionen. Und die Herkunft einer Äquivalenz aus einer Definition anzudeuten, schreibt man $A :\Leftrightarrow B$ und liest: A ist definitionsgemäß genau dann wahr, wenn B wahr ist. Dabei ist B eine bereits wohlumrissene Aussage und A eine neugebildete Aussage.

Der Allquantor und der Existenzquantor Das Rechnen mit Aussagen wird erst interessant, wenn man Scharen von Aussagen zur Verfügung hat. Sei z.B. $A(x, y)$ eine Aussage für jedes Paar $(x, y) \in I \times J$, wobei I und J irgendwelche Indexmengen sind.

Mit $(\forall y \ A(x, y))$ oder $\left(\bigwedge_y A(x, y)\right)$ bezeichnet man diejenige Aussage $B(x)$, welche genau dann wahr ist, wenn $A(x, y)$ für alle y wahr ist.

Mit $(\exists y \ A(x, y))$ oder $\left(\bigvee_y A(x, y)\right)$ bezeichnet man die Aussage $C(x)$, welche genau dann wahr ist, wenn $A(x, y)$ für mindestens ein y wahr ist.

Die Aussage, die genau dann wahr ist, wenn $A(x, y)$ für alle (x, y) wahr ist, kann man auch mit $\forall x \ B(x)$ notieren. Interessanter sind Konstruktionen, in welchen sowohl Allquantoren als auch Existenzquantoren vorkommen, z.B.

$$\forall x \ \exists y \ A(x, y) .$$

Diese Aussage C ist genau dann wahr, wenn zu jedem x ein y existiert, so daß $A(x, y)$ wahr ist.

$$\exists y \ \forall x \ A(x, y) .$$

hingegen ist die Aussage D , die genau dann wahr ist, wenn ein y existiert, so daß $A(x, y)$ für alle x wahr ist.

Beispiele

- 1) Sei $\exp(\cdot)$ die bekannte Exponentialfunktion im Reellen

$$A(x, y) = (x \leq y \rightarrow e^x \leq e^y)$$

ist eine wahre Aussage $A(x, y)$ für jedes $(x, y) \in \mathbb{R} \times \mathbb{R}$.

Die Aussage $A = \forall (x, y) \ A(x, y)$ ist also eine wahre Aussage.

- 2) Sei $f(\cdot)$ irgendeine reellwertige Funktion auf $\mathbb{R}$.

$$(x \leq y \rightarrow f(x) \leq f(y))$$

kann für gewisse (x, y) wahr sein und für andere falsch.

Damit haben wir ein Beispiel für eine Äquivalenz

$$f(\cdot) \text{ isoton} \iff \forall x, y : (x \leq y \rightarrow f(x) \leq f(y)) .$$

3) Betrachte für $n = 1, 2, \dots$ die Aussage

$$\forall x \geq 0 \quad \left[\left(1 + \frac{x}{n}\right)^n \leq \left(1 + \frac{x}{n+1}\right)^{n+1} \right] .$$

Wir werden sehen, daß diese Aussage A_n wahr ist für jedes n , daß also

$$\forall n \in \mathbb{N} \forall x \geq 0 \quad \left(1 + \frac{x}{n}\right)^n \leq \left(1 + \frac{x}{n+1}\right)^{n+1}$$

eine wahre Aussage ist.

4) $f_1(\cdot), f_2(\cdot), \dots$ seien reellwertige Funktionen (auf einen Definitionsbereich D).

Was bedeutet die Aussage, daß die Funktionenfolge monoton ansteigend ist? Wir definieren

$$(f_n(\cdot))_n \text{ monoton ansteigend} : \iff \forall n \in \mathbb{N} \forall x \in D \quad (f_n(x) \leq f_{n+1}(x)) .$$

Die weiteren Beispiele führen Begriffe ein, die wir im weiteren Verlauf der Vorlesung entwickeln werden. Wir werden sehen, daß es sich um fruchtbare Begriffe handelt, Begriffe, mit denen man arbeiten kann. Hier kommt es uns zunächst nur darauf an, daß es sich (prima vista gesehen) um wohlgebildete Begriffe handelt.

5) $x_1, x_2, \dots$ sei eine Folge von reellen Zahlen

$$(x_n)_n \text{ Nullfolge} : \iff \forall \varepsilon > 0 \exists N \forall n \geq N \quad (|x_n| < \varepsilon)$$

Noch formeller könnte man schreiben

$$(a) \quad (x_n)_n \text{ Nullfolge} : \iff \forall \varepsilon > 0 \exists N \forall n \quad (n \geq N \rightarrow |x_n| < \varepsilon)$$

$$(b) \quad \lim x_n = x^* : \iff \forall \varepsilon > 0 \exists N \forall n \geq N \quad (|x_n - x^*| < \varepsilon)$$

$$(c) \quad (x_n)_n \text{ Cauchy-Folge} : \iff \forall \varepsilon > 0 \exists N \forall n, m \geq N \quad (|x_n - x_m| < \varepsilon) .$$

Negation von Aussagen mit Quantoren Für jedes x aus einer Indexmenge sei $B(x)$ eine Aussage; und es sei B die Aussage $\forall x \quad B(x)$

$\neg B$ ist die Aussage, die genau dann wahr ist, wenn B falsch ist.

$$\neg B = \neg(\forall x \quad B(x)) = \exists x \quad (\neg B(x)) .$$

B ist nämlich genau dann falsch, wenn es ein x gibt, für welches $B(x)$ falsch, d.h. $\neg B(x)$ wahr ist.

Entsprechend gilt

$$\neg(\exists x \quad B(x)) = \forall x \quad (\neg B(x)) .$$

$$6) (x_n)_n \text{ ist nicht Nullfolge : } \iff \exists \varepsilon > 0 \quad \forall N \quad \exists n \geq N \quad (|x_n| \geq \varepsilon) .$$

Bemerke

$$\begin{aligned} \neg(\forall n \geq N : |x_n| < \varepsilon) &= \neg(\forall n \quad (n \geq N) \rightarrow (|x_n| < \varepsilon)) \\ &= \exists n \quad (n \geq N) \wedge \neg(|x_n| < \varepsilon) \\ &= \exists n \geq N \quad (|x_n| \geq \varepsilon) . \end{aligned}$$

$$7) f(\cdot) \text{ ist nicht isoton } \iff \exists (x, y) \quad (x \leq y) \wedge (f(x) > f(y))$$

Für die Aussage

$$A(x, y) = (x \leq y) \rightarrow (f(x) \leq f(y)) = \neg(x \leq y) \vee (f(x) \leq f(y))$$

haben wir nämlich

$$\neg A(x, y) = (x \leq y) \wedge \neg(f(x) \leq f(y)) = (x \leq y) \wedge (f(x) > f(y)) .$$

Bemerkung zum Auswahlaxiom Für jedes $(x, y) \in I \times J$ sei $A(x, y)$ eine Aussage. Nehmen wir an, die Aussage

$$A = \forall x \in I \quad \exists y \in J \quad A(x, y)$$

sei wahr. Zu jedem $x \in I$ existiert dann mindestens ein y , so daß $A(x, y)$ wahr ist. Damit ist uns aber noch keine Methode gegeben, wie wir zu jedem $x \in I$ ein $y(x) \in J$ finden können, so daß

$$A(x, y(x)) \text{ wahr ist für alle } x .$$

Das *Auswahlaxiom* der Mengenlehre (von E. ZERMELO (1871–1953) als Axiom eingeführt), besagt nun aber, daß es eine solche Auswahlfunktion

$$I \ni x \longmapsto y(x) \in J$$

gibt, daß also

$$A \iff \exists y(\cdot) \quad \forall x \quad (A(x, y(x))) .$$

Das Auswahlaxiom in dieser scharfen Fassung, d.h. ohne alle Annahmen über $I \times J$ und über die Art der Familie $A(x, y)$ ist unter Logikern umstritten. Die meisten Mathematiker benützen es aber ziemlich ungeniert, wenn sie es brauchen, obwohl sie wissen, daß es in manchen Situationen zu ziemlich paradoxen Aussagen führt. Es gilt heute als eine Frage des mathematischen Geschmacks, wie man das Auswahlaxiom benützt und welche Bedeutung man den damit gewonnenen mathematischen Aussagen zumißt. Während der Grundlagenstreit bis in die 30er Jahre dieses Jahrhunderts teilweise erbittert geführt worden ist, koexistiert heute eine Logik der konstruktiven Beweise, die den Verfahrensaspekt der Mathematik pflegt, mit dem von Hilbert immer wieder vertretenen Erkenntnisoptimismus des freien Gebrauchs von nichtkonstruktiven Methoden.

Banachs Beispiel einer nichtmeßbaren Menge In der Menge $[0, 1)$ aller reellen Zahlen zwischen 0 und 1 führen wir eine Äquivalenzrelation ein

$$x_1 \sim x_2 : \Longleftrightarrow x_2 - x_1 \in \mathbb{Q}.$$

Wir betrachten die Menge $(E, =)$ der Äquivalenzklassen; jedes $X \in E$ ist also eine abzählbare nichtleere Menge reeller Zahlen. Die Aussage $(\forall X \in E \exists y \in [0, 1) \quad y \in X)$ ist wahr.

Denken wir uns eine Auswahlfunktion $y(X) : E \rightarrow [0, 1)$ und betrachten wir die Menge M der ausgewählten Elemente. Dieses M ist nun eine sehr merkwürdige überabzählbare Menge. Eine Merkwürdigkeit ist die folgende:

für r rational, $0 \leq r < 1$, betrachte

$$M_r := \{x : x - r \in M \text{ oder } x - r + 1 \in M\}.$$

Man schreibt $M_r = (M+r) \pmod{1}$. Die abzählbar vielen M_r sind disjunkt und überdecken $[0, 1)$; sie sind andererseits einander recht ähnlich. Daraus ergibt sich, daß ihnen kein Maß im Lebesgueschen Sinne zukommt. Solche Mengen M wurden von Banach als Beispiele für eine nichtborelmeßbare Menge angegeben. Es handelt sich aber natürlich nicht um eine wirkliche Konstruktion; niemand kann eine Auswahlfunktion $y(\cdot)$ wirklich angeben. Jedoch ist die Annahme, daß es solche Auswahlfunktionen gibt, nicht zu widerlegen. Cantor sagte „Das Wesen der Mathematik liegt in ihrer Freiheit“. Er spürte einen harten Druck durch den starken Existenzbegriff von Kronecker (Existenz = Konstruierbarkeit). Hilbert brachte seinen eigenen Existenzbegriff später auf die Formel „Existenz = Widerspruchsfreiheit“. (vgl. Purkert, insbesondere S. 174.)

Weitere Beispiele für die logische Notation Betrachten wir zweistellige Relationen R auf einer Menge E . In logischer Notation lauten die Definitionen von oben wie folgt

Definition

1) R Äquivalenzrelation: $\Longleftrightarrow$

$$\Longleftrightarrow (\forall x \quad xRx) \wedge \forall x, y, z \quad (xRz \wedge yRz \rightarrow xRy) .$$

2) R Ordnungsrelation: $\Longleftrightarrow$

$$\begin{aligned} \Longleftrightarrow (\forall x \quad xRx) \wedge \forall x, y, z \quad (xRy \wedge yRz \rightarrow xRz) \\ \wedge \forall x, y \quad (xRy \wedge yRx \rightarrow x = y) . \end{aligned}$$

3) R lineare Ordnung: $\Longleftrightarrow$

$$\Longleftrightarrow (R \text{ Ordnung}) \wedge \forall x, y \quad (xRy \vee yRx) .$$

Eine Abbildung $\varphi : E \rightarrow F$ wird durch ihren Funktionsgraphen

$$\Gamma = \{(x, y) : x \in E, y \in F, y = \varphi(x)\} \subseteq E \times F$$

beschrieben. Die folgende Bemerkung charakterisiert diejenigen Teilmengen von $E \times F$, die der Funktionsgraph (oder Abbildungsgraph) einer Abbildung sind.

Bemerkung Sei $\Gamma \subseteq E \times F$.

Γ Funktionsgraph $\Longleftrightarrow \forall x \in E \exists^1 y \in F \quad (x, y) \in \Gamma$

(„Zu jedem $x \in E$ existiert *genau ein* $y \in F$, so daß $(x, y) \in \Gamma$ “).

Ausführlich geschrieben

$$\begin{aligned} \Longleftrightarrow (\forall x \in E \exists y \in F \quad (x, y) \in \Gamma) \\ \wedge \forall x \in E \forall y_1, y_2 \in F \quad ((x, y_1) \in \Gamma \wedge (x, y_2) \in \Gamma) \rightarrow (y_1 = y_2) \end{aligned}$$

Definition Für eine Abbildung $\varphi : E \rightarrow F$ haben wir

1) φ injektiv $\Longleftrightarrow \forall x_1, x_2 \in E : \varphi(x_1) = \varphi(x_2) \rightarrow x_1 = x_2$

2) φ surjektiv $\Longleftrightarrow \forall y \in F \exists x \in E : \varphi(x) = y$

Bemerke

$$\varphi \text{ nicht injektiv} \iff \exists x_1, x_2 \in E : (x_1 \neq x_2) \wedge (\varphi(x_1) = \varphi(x_2))$$

$$\varphi \text{ nicht surjektiv} \iff \exists y \in F \forall x \in E : \varphi(x) \neq y .$$

Eine Abbildung heißt bekanntlich bijektiv, wenn sie sowohl injektiv als auch surjektiv ist. Mit Hilfe des Symbols $\exists^1$ ist das leicht kurz zu schreiben.

$$3) \varphi \text{ bijektiv} \iff \forall y \exists^1 x : \varphi(x) = y$$

Zu einer bijektiven Abbildung φ gibt es eine Umkehrabbildung ψ . Es gilt

$$\forall x : \psi \circ \varphi(x) = x \quad \text{und} \quad \forall y : \varphi(\psi(y)) = y .$$

Anwendungen der logischen Notation auf den Gegenstandsbereich des B–Zuges

Reelle Zahlen, Suprema von Zahlenmengen, obere und untere Limiten von Zahlenfolgen ...

1) „Die klassische Schlußweise der Analysis“

Um von einer nichtnegativen Zahl a zu beweisen, daß sie 0 ist, genügt es zu zeigen, daß sie kleiner ist als jede echtpositive (rationale) Zahl ε . Für eine reelle Zahl a gilt

$$\forall \varepsilon > 0 : a < \varepsilon \iff a \leq 0 .$$

Beispiel : Die Zahl $1,999\dots$ ist gleich 2 und nicht nur in irgendeinem Sinne beliebig nahe an 2. Für $a = 2,000\dots - 1.999\dots$ gilt nämlich $a < \varepsilon$ für alle $\varepsilon > 0$.

2) Sei A eine (nach oben und unten) beschränkte Menge von reellen Zahlen und b eine reelle Zahl. Es gilt

$$\sup A > b \iff \exists a \in A : a > b \quad \text{Beweis !}$$

Durch Negation erhält man

$$\sup A \leq b \iff \forall a \in A : a \leq b ,$$

Wenn man an die Spiegelung der Zahlenmenge am Nullpunkt denkt, dann ergibt sich

$$\inf A \geq b \iff \forall a \in A : a \geq b$$

$$\inf A < b \iff \exists a \in A : a > b .$$

3) Sei A ein (nach oben und unten beschränkte) Menge von reellen Zahlen und b eine reelle Zahl. Es gilt

$$\sup A \geq b \iff \forall \varepsilon > 0 \exists a \in A : a > b - \varepsilon \quad \text{Beweis !}$$

Denken Sie nach, ob Sie $\sup A \geq b$ auch auf andere (vielleicht einfachere) Weise zum Ausdruck bringen können.

$$\begin{aligned}\text{Vorschlag: } \sup A \geq b &\iff \forall c < b \quad \sup A > c \\ &\iff \forall c < b \quad \exists a \in A : a > c\end{aligned}$$

Durch Negieren bzw. Spiegeln erhalten wir

$$\begin{aligned}\sup A < b &\iff \exists \varepsilon > 0 \quad \forall a \in A : a \leq b - \varepsilon \\ \inf A > b &\iff \exists \varepsilon > 0 \quad \forall a \in A : a > b + \varepsilon \\ \inf A \leq b &\iff \forall \varepsilon > 0 \quad \exists a \in A : a \leq b + \varepsilon\end{aligned}$$

4) Sei $(a_n)_n$ eine Folge reeller Zahlen. Es gilt

$$\liminf a_n > b \iff \exists \varepsilon > 0 \exists N : \forall n \geq N \quad a_n \geq b + \varepsilon \quad \text{Beweis ! .}$$

„Schließlich alle a_n sind größer als $b + \varepsilon$ mit einem echt positiven ε .“

$$\liminf a_n \geq b \iff \forall \varepsilon > 0 \exists N \forall n \geq N \quad a_n > b - \varepsilon \quad \text{Beweis ! .}$$

„Schließlich alle a_n sind größer als $b - \varepsilon$, für beliebiges $\varepsilon > 0$.“

Bemerke : Bekanntlich gilt

$$\liminf a_n = \lim_N \uparrow \inf\{a_n : n \geq N\} = \lim_N \uparrow b_N .$$

$$\begin{aligned}\liminf a_n > b &\iff \lim \uparrow b_N > b \iff \exists \varepsilon > 0 \exists N : b_N \geq b + \varepsilon \\ &\iff \exists \varepsilon > 0 \exists N : \inf\{a_n : n \geq N\} \geq b + \varepsilon \\ &\iff \exists \varepsilon > 0 \exists N : \forall n \geq N \quad a_n \geq b + \varepsilon .\end{aligned}$$

Als Übung betrachte man die Aussagen

$$\begin{aligned}\liminf a_n &\leq b \\ \limsup a_n &\geq b \\ \limsup a_n &> b .\end{aligned}$$

C.3 Metrische Räume

Eine Menge $(E, =)$ wird zu einem metrischen Raum, wenn man eine Metrik $d(\cdot, \cdot)$ auszeichnet.

Definition Eine Metrik auf der Menge $(E, =)$ ist eine nichtnegative Funktion $d(\cdot, \cdot)$ auf $E \times E$ mit

- (i) $d(x, y) = 0 \iff x = y$ („Verträglichkeit mit der Gleichheit“)
- (ii) $d(x, y) = d(y, x)$ für alle x, y („Symmetrie“)
- (iii) $d(x, z) \leq d(x, y) + d(y, z)$ für alle x, y, z („Dreiecksungleichung“)

Beispiel Man denke an die komplexe Zahlenebene $\mathbb{C}$, vergesse alle algebraische Struktur und behalte nur die Struktur des Abstands

$$d(z, w) := |w - z|.$$

Wir erhalten einen vollständigen metrischen Raum. (Was vollständig heißt, werden wir unten definieren.)

Die Menge $\mathbb{C}_{\text{rat}}$ der komplexen Zahlen mit rationalem Real- und Imaginärteil ist eine abzählbare überall dichte Teilmenge dieses vollständigen metrischen Raums. (Was Dichtliegen heißt, werden wir unten definieren.)

Definition $(E, d(\cdot, \cdot))$ sei ein metrischer Raum. Eine Teilmenge $U \subseteq E$ heißt offen, wenn es zu jedem $x \in U$ eine Kugel mit dem Mittelpunkt x gibt, welche in U enthalten ist.

$$U \text{ offen} : \iff \forall x \in U \exists r > 0 : \forall y \ (d(x, y) < r) \rightarrow (y \in U).$$

Satz Das System $\mathfrak{U}$ aller offenen Mengen hat die Eigenschaften

- (i) $\emptyset \in \mathfrak{U}, \quad E \in \mathfrak{U}$
- (ii) $U_1, U_2 \in \mathfrak{U} \implies U_1 \cap U_2 \in \mathfrak{U}$
- (iii) Die Vereinigung von beliebig vielen Mengen aus $\mathfrak{U}$ gehört zu $\mathfrak{U}$.

Der Beweis ist trivial.

Sprechweise Die Metriken $d(\cdot, \cdot)$ und $\rho(\cdot, \cdot)$ auf $(E, =)$ heißen äquivalent, wenn sie zum gleichen System von offenen Mengen $\mathfrak{U}$ führen.

Beispiel Sei $(E, d(\cdot, \cdot))$ ein metrischer Raum und

$$\rho(x, y) = \frac{d(x, y)}{1 + d(x, y)} .$$

Dann ist $\rho(\cdot, \cdot)$ eine äquivalente Metrik.

Beweis

- (1) Zunächst ist zu zeigen, daß $\rho(\cdot, \cdot)$ eine Metrik ist. Die Eigenschaften (i) und (ii) sind offensichtlich. Mühe macht nur die Dreiecksungleichung.

Ihr Nachweis ist eine nette Übung in elementarer Analysis. Seien $x, y, z \in E$ und

$$a = d(x, y) , \quad b = d(y, z) , \quad c = d(x, z) .$$

Wir wissen $c \leq a + b$ und müssen beweisen

$$\frac{c}{1 + c} \leq \frac{a}{1 + a} + \frac{b}{1 + b} .$$

Es genügt, zu beweisen, daß

$$\frac{a + b}{1 + (a + b)} \leq \frac{a}{1 + a} + \frac{b}{1 + b} \quad \text{für alle } a, b > 0 .$$

Wir haben

$$\begin{aligned} & (a + b)(1 + a)(1 + b) - (a(1 + b) + b(1 + a))(1 + a + b) \\ &= (a + b)[1 + a + b + ab - (a + b + 2ab)] - (a + b + 2ab) \\ &= (a + b)(1 - ab) - (a + b) - 2ab \leq 0 . \end{aligned}$$

- (2) Wenn U offen ist bzgl. der Metrik $d(\cdot, \cdot)$, dann ist U auch offen bzgl. der Metrik $\rho(\cdot, \cdot)$ und umgekehrt. ■

Definition $(E, =)$ sei eine Menge. Unter einer E -wertigen *Folge* versteht man eine Abbildung

$$\mathbb{N} \rightarrow E .$$

Man notiert $(x_n)_n$ oder $(x_1, x_2, \dots)$ und nennt x_n das n -te *Folgenelement*.

Für späteren Gebrauch definieren wir hier auch gleich den Begriff der *Teilfolge* einer gegebenen Folge $(x_1, x_2, \dots)$. Wenn $n_1 < n_2 < \dots$ eine strikt aufsteigende Folge natürlicher Zahlen ist, dann heißt $(x_{n_1}, x_{n_2}, \dots)$ die dazugehörige Teilfolge.

Definition $(E, d(\cdot, \cdot))$ sei ein metrischer Raum. $(x_n)_n$ bezeichne eine E -wertige Folge, x^* ein Element von E . Wir definieren:

$(x_n)_n$ konvergiert gegen x^* : $\Longleftrightarrow$

$$\Longleftrightarrow \quad \forall \varepsilon > 0 \exists N \forall n \quad (n \geq N \rightarrow d(x_n, x^*) < \varepsilon) .$$

Man schreibt in einem solchen Fall auch

$$x^* = \lim_{n \rightarrow \infty} x_n \quad \text{oder auch} \quad (x_n \rightarrow x^* \text{ für } n \rightarrow \infty) .$$

(Der Konvergenzpfad $\rightarrow$ darf nicht mit dem Abbildungspfeil und auch nicht mit dem Subjunktionspfad verwechselt werden; die Kontexte sind so verschieden, daß eine Verwechslung nicht zu befürchten ist.)

Definition Eine Folge $(x_n)_n$ heißt eine *konvergente Folge*, wenn es ein x^* gibt, so daß $x_n \rightarrow x^*$ für $n \rightarrow \infty$.

Hinweis Konvergenzbeweise für Folgen haben häufig zwei Komponenten

- (i) $(x_n)_n$ ist eine konvergente Folge
- (ii) Identifikation des Limes

Beispiel

- 1) Die Zahlenfolge $(1 + \frac{1}{n})^n$ konvergiert in $\mathbb{R}$, weil sie monoton und beschränkt ist. Der Limes ist die Eulersche Zahl e .
- 2) Betrachte die rekursiv definierte Zahlenfolge

$$x_1 = \sqrt{2}, \quad \sqrt{2 + \sqrt{2}}, \dots, x_{n+1} = \sqrt{2 + x_n}, \dots$$

Die Folge $(x_n)_n$ konvergiert, weil sie monoton und beschränkt ist. Für den Limes x^* gilt

$$x^* = \sqrt{2 + x^*}, \quad \text{also} \quad x^* = 2 .$$

Daraus ergibt sich $x_n \rightarrow 2$ für $n \rightarrow \infty$.

Satz Für Folgen reeller Zahlen $(x_n)_n$ gilt

$$(x_n)_n \text{ konvergent} \iff \liminf_{n \rightarrow \infty} x_n = \limsup_{n \rightarrow \infty} x_n \in \mathbb{R} .$$

Beweis

- (1) Wir brauchen zunächst den Begriff des oberen und des unteren Limes einer Folge reeller Zahlen. Es ist bequem, auch die Werte $+\infty$ und $-\infty$ als mögliche Werte des limes superior und des limes inferior einer Zahlenfolge zuzulassen. Jede Zahlenfolge $(x_n)_n$ hat dann nämlich einen wohlbestimmten oberen und unteren Limes. Man definiert

$$\begin{aligned} \limsup x_n &= \lim_{n \rightarrow \infty} \searrow \sup\{x_m : m \geq n\} \\ \liminf x_n &= \lim_{n \rightarrow \infty} \nearrow \inf\{x_m : m \geq n\} \end{aligned}$$

mit der Maßgabe, daß $\limsup x_n = +\infty$, wenn $(x_n)_n$ nach oben unbeschränkt ist, $\liminf x_n = -\infty$, wenn $(x_n)_n$ nach unten unbeschränkt ist.

Beachte Für jede Zahlenfolge gilt

$$\liminf x_n \leq \limsup x_n .$$

- (2) Sei $(x_n)_n$ konvergent mit $\lim x_n = x^* \in \mathbb{R}$. Von einer gewissen Stelle $N = N(\varepsilon)$ an, liegen alle x_n in der ε -Umgebung von x^* . Also gilt

$$\begin{aligned} \limsup x_n &\leq \sup\{x_m : m \geq N(\varepsilon)\} \leq x^* + \varepsilon \\ \liminf x_n &\geq \inf\{x_m : m \geq N(\varepsilon)\} \geq x^* - \varepsilon . \end{aligned}$$

$\varepsilon > 0$ kann beliebig gewählt werden. Daher gilt

$$\limsup x_n \leq x^* \leq \liminf x_n .$$

- (3) Sei $\liminf x_n = x^* = \limsup x_n$. Zu jedem $\varepsilon > 0$ existiert $N(\varepsilon)$, so daß

$$\sup\{x_m : m \geq N(\varepsilon)\} < x^* + \varepsilon , \quad \inf\{x_m : m \geq N(\varepsilon)\} > x^* - \varepsilon$$

d.h. $-\varepsilon < x_m - x^* < \varepsilon$ für alle $m \geq N(\varepsilon)$. Dies zeigt $\lim x_n = x^*$.

Sprechweisen

- a) Von einer Folge reeller Zahlen $(x_n)_n$ sagt man, daß sie nach $+\infty$ konvergiert, wenn $\liminf x_n = +\infty$. In Formeln

$$\lim x_n = +\infty : \Longleftrightarrow \forall M \exists N \forall n \geq N \quad (x_n \geq M) .$$

Entsprechend

$$\lim x_n = -\infty : \Longleftrightarrow \forall M \exists N \forall n \geq N \quad (x_n \leq -M) .$$

- b) Manche Bücher nennen die Konvergenz gegen $+\infty$ ($-\infty$) die *bestimmte Divergenz* gegen $+\infty$ (bzw. $-\infty$). Wir werden diese Sprechweise nicht benutzen.

Wir definieren stattdessen für eine Zahlenfolge $(x_n)_n$

$$(x_n)_n \text{ divergent} \Longleftrightarrow \liminf x_n \neq \limsup x_n$$

- c) Von einer Folge komplexer Zahlen $(z_n)_n$ sagen wir, daß sie gegen ∞ konvergiert, wenn $(|z_n|)_n$ nach $+\infty$ konvergiert.

Hinweis Für eine Folge von Zahlen d -Tupeln bedeutet die Konvergenz, daß die Komponenten konvergieren. Das ist es, was man sich unter Konvergenz im $\mathbb{R}^d$ bzw. im $\mathbb{C}^d$ vorstellen sollte. Konvergenz gegen Unendlich hat für $d > 1$ keinen Sinn. Die Konvergenz in $\mathbb{R}^d$ bzw. $\mathbb{C}^d$ wollen wir nun aber im allgemeinen Rahmen der Konvergenz in einem metrischen Raum behandeln. Diese Behandlung soll auf allgemeinere Situationen vorbereiten wie z.B. auf den Begriff der gleichmäßigen Konvergenz einer Funktionenfolge.

Definition Seien $f_1(\cdot), f_2(\cdot), \dots$ beschränkte reellwertige Funktionen auf einem abstrakten Definitionsbereich Ω und $f^*(\cdot)$ ebenfalls beschränkt reellwertig. Man definiert

$$(f_n)_n \rightarrow f^* \text{ gleichmäßig auf } \Omega : \Longleftrightarrow$$

$$\Longleftrightarrow \forall \varepsilon > 0 \exists N \forall n \geq N \quad \sup\{|f_n(\omega) - f^*(\omega)| : \omega \in \Omega\} < \varepsilon .$$

Die Definition paßt in unseren allgemeinen Rahmen. Man betrachte nämlich auf einer Menge E von reellwertigen Funktionen auf Ω die Metrik

$$d(f, g) = \sup\{|g(\omega) - f(\omega)| : \omega \in \Omega\} .$$

(Den Nachweis, daß $d(\cdot, \cdot)$ wirklich eine Metrik ist, überlassen wir den Übungen). Gleichmäßige Konvergenz bedeutet Konvergenz bzgl. dieses Abstands in E ; man spricht auch von Konvergenz in der Supremumsnorm.

Beispiel Betrachte $f^*(\omega) = \sqrt{\omega}$ für $0 \leq \omega \leq M < \infty$. $\Omega = [0, M]$

$$f_1(\omega) = \frac{1}{2}(1 + \omega), \dots \quad f_{n+1}(\omega) = \frac{1}{2} \left(f_n(\omega) + \frac{\omega}{f_n(\omega)} \right), \dots$$

Man kann leicht zeigen, daß die Funktionenfolge $(f_n)_n$ gleichmäßig gegen f^* konvergiert.

Wir haben in B2 nur gezeigt, daß f_n punktweise (antiton) gegen f^* konvergiert. Die Zusatzüberlegungen, die erforderlich sind, um die gleichmäßige Konvergenz nachzuweisen, wollen wir an geeigneter Stelle explizieren.

Wir nehmen nun wieder den abstrakten Standpunkt ein. $(E, d(\cdot, \cdot))$ ist nichts weiter als ein metrischer Raum. Wer sich durchaus konkrete Vorstellungen machen will, denke an $\mathbb{C}$ oder $\mathbb{C}_{\text{rat}}$ mit dem üblichen Abstand. Etwas gefährlich ist es, sich $\mathbb{R}$ oder $\mathbb{Q}$ als Modellfall vorzustellen; die Anordnung in $\mathbb{R}$ kann allzuleicht nichtverallgemeinerbare Assoziationen wecken.

Generalvoraussetzung Die folgende Annahme über unsere metrischen Räume $(E, d(\cdot, \cdot))$ wird zwar zunächst kaum eine Rolle spielen. Sie erweist sich erst in diffizileren Überlegungen als angebracht. Wir setzen zwar nicht voraus, daß E abzählbar ist (dies würde ja $\mathbb{R}$ oder $\mathbb{C}$ oder $\mathbb{R}^d$ oder $\mathbb{C}^d$ ausschließen); wir nehmen aber an, daß es eine abzählbare Teilmenge $S \subseteq E$ gibt, welche in dem Sinne dicht in E liegt, daß jedes $x^* \in E$ als Limes einer S -wertigen Folge gewonnen werden kann. Die metrischen Räume $\mathbb{R}^d$ und $\mathbb{C}^d$ haben offenbar diese Eigenschaft, die in der allgemeinen Topologie als die Eigenschaft der Separabilität bezeichnet wird. In der gesamten Vorlesung bezeichnet $(E, d(\cdot, \cdot))$ eine separablen metrischen Raum.

Definition Für E -wertige Folge $(x_n)_n$ definieren wir

$$(x_n)_n \text{ Cauchy-Folge : } \iff \forall \varepsilon > 0 \exists N \forall n, m \geq N \quad d(x_n, x_m) < \varepsilon .$$

Satz Jede konvergente Folge ist Cauchy-Folge.

Beweis Wenn $\lim x_n = x^*$, dann existiert zu jedem $\varepsilon > 0$ ein N , so daß $d(x_n, x^*) < \frac{\varepsilon}{2}$ für alle $n \geq N$. Für alle $n, m \geq N$ gilt daher

$$d(x_n, x_m) \leq d(x_n, x^*) + d(x^*, x_m) < \varepsilon .$$

$(x_n)_n$ ist also Cauchy-Folge.

Bemerkung In $\mathbb{Q}$ (oder $\mathbb{C}_{\text{rat}}$) gibt es Cauchy-Folgen, die nicht konvergieren. In $\mathbb{R}$ (oder $\mathbb{C}$) ist jede Cauchy-Folge konvergent, wie wir sehen werden.

Definition Ein metrischer Raum $(E, d(\cdot, \cdot))$ heißt vollständig, wenn jede Cauchy-Folge konvergiert.

Satz Die reelle Achse $\mathbb{R}$ mit dem üblichen Abstand ist vollständig.

Beweis Sei $(x_n)_n$ eine Cauchy-Folge reeller Zahlen. Wir finden eine reelle Zahl x^* , gegen welche die Folge konvergiert. Es genügt zu zeigen

$$x^* = \limsup_{n \rightarrow \infty} x_n = \liminf_{n \rightarrow \infty} x_n$$

(siehe den obigen Satz). Dies ist eine leichte Übung.

Korollar Die komplexe Ebene $\mathbb{C}$ mit dem üblichen Abstand ist vollständig.

Beweis Eine Folge komplexer Zahlen $z_n = x_n + iy_n$ ist genau dann Cauchy-Folge, wenn die Realteile x_n und die Imaginärteile y_n Cauchy-Folgen sind.

Ergänzung :**Die Vervollständigung eines metrischen Raums**

$(E, d(\cdot, \cdot))$ sei ein metrischer Raum. (Man denke etwa an $E = \mathbb{C}_{\text{rat}}$.) Wir nennen zwei Cauchy-Folgen $X = (x_n)_n$ und $Y = (y_n)_n$ äquivalent, wenn sie sich schließlich beliebig wenig unterscheiden. Formell

$$X \sim Y : \Longleftrightarrow \forall \varepsilon > 0 \exists N \forall n, m \geq N \quad d(x_n, y_m) < \varepsilon .$$

Satz Es handelt sich wirklich um eine Äquivalenzrelation.

Beweis

- 1) Wir bemerken zunächst $X \sim X$ für jede Cauchy-Folge.
- 2) Aus $X \sim Z$, $Y \sim Z$ folgern wir $X \sim Y$.

Wähle N so groß, daß

$$\begin{aligned} \forall n, k \geq N \quad d(x_n, z_k) &< \varepsilon/2 \\ \forall m, k \geq N \quad d(y_m, z_k) &< \varepsilon/2 \end{aligned}$$

Es gilt dann $\forall n, m \geq N \quad d(x_n, y_m) < \varepsilon$. Damit haben wir $X \sim Y$.

Satz Sei $\overline{E}$ die Menge alle Äquivalenzklassen von Cauchy-Folgen. Setze

$$\overline{d}(X, Y) = \limsup_{n, m \rightarrow \infty} d(x_n, y_m) = \lim_{N \rightarrow \infty} \searrow \sup\{d(x_n, y_m) : n, m \geq N\} .$$

Dann ist $(\overline{E}, \overline{d}(\cdot, \cdot))$ ein vollständiger metrischer Raum.

Beweis

- 1) Zunächst zeigen wir, daß $\overline{d}(\cdot, \cdot)$ wohldefiniert ist. Für jedes $x^* \in E$ bilden die gegen x^* konvergierenden Folgen eine Äquivalenzklasse $\overline{x^*} \in \overline{E}$.

Für jede Cauchy-Folge $(y_m)_m$ und jedes x^* ist die Zahlenfolge $(d(x^*, y_m))_m$ eine Cauchy-Folge. Ihr Grenzwert ändert sich nicht, wenn man von $(y_m)_m$ zu einer äquivalenten Cauchy-Folge übergeht. Wir schreiben

$$\lim_{m \rightarrow \infty} d(x^*, y_m) = \overline{d}(\overline{x^*}, Y) .$$

Für jede gegen x^* konvergierende Folge $(x_n)_n$ gilt

$$\forall \varepsilon > 0 \exists N \forall n, m \geq N \left| d(x_n, y_m) - \bar{d}(\bar{x}^*, Y) \right| < \varepsilon .$$

Für jede Cauchy-Folge $(x_n)_n$ ist $(\bar{d}(\bar{x}_n, Y))_n$ eine Cauchy-Folge. Ihr Grenzwert ändert sich nicht, wenn man von $(x_n)_n$ zu einer äquivalenten Cauchy-Folge übergeht; dieser Grenzwert wird mit $\bar{d}(X, Y)$ bezeichnet. Es gilt

$$\forall \varepsilon > 0 \exists N \forall n, m \geq N \left| d(x_n, y_m) - \bar{d}(X, Y) \right| < \varepsilon .$$

Dann ist gezeigt, daß $\bar{d}(X, Y)$ wohldefiniert ist mit

$$\bar{d}(\bar{x}^*, \bar{y}^*) = d(x^*, y^*) \quad \text{für alle } x^*, y^* \in E .$$

- 2) Wir zeigen, daß $\bar{d}(\cdot, \cdot)$ eine Metrik auf $\bar{E}$ ist. Offenbar gilt $\bar{d}(X, Y) = 0$ genau dann, wenn $X \sim Y$. Für alle X, Y gilt $\bar{d}(X, Y) = \bar{d}(Y, X)$.

Die Dreiecksungleichung für $\bar{d}(\cdot, \cdot)$ ergibt sich folgendermaßen:

Sei $(x_n)_n$ ein Repräsentant von X , $(y_m)_m$ ein Repräsentant von Y und $(z_k)_k$ ein Repräsentant von Z .

Zu $\varepsilon > 0$ wählen wir N so groß, daß für alle $n, m, k \geq N$

$$\begin{aligned} d(x_n, y_m) &\leq \bar{d}(X, Y) + \frac{\varepsilon}{2} \\ d(y_m, z_k) &\leq \bar{d}(Y, Z) + \frac{\varepsilon}{2} . \end{aligned}$$

Dann gilt für alle $n \geq N$ nach der Dreiecksungleichung für $d(\cdot, \cdot)$

$$d(x_n, z_k) \leq \bar{d}(X, Y) + \frac{\varepsilon}{2} + \bar{d}(Y, Z) + \varepsilon 2 .$$

Dies beweist

$$\bar{d}(X, Z) \leq \bar{d}(X, Y) + \bar{d}(Y, Z) + \varepsilon .$$

Da die Ungleichung für alle $\varepsilon > 0$ gilt, haben wir die Dreiecksungleichung für $\bar{d}(\cdot, \cdot)$.

- 3) Wir zeigen die Vollständigkeit von $(\bar{E}, \bar{d}(\cdot, \cdot))$. Sei $(X_n)_n$ eine Cauchy-Folge in $\bar{E}$ mit X_n repräsentiert durch $(x_{n,m})_m$. Wählen wir $m(n)$ so groß, daß $\bar{d}(x_{n,m(n)}, X_n) < \frac{1}{n}$, dann liefert

$$(\tilde{x}_n)_n = (x_{n,m(n)})_n$$

eine Cauchy-Folge. Für das durch sie repräsentierte $\tilde{X} \in \overline{E}$ gilt

$$\lim_{n \rightarrow \infty} \overline{d}(X_n, \tilde{X}) = 0.$$

$\tilde{X}$ ist also der Grenzwert von $(X_n)_n$. $(\overline{E}, \overline{d}(\cdot, \cdot))$ ist vollständig.

Damit ist der Satz bewiesen.

Wir haben auch bewiesen, daß die Abbildung

$$E \ni x^* \longmapsto \overline{x}^* \in \overline{E}$$

injektiv ist.

Den Satz von der Vervollständigung eines metrischen Raums formulieren wir nun rein begrifflich, d.h. ohne Verweis auf die Konstruktion, als einen Existenzsatz.

Satz Zu jedem metrischen Raum $(E, d(\cdot, \cdot))$ gibt es eine abstandserhaltende Abbildung in einen vollständigen metrischen Raum, in welchem das Bild von E dicht liegt. Dieser vollständige metrische Raum ist bis auf Isomorphie eindeutig bestimmt.

$$i : (E, d(\cdot, \cdot)) \longrightarrow (\overline{E}, \overline{d}(\cdot, \cdot))$$

$$1) \quad \overline{d}(i(x), i(y)) = d(x, y) \quad \text{für alle } x, y$$

$$2) \quad \text{Zu jedem } X \in \overline{E} \text{ gibt es } E\text{-wertige Cauchy-Folgen } (x_n)_n \text{ mit}$$

$$\lim_{n \rightarrow \infty} i(x_n) = X.$$

Beispiele und Hinweise Bei der Konstruktion der Vervollständigung spielt es keine Rolle, welche Art von Objekten wir zu einer Menge $(D, =)$ zusammengefaßt und mit einem Abstand versehen haben. Wir kennen auch noch nicht viele Mengen interessanter mathematischer Objekte; der Schatz der Beispiele ist daher klein. Wenn wir ein Gefühl bekommen wollen, wie weitreichend der Vervollständigungssatz genutzt werden kann, müssen wir also etwas vorausschauen:

- 1) Sei $(\mathbb{Q}, |\cdot|)$ die Menge der rationalen Zahlen mit dem Abstand $d(r', r'') = |r'' - r'|$. Die Vervollständigung ergibt die Menge der reellen Zahlen mit dem Abstand $d(x, y) = |y - x|$. Viele Lehrbücher benützen das Vervollständigungsargument zur „Konstruktion“ der reellen Zahlen. Während die Dedekindsche „Konstruktion“ die Ordnung auf $\mathbb{Q}$ zum Ausgangspunkt nimmt, wird in dieser Konstruktion die Metrik auf $\mathbb{Q}$ ins Zentrum der Konstruktion der reellen Zahlen gestellt.

- 2) Sei $\mathbb{C}_{\text{rat}}$ die Menge der komplexen Zahlen mit rationalem Real- und Imaginärteil. Die Vervollständigung ergibt die Menge $\mathbb{C}$ aller komplexen Zahlen mit dem üblichen Abstand $d(z, w) = |w - z|$.
- 3) Fassen wir die Polynome mit reellen Koeffizienten zu einer Menge $\mathfrak{P}$ zusammen und definieren wir den Abstand durch die Supremumsnorm über $[0, 1]$

$$d(p, q) = \sup\{|q(x) - p(x)| : x \in [0, 1]\} .$$

Es ist eine bemerkenswerte Tatsache, daß man die Punkte der Vervollständigung als Funktionen über dem Einheitsintervall verstehen kann mit

$$\overline{d}(f, g) = \sup\{|g(x) - f(x)| : x \in [0, 1]\} .$$

Die Punkte der Vervollständigung sind diejenigen $f(\cdot)$, die man gleichmäßig durch Polynome approximieren kann

$$f \in \mathfrak{P} \iff \forall \varepsilon > 0 \exists p(\cdot) \in \mathfrak{P} \sup\{|f(x) - p(x)| : x \in [0, 1]\} < \varepsilon .$$

Ein berühmter Satz von Weierstraß besagt, daß man jede *stetige Funktion* über $[0, 1]$ gleichmäßig durch Polynome approximieren kann.

Auf der anderen Seite ist jeder gleichmäßige Limes stetiger Funktionen stetig, wie wir im nächsten Kapitel beweisen werden. Somit ist die Vervollständigung $\overline{\mathfrak{P}}$ der metrische Raum $C([0, 1])$ aller stetigen Funktionen über $[0, 1]$.

Hinweis Eine beliebte Übungsaufgaben in der elementaren Wahrscheinlichkeitstheorie ist der Beweis, daß die folgende Folge von Polynomen $(p_n)_n$ für jedes stetige $f(\cdot)$ auf $[0, 1]$ gleichmäßig gegen $f(\cdot)$ konvergiert

$$p_n(x) = \sum_{k=0}^n f\left(\frac{k}{n}\right) \binom{n}{k} x^k (1-x)^{n-k}$$

(„Bernsteinsche Polynome“)

- 4) Trigonometrische Polynome sind spezielle 2π -periodische Funktionen, nämlich die Funktionen der Gestalt

$$f(x) = \frac{a_0}{2} + \sum_{k=1}^n (a_k \cos kx + b_k \sin kx) , \quad n \text{ beliebig} .$$

$F(\cdot)$ ist reellwertig, wenn die Koeffizienten a_k, b_k reell sind. Man betrachtet aber auch den Fall komplexer Koeffizienten. In diesem Falle schreibt man bequemer

$$f(x) = \sum_{k=-n}^{+n} c_k e^{ikx} \quad \text{mit} \quad c_k = \frac{1}{2} (a_k - ib_k) .$$

Es gibt mehrere interessante Abstandsbegriffe auf der Menge $\mathfrak{T}$ der trigonometrischen Polynome. Die wichtigsten sind die Supremumsnorm und die L_2 -Norm. Man definiert

$$\begin{aligned} d_\infty(f, g) &= \sup\{|g(x) - f(x)| : x \in \mathbb{R}\} \\ &= \sup\{|g(x) - f(x)| : x \in [0, 2\pi]\} \\ d_2(f, g) &= \left(\int_0^{2\pi} |f(x) - g(x)|^2 dx \right)^{1/2} = \|g - f\|_2 \end{aligned}$$

Bemerke Der Abstand von f zu g ist der Abstand von $g - f$ zur Nullfunktion.

Es stellt sich heraus, daß die Vervollständigung von $(\mathfrak{T}, d_\infty(\cdot, \cdot))$ die Menge aller stetigen 2π -periodischen Funktionen ist. Den Beweis gewinnt man üblicherweise als einen Spezialfall des berühmten Satzes von Stone–Weierstraß, der im nächsten Semester bewiesen werden soll.

Etwas anders steht es um die Vervollständigung von $(\mathfrak{T}, d_2(\cdot, \cdot))$. Die Elemente der Vervollständigung kann man nicht als Funktionen interpretieren, wohl aber als Äquivalenzklassen 2π -periodischer Funktionen. (Dies ist der berühmte Satz von Fischer–Riesz.)

Hierbei heißen zwei Funktionen äquivalent, wenn sie sich nur auf einer Lebesgueschen Nullmenge unterscheiden. Der Begriff der Lebesgueschen Nullmenge wird in der Integrationstheorie im nächsten Semester eingehend diskutiert werden.

Man kann die Punkte der Vervollständigung von $(\mathfrak{T}, d_2(\cdot, \cdot))$ auch noch auf eine andere Weise charakterisieren. Zunächst charakterisiert man eine trigonometrisches Polynom nicht durch eine 2π -periodische Funktion sondern durch die Folge der Koeffizienten

$$f \longleftrightarrow (\dots, c_{-2}, c_{-1}, c_0, c_1, c_2, \dots) .$$

Bei einem trigonometrischen Polynom sind nur endlich viele der c_n von Null verschieden. Den Abstand von f zur Nullfunktion kann man aus der Koeffizientenfolge bequem ablesen. Es gilt

$$d_2(f, 0) = \|f\|_2 = \left(\sum |c_n|^2 \right)^{1/2} .$$

Der Beweis ist einfach und wird später hergeleitet. Es stellt sich heraus, daß man die Punkte der Vervollständigung von $(\mathfrak{T}, d_2(\cdot, \cdot))$ durch Zahlenfolgen $(\dots, d_{-2}, d_{-1}, d_0, d_1, d_2, \dots)$ charakterisieren kann, wobei $\sum_{-\infty}^{\infty} |d_n|^2 < \infty$.

Der Beweis ist einfach, paßt aber nicht recht in den gegenwärtigen Zusammenhang.

C.4 Der Begriff der Stetigkeit

Wir betrachten Abbildungen von einem metrischen Raum in einen weiteren metrischen Raum

$$\varphi : (D, d(\cdot, \cdot)) \longrightarrow (E, e(\cdot, \cdot))$$

und definieren

$$\varphi \text{ stetig in } x^* : \Longleftrightarrow \quad \forall \varepsilon > 0 \exists \delta > 0 \forall x (d(x, x^*) < \delta \rightarrow e(\varphi(x), \varphi, x^*) < \varepsilon)$$

$$\varphi \text{ stetig in } D : \Longleftrightarrow \quad \forall x^* (\varphi \text{ stetig in } x^*)$$

Geschichtliches Der Begriff der Stetigkeit hat eine lange Tradition. B. BOLZANO hat ihn 1817 als erster scharf gefaßt. Er dachte an reellwertige Funktionen auf einem Intervall. Er kritisierte die damals übliche Verbindung des Stetigkeitsbegriffs mit der Vorstellung von einer Bewegung und schreibt (Übersetzung in Ostwalds Klassiker 1905):

„ Nach einer richtigen Erklärung nähmlich versteht man unter der Redensart, daß eine Funktion $f(x)$ für alle Werte von x , die inner- oder außerhalb gewisser Grenzen liegen, nach dem Gesetze der Stetigkeit sich ändere, nur so viel, daß wenn x irgend ein solcher Werth ist, der Unterschied $f(x + \omega) - f(x)$ kleiner als jede gegebene Größe gemacht werden könne, wenn man ω so klein als man nur immer will, annehmen kann. “

Cauchy übernahm die Definition in seinen „Cours d’analyse“ von 1821, ohne Bolzano zu zitieren. Cauchy entwickelte eine umfassende Theorie, die aber an vielen Stellen Schwächen zeigte. Einige knappe Ausführungen dazu finden sich z.B. in Walter, Analysis 2, S. 39/40 und Heuser, Lehrbuch der Analysis 2, S. 694.

Eine präzise Vorstellung von Stetigkeit hat sich im 19. Jahrhundert recht langsam durchgesetzt. Noch 1864 führte der hochangesehene JOSEPH BERTRAND in seinem Lehrbuch über Differential- und Integralrechnung einen (natürlich falschen) Beweis der Aussage, daß jede stetige Funktion einer reellen Veränderlichen mit Ausnahme isoliert liegender Stellen überall differenzierbar ist. (Vgl. R. Bölling „Karl Weierstraß — Stationen eines Lebens“, J. ber. d. Dt. Math.-Verein. 1994, Teubner-Verlag.) Weierstraß erregte Erstaunen, Unglauben und sogar Abscheu, als er 1872 in der Berliner Akademie einen Vortrag über eine Funktion hielt, die für alle reellen Zahlen definiert und stetig ist, aber an keiner Stelle differenziert werden kann. Noch zwanzig Jahre später klagte Hermite (in einem Brief an Th. Stieltjes vom 20.5.1893):

„ Aber diese so eleganten Entwicklungen sind mit einem Fluch belegt; [...] Die Analysis nimmt mit der einen Hand zurück, was die andere gibt. Ich wende mich mit Abscheu und Schrecken von dieser beklagenswerten Wunde der stetigen nirgendsdifferenzierbaren Funktionen ab [...]“
(zitiert nach Bölling)

Die auf absolute Strenge bedachten Konstruktionen und Deduktionen von Weierstraß führten weg von einem in vielen Bereichen erfolgreichen heuristischen „Prinzip der Stetigkeit“, welches der Geometer V. PONCELET (1788–1867) folgendermaßen formuliert hat:

„ Wenn eine Figur aus einer anderen durch eine stetige Veränderung hervorgeht und ebenso allgemein ist wie die erste, dann kann eine für die erste Figur bewiesene Eigenschaft ohne erneute Untersuchung auf die andere übertragen werden.“
(zitiert nach Struik, S. 191)

Um die Stetigkeit zu dem heute geläufigen Grundbegriff der Analysis zu entwickeln, bedurfte es flankierender Begriffe wie z.B.

gleichmäßige Konvergenz einer Funktionenfolge

gleichmäßige Stetigkeit einer Funktion (bzw. Abbildung)

gleichgradige Stetigkeit einer Familie von Funktionen.

Ferner mußte der Begriff der Kompaktheit herausgebildet werden. Schließlich mußte verstanden werden, für welche Folgen stetiger Funktionen eine konvergente Teilfolge existiert. Hier erwies sich der Begriff der gleichgradigen Stetigkeit einer Familie von Funktionen als Schlüsselbegriff (vgl. Satz von Arzelà–Ascoli in Analysis II). Die „Weierstraßsche Strenge“ wurde nicht von heute auf morgen geboren, und dem Programm von Weierstraß stellten sich mächtige Gegner in den Weg (siehe Bölling in DMV Jahresberichte 1994). Es mußten einige Jahrzehnte ins Land gehen, bis D. Hilbert in den Mathematischen Annalen von 1926 feststellen konnte: „*Wenn heute in Verfolgung von Schlußweisen, die auf dem Begriff der Irrationalzahl und überhaupt des Limes beruhen, in der Analysis volle Übereinstimmung und Sicherheit herrscht und in den verwickeltsten Fragen, die die Theorie der Differential- und Integralrechnung betreffen, trotz der kühnsten und mannigfaltigsten Kombinationen unter Anwendung von Über- Neben- und Durcheinanderhäufung der Limites doch Einhelligkeit aller Ergebnisse statthat, so ist das wesentlich ein Verdienst der wissenschaftlichen Tätigkeit von Weierstraß.*“

Gleichmäßige Limiten stetiger Abbildungen sind stetig

Satz Wir betrachten E -wertige stetige Funktionen auf $(D, d(\cdot, \cdot))$.

- a) Wenn eine Folge E -wertiger Abbildungen $\varphi_1, \varphi_2, \dots$ punktweise konvergiert, dann ist die Limesfunktion φ^* im allgemeinen nicht stetig.
- b) Wenn die Funktionenfolge gleichmäßig konvergiert, dann ist die Grenzfunktion stetig.

Beweis

- 1) a) wird durch ein Beispiel erledigt

$$f_n(x) := \frac{x^2}{1+x^2} + \frac{x^2}{(1+x^2)^2} + \dots + \frac{x^2}{(1+x^2)^n} \quad \text{für } |x| \leq 1.$$

Es gilt $f_n(0) = 0$ für alle n . Für $x \neq 0$ erhalten wir durch Summation der geometrischen Reihe

$$f_n(x) = x^2 \cdot \frac{1}{1+x^2} \left(1 - \frac{1}{1+x^2}\right)^{-1} \left[1 - \frac{1}{(1+x^2)^n}\right] \nearrow 1$$

$$f^*(x) = \begin{cases} 0 & \text{für } x = 0 \\ 1 & \text{für } x \neq 0. \end{cases}$$

Der punktweise Grenzwert ist also unstetig.

- 2) $\varphi_n \rightarrow \varphi^*$ gleichmäßig auf $D \iff$

$$\iff \forall \varepsilon > 0 \exists N \forall n \geq N \sup\{e(\varphi_n(x), \varphi^*(x)) : x \in D\} < \varepsilon.$$

Wir zeigen die Stetigkeit von φ^* im Punkte $\tilde{x} \in D$, indem wir zu $\varepsilon > 0$ ein $\delta > 0$ angeben mit

$$\forall x (d(x, \tilde{x}) < \delta \longrightarrow e(\varphi^*(x), \varphi^*(\tilde{x})) < \varepsilon).$$

Dazu wählen wir N so groß, daß

$$\sup\{e(\varphi_N(x), \varphi^*(x)) : x \in D\} < \frac{\varepsilon}{3}$$

und δ so klein, daß

$$\forall x : (d(x, \tilde{x}) < \delta \longrightarrow e(\varphi_n(x), \varphi_N(\tilde{x})) < \frac{\varepsilon}{3}).$$

Wir haben dann für alle x mit $d(x, \tilde{x}) < \delta$

$$\begin{aligned} e(\varphi^*(x), \varphi^*(\tilde{x})) &\leq e(\varphi^*(x), \varphi_N(x)) + e(\varphi_N(x), \varphi_N(\tilde{x})) + e(\varphi_N(\tilde{x}), \varphi^*(\tilde{x})) \\ &< \frac{\varepsilon}{3} + \frac{\varepsilon}{3} + \frac{\varepsilon}{3} = \varepsilon. \end{aligned}$$

Damit ist gezeigt, daß φ^* im Punkte $\tilde{x}$ stetig ist. Das Argument funktioniert für alle $\tilde{x} \in D$. ■

Einige Schlüsse bei Cauchy können als lückenhaft (oder fehlerhaft) gelten, weil Cauchy bei seinen Konstruktionen „gleichmäßige Stetigkeit“ benutzt, aber nur die Stetigkeit als Voraussetzung genannt hat.

Definition φ gleichmäßig stetig in $D : \iff$

$$\iff \forall \varepsilon > 0 \exists \delta > 0 \forall x, y (d(x, y) < \delta \rightarrow e(\varphi(x), \varphi(y)) < \varepsilon)$$

Cauchy hatte hauptsächlich stetige Funktionen auf Intervallen $\subseteq \mathbb{R}$ im Auge; und diese Situation ist kaum geeignet, die Bedeutung des Begriffs der gleichmäßigen Stetigkeit herauszustellen. Nach ein paar scholastischen Beispielen kommen wir im letzten Beispiel und im Satz über die Fortsetzbarkeit gleichmäßig stetiger Funktionen zum Kern.

Beispiele für stetige reellwertige Funktionen auf einem Intervall

- 1) Die Exponentialfunktion ist gleichmäßig stetig auf jedem beschränkten Intervall, aber nicht gleichmäßig stetig auf $\mathbb{R}$.

Zum Beweis beweisen wir zunächst

Lemma $|h| \leq |e^h - 1| \leq \frac{|h|}{1 - |h|} \quad \text{für } |h| < 1$

Beweis Wir wissen $e^h \geq 1 + h$ für alle h

$$1 + h \leq e^h = \frac{1}{e^{-h}} \leq \frac{1}{1 - h}.$$

Für $h \geq 0$ haben wir

$$|h| \leq e^h - 1 = |e^h - 1| \leq \frac{1}{1 - h} - 1 = \frac{h}{1 - |h|}$$

und andererseits

$$\begin{aligned}\frac{-h}{1+h} &= \frac{1}{1+h} - 1 \leq e^{-h} - 1 \leq (1-h) - 1 = -h \\ h &\leq |e^{-h} - 1| = -(e^{-h} - 1) \leq \frac{h}{1-h} \\ \left| e^{x^*+h} - e^{x^*} \right| &= e^{x^*} \cdot \left| e^h - 1 \right|.\end{aligned}$$

Das Lemma zeigt für $f(x) = e^x$

$$|h| \cdot e^{x^*} \leq |f(x^*+h) - f(x^*)| \leq \frac{|h|}{1-|h|} \cdot e^{x^*}.$$

Zu jedem $\varepsilon > 0$ können wir $|h|$ so klein wählen, daß für alle $|x^*| \leq M < \infty$ der Zuwachs kleiner als ε ist. Zu jedem h können wir aber x^* so groß wählen, daß der Zuwachs groß ist.

- 2) Die Funktion $x \mapsto |x|^\beta$ ist für $0 \leq \beta \leq 1$ gleichmäßig stetig in $\mathbb{R}$. Für $\beta \neq 1$ ist sie nicht Lipschitz-stetig. (Definition der Lipschitz-Stetigkeit unten.)

Für $\beta > 0$ ist die Funktion Lipschitz-stetig auf jedem beschränkten Intervall.

- 3) Durch keine Wahl des Funktionswerts im Nullpunkt kann die Funktion $\sin \frac{1}{x}$ (für $x \neq 0$) zu einer im Nullpunkt stetigen Funktion gemacht werden.

Für jedes $\beta > 0$ ist

$$f(x) := |x|^\beta \cdot \sin \frac{1}{x}, \quad f(0) = 0$$

eine stetige Funktion. Für $0 < \beta \leq 1$ ist sie gleichmäßig stetig

- 4) $f(x) = \frac{1}{x} \cdot \sin x$, $f(0) = 1$
ist Lipschitz-stetig auf $\mathbb{R}$.

- 5) $f(x) = \sin(x^2)$

ist stetig, aber nicht gleichmäßig stetig auf $\mathbb{R}$. Es gibt nämlich Intervalle beliebig kleiner Länge, in welchen $f(\cdot)$ die Werte $+1$ und -1 annimmt.

- 6) $f(x) = \frac{1}{x} \sin(x^2)$, $f(0) = 0$
ist Lipschitz-stetig auf $\mathbb{R}$.

- 7) Die auf $\mathbb{Q}$ definierte reellwertige Funktion

$$\mathbb{Q} \ni x \mapsto \frac{1}{x - \sqrt{2}}$$

ist auf $\mathbb{Q}$ stetig, aber nicht gleichmäßig stetig.

Die Fortsetzung gleichmäßig stetiger Funktionen

Wir betrachten gleichmäßig stetige Funktionen auf einem abstrakten metrischen Raum $(D, d(\cdot, \cdot))$. Der wichtigste Fall ist der der reellwertigen Funktionen. Der Wertebereich der Funktionen (oder „Abbildungen“) kann aber auch irgendein anderer vollständiger metrischer Raum $(E, e(\cdot, \cdot))$ sein; die Überlegungen ändern sich dadurch nicht. Der Anfänger, der einen konkreten Fall vor Augen haben möchte, möge etwa an gleichmäßig stetige reellwertige Funktionen auf einem Definitionsbereich wie $D = \mathbb{C}_{\text{rat}}$ (mit der üblichen Metrik) denken.

Theorem Jede auf $(D, d(\cdot, \cdot))$ gleichmäßig stetige reellwertige Funktion kann in eindeutiger Weise zu einer stetigen Funktion auf der Vervollständigung fortgesetzt werden.

Beweis Wir beweisen die Fortsetzbarkeit in dem allgemeineren Fall einer gleichmäßig stetigen Abbildung in einen vollständigen metrischen Raum $(E, e(\cdot, \cdot))$.

$$\varphi : (D, d(\cdot, \cdot)) \longrightarrow (E, e(\cdot, \cdot))$$

wird fortgesetzt zu

$$\overline{\varphi} : (\overline{D}, \overline{d}(\cdot, \cdot)) \longrightarrow (E, e(\cdot, \cdot))$$

Wir haben gesehen, daß jedes $\overline{x}$ aus der Vervollständigung als Limes einer D -wertigen Cauchy-Folge gewonnen werden kann. Sei $(x_n)_n$ eine Cauchy-Folge mit dem Grenzwert $\overline{x}$. Wegen der gleichmäßigen Stetigkeit von $\varphi(\cdot)$ ist $(\varphi(x_n))_n$ eine Cauchy-Folge; sie hat einen Limes in E . Für jede Folge $(y_n)_n$ mit $\lim y_n = \overline{x}$ erhalten wir denselben Limes; wir bezeichnen ihn mit $\overline{\varphi}(\overline{x})$. $\overline{\varphi}(\cdot)$ setzt $\varphi(\cdot)$ auf $\overline{D}$ fort. $\overline{\varphi}(\cdot)$ ist ebenso wie $\varphi(\cdot)$ gleichmäßig stetig; denn, wenn $\delta > 0$ zu $\varepsilon > 0$ so gewählt ist, daß

$$\forall (x, y) \ (d(x, y) < \delta \longrightarrow e(\varphi(x), \varphi(y)) < \varepsilon)$$

wahr ist, dann ist auch die folgende Aussage wahr

$$\forall (\overline{x}, \overline{y}) \ (\overline{d}(\overline{x}, \overline{y}) < \delta \longrightarrow e(\overline{\varphi}(\overline{x}), \overline{\varphi}(\overline{y})) < \varepsilon)$$

Damit ist der Fortsetzungssatz bewiesen. ■

Die Überlegungen beweisen auch den

Satz Sei φ eine Abbildung

$$\varphi : (D, d(\cdot, \cdot)) \longrightarrow (E, e(\cdot, \cdot)) .$$

- a) Wenn φ gleichmäßig stetig ist, dann ist die Bildfolge $(\varphi(x_n))_n$ einer beliebigen Cauchy-Folge $(x_n)_n$ selbst eine Cauchy-Folge.
- b) Wenn φ in x^* stetig ist, dann gilt

$$\lim x_n = x^* \implies \lim \varphi(x_n) = \varphi(x^*) .$$

Lipschitz-Stetigkeit und der Fixpunktsatz für Kontraktionen

Definition φ Lipschitz-stetig mit einem Dehnungsfaktor $\leq K : \iff$

$$\iff \forall x, y \quad e(\varphi(x), \varphi(y)) \leq K \cdot d(x, y)$$

Lipschitz-stetige Abbildungen sind offenbar gleichmäßig stetig.

Satz („Fixpunktsatz für Kontraktionen“)

Sei $(D, d(\cdot, \cdot))$ ein vollständiger metrischer Raum, und sei

$$\varphi : (D, d(\cdot, \cdot)) \longrightarrow (D, d(\cdot, \cdot))$$

eine Abbildung, so daß für ein $\alpha < 1$ gilt

$$\forall x, y : d(\varphi(x), \varphi(y)) \leq \alpha \cdot d(x, y) \quad (, \varphi \text{ ist eine } \alpha\text{-Kontraktion}) .$$

Dann existiert genau ein $x^* \in D$ mit

$$\varphi(x^*) = x^* , \quad \text{ein „Fixpunkt für } \varphi(\cdot) \text{“} .$$

Beweis Man bestimmt x^* als Limes einer Cauchy-Folge $(x_1, x_2, \dots)$:

Wähle x_1 irgendwie und setze

$$x_{n+1} = \varphi(x_n) \quad \text{für } n = 1, 2, \dots$$

Für $n > N$ gilt nach der Dreiecksungleichung

$$\begin{aligned} d(x_N, x_n) &\leq d(x_N, x_{N+1}) + d(x_{N+1}, x_{N+2}) + \dots + d(x_{n-1}, x_n) \\ &\leq d(x_N, x_{N+1}) [1 + \alpha + \alpha^2 + \dots] \\ &\leq \frac{1}{1 - \alpha} \cdot d(x_N, x_{N+1}) \leq \frac{\alpha^{N-1}}{1 - \alpha} \cdot d(x_1, x_2) . \end{aligned}$$

Für genügend großes N sind also alle diese Abstände kleiner als jedes vorgegebene $\varepsilon > 0$. $(x_n)_n$ ist also eine Cauchy-Folge. Bezeichne mit x^* den Grenzwert dieser Cauchy-Folge. Es gilt

$$x^* = \varphi(x^*)$$

Da x_1 beliebig war, ist x^* der einzige Fixpunkt von $\varphi(\cdot)$.

Hinweis Der Begriff der Stetigkeit und Folgenstetigkeit wird in Analysis II wieder aufgenommen. Die wahre Bedeutung dieser Begriffe wird nämlich erst in allgemeineren Situationen klar, in Situationen nämlich, in welchen keine Metrik zugrundegelegt wird.

C.5 Die Eigenschaft von Bolzano–Weierstraß: Sätze über stetige Funktionen

Der klassische Satz von Bolzano–Weierstraß besagt

Satz Jede beschränkte Folge reeller Zahlen besitzt eine konvergente Teilfolge.
Er wird üblicherweise mit dem „Intervallschachtelungsprinzip“ bewiesen.

Wir werden unten einen allgemeineren Satz beweisen.

Definition Ein metrischer Raum $(D, d(\cdot, \cdot))$ heißt ein **BW–Raum**, wenn jede Folge $(x_n)_n$ eine konvergente Teilfolge besitzt.

Hinweis Die BW–Eigenschaft ist in der allgemeinen Topologie als Folgenkompaktheit bekannt. Sie wird dort mit dem Begriff der Kompaktheit konfrontiert. Auf den für diese Begriffe passenden Rahmen wollen wir uns hier nicht einlassen.

Definition Ein metrischer Raum $(D, d(\cdot, \cdot))$ heißt **totalbeschränkt**, wenn er für jedes $\delta > 0$ mit endlich vielen δ –Kugeln überdeckt werden kann.

Bemerkung Wenn $(D, d(\cdot, \cdot))$ nicht totalbeschränkt ist, dann existiert ein $\delta^* < 0$ und eine unendliche Folge $x_1, x_2, \dots$, so daß $d(x_i, x_j) \geq \delta^*$ für alle $i \neq j$.

Der Beweis liegt auf der Hand. Wähle $\delta^* > 0$ geeignet und beginne mit irgendeinem x_1 . Wähle x_2 außerhalb der δ^* –Kugel mit dem Mittelpunkt x_1 usw.

Hinweis In der Theorie der normierten (oder allgemeiner: topologischen) Vektorräume gibt es den Begriff der beschränkten Teilmenge. Eine in diesem Sinne beschränkte Menge ist im allgemeinen nicht totalbeschränkt.

Betrachten wir z.B. die Menge der trigonometrischen Polynome, ausgestattet mit der Supremumsnorm. Die Folge $(\sin nx)_n$ besitzt keine Teilfolge, die Cauchy–Folge ist.

Satz Ein metrischer Raum ist genau dann ein BW–Raum, wenn er vollständig und totalbeschränkt ist.

Beweis

- 1) Aus der BW-Eigenschaft folgt die Totalbeschränktheit aufgrund der obigen Bemerkung. Eine Cauchy-Folge, welche eine konvergente Teilfolge besitzt, ist selbst konvergent. Somit ist ein BW-Raum notwendigerweise vollständig.
- 2) Sei $(D, d(\cdot, \cdot))$ totalbeschränkt und $x_1, x_2, \dots$ irgendeine D -wertige Folge. Wir konstruieren dazu eine Teilfolge $z_1, z_2, \dots, z_n = x_{m(n)}, m(1) < m(2) < \dots$ welche Cauchy-Folge ist.

Überdecken wir D mit endlich vielen Kugeln mit dem Radius 1. Mindestens eine von ihnen, nennen wir sie $B^{(1)}$, enthält unendlich viele Folgeelemente; wählen wir ein solches: $z_1 = x_{m(1)}$.

Überdecken wir D (und damit auch $B^{(1)}$) mit endlich vielen Kugeln vom Radius $\frac{1}{2}$. Mindestens eine von ihnen, nennen wir sie $B^{(2)}$, enthält unendlich viele in $B^{(1)}$ gelegene Folgeelemente x_k mit $k > m(1)$. Wählen wir ein solches: $z_2 = x_{m(2)}$.

So konstruieren wir $z_n = x_{m(n)} \in B^{(1)} \cap B^{(2)} \cap \dots \cap B^{(n)}$, so daß $B^{(n)}$ eine Kugel vom Radius $\frac{1}{n}$ ist und in $B^{(1)} \cap \dots \cap B^{(n)}$ unendlich viele Folgeelemente x_k mit $k > m(n)$ liegen. Die Folge $z_1, z_2, \dots$ ist dann eine Cauchy-Folge.

Korollar Die Vervollständigung eines totalbeschränkten metrischen Raums ist ein BW-Raum.

Der klassische Satz von Bolzano–Weierstraß ergibt sich als Spezialfall des Satzes. Man formuliert ihn auch manchmal so: Eine Teilmenge von $\mathbb{R}$ (mit der üblichen Metrik ausgestattet) ist genau dann folgenkompakt, wenn sie beschränkt und abgeschlossen ist.

Sätze über stetige Funktionen auf einem BW-Raum

Satz 1 Jede stetige Funktion auf einem BW-Raum ist gleichmäßig stetig.

Satz 2 Jede stetige Funktion auf einem BW-Raum nimmt ihr Maximum an.

Satz 3 (Satz von Dini) Wenn eine Funktionfolge auf einem BW-Raum monoton gegen die Nullfunktion konvergiert, dann konvergiert sie gleichmäßig gegen die Nullfunktion.

Beweis von Satz 1 $(D, d(\cdot, \cdot))$ sei ein BW-Raum und $f(\cdot)$ eine reellwertige Funktion auf D , welche nicht gleichmäßig stetig ist.

$$\exists \varepsilon > 0 \quad \forall \delta > 0 \quad \exists (x, y) \quad (d(x, y) < \delta) \wedge |f(y) - f(x)| \geq \varepsilon .$$

Wir finden nun einen Punkt x^* , in welchem $f(\cdot)$ unstetig ist. Dazu wählen wir $\tilde{\varepsilon} > 0$ genügend klein und zu $\delta = \frac{1}{n}$ ein Punktepaar (x_n, y_n) mit $d(x_n, y_n) < \frac{1}{n}$ und $|f(y_n) - f(x_n)| \geq \tilde{\varepsilon}$. $(x_n)_n$ besitzt eine konvergente Teilfolge. In ihrem Limespunkt x^* ist $f(\cdot)$ unstetig; denn in jeder Umgebung von x^* gibt es Punkte x, y , so daß $|f(x^*) - f(x)| \geq \frac{\tilde{\varepsilon}}{2}$ oder $|f(x^*) - f(y)| \geq \frac{\tilde{\varepsilon}}{2}$.

Bemerke Der Beweis funktioniert auch, wenn $f(\cdot)$ Werte in irgendeinem metrischen Raum annimmt. Die nachfolgenden Beweise der Sätze 2 und 3 funktionieren auch für oberhalbstetige reellwertige Funktionen. Zum Begriff von Unterhalb- und Oberhalbstetigkeit siehe Anhang.

Beweis von Satz 2 Sei $f(\cdot)$ eine oberhalbstetige reellwertige Funktion auf dem BW-Raum $(D, d(\cdot, \cdot))$. Sei

$$s = \sup\{f(x) : x \in D\}.$$

Es existiert eine Folge $(x_n)_n$ mit $f(x_n) \nearrow s$. o.B.d.A. können wir annehmen, daß $(x_n)_n$ eine konvergente Folge ist. In ihrem Limespunkt x^* haben wir $f(x^*) = s$. $f(x^*) \leq s$ ist ohnehin klar. Andererseits gilt aber wegen der Oberhalbstetigkeit

$$f(x^*) \geq \limsup f(x_n) = s.$$

Beweis des Satzes von Dini Seien $f_1 \geq f_2 \geq \dots$ nichtnegative oberhalbstetige Funktionen. Gleichmäßige Konvergenz gegen die Nullfunktion bedeutet

$$\forall \varepsilon > 0 \quad \exists N \quad \forall n \geq N \quad \sup\{f_n(x) : x \in D\} < \varepsilon.$$

Nehmen wir an, daß $(f_n)_n$ nicht gleichmäßig gegen die Nullfunktion konvergiert und folgern wir daraus, daß es einen Punkt x^* gibt mit

$$\lim \searrow f_k(x^*) > 0.$$

Nehmen wir also an

$$\exists \varepsilon > 0 \quad \forall N \quad \exists n \geq N \quad \exists x \quad (f_n(x) \geq \varepsilon).$$

Wählen wir $\tilde{\varepsilon} > 0$ genügend klein und zu $N = 1, 2, \dots$ einen Index $n \geq N$ und ein x_n mit $f_n(x_n) \geq \tilde{\varepsilon}$. o.B.d.A. können wir annehmen, daß $(x_n)_n$ konvergiert. Der Grenzwert sei x^* . Wir haben wegen der Oberhalbstetigkeit von $f_k(\cdot)$

$$f_k(x^*) \geq \limsup_{n \rightarrow \infty} f_k(x_n)$$

und wegen $f_1 \geq f_2 \geq \dots$ für jedes feste k

$$\tilde{\varepsilon} \leq \limsup_{n \rightarrow \infty} f_n(x_n) \leq \limsup_{n \rightarrow \infty} f_k(x_n) \leq f_k(x^*) .$$

Hinweise

- 1) Im $\mathbb{R}^d$ mit irgendeiner Norm ist es klar, was es bedeutet, daß eine Teilmenge totalbeschränkt ist. Wenn es auch nicht falsch ist, so weist es doch in eine falsche Richtung, wenn man sagt: Eine Teilmenge des $\mathbb{R}^d$ ist genau dann ein BW–Raum, wenn sie beschränkt und abgeschlossen ist.
- 2) In komplizierten metrischen Räumen ist es meistens nicht so einfach, die totalbeschränkten Mengen mit handlichen Bedingungen zu charakterisieren. Im nächsten Semester werden uns gewissen Funktionenräume und Räume von Maßen in dieser Hinsicht beschäftigen.
- 3) E. HEINE (1821–1881) hat eher beiläufig die Beobachtung ausgesprochen, daß jede abzählbare offene Überdeckung eines abgeschlossenen endlichen Intervalls $[a, b]$ schon eine endliche Teilüberdeckung besitzt. E. BOREL (1871–1956) hat diese Beobachtung als eigenständigen Satz formuliert.

Die Beobachtung von Heine kann als Satz über BW–Räume formuliert und bewiesen werden.

Satz (Heine–Borel) Wenn ein abzählbares System von offenen Mengen den BW–Raum $(D, d(\cdot, \cdot))$ überdeckt, dann gibt es ein endliches Teilsystem, welches D überdeckt.

Beweis durch Widerspruch Seien $U_1, U_2, \dots$ offen mit $\bigcup_{i=1}^{\infty} U_i = D$.

Angenommen, zu jedem n existiert $x_n \in D$ mit

$$x_n \notin \bigcup_{i=1}^n U_i .$$

Wählen wir eine konvergente Teilfolge und nennen wir ihren Limes x^* . Es gibt ein N , so daß $x^* \in U_N$. Da U_N offen ist, liegen schließlich alle Elemente der Teilfolge in U_N . Das ist ein Widerspruch.

Hinweis Die Beziehungen zwischen den Begriffen
Kompaktheit
Folgenkompaktheit
Eigenschaft von Heine–Borel
Eigenschaft von Bolzano–Weierstraß
werden in Analysis II für allgemeine topologische Räume (ohne Metrik) diskutiert.

Anhang :

Stetigkeit, Folgenstetigkeit, Halbstetigkeit

a) Stetigkeit mit den Urbildern offener Mengen erklärt

Notation $\varphi : D \rightarrow E$. Für $B \subseteq E$ definieren wir

$$\varphi^{-1}(B) := \{x : x \in D, \varphi(x) \in B\}$$

und nennen diese Menge das *volle Urbild* von B bzgl. φ .

Bemerke

- (a) $\varphi^{-1}(E) = D$
- (b) $\varphi^{-1}(E \setminus B) = D \setminus \varphi^{-1}(B)$ für alle $B \subseteq E$
- (c) $\varphi^{-1}\left(\bigcup_{\alpha} B_{\alpha}\right) = \bigcup_{\alpha} \varphi^{-1}(B_{\alpha})$.

Seien nun $(D, d(\cdot, \cdot))$ und $(E, e(\cdot, \cdot))$ metrische Räume. $\mathfrak{U}$ bezeichne das System der offenen Mengen in D , $\mathfrak{V}$ das System der offenen Mengen in E .

Satz φ stetig in $D \iff \forall V \in \mathfrak{V} \quad (\varphi^{-1}(V) \in \mathfrak{U})$

(„ φ ist genau dann stetig, wenn die vollen Urbilder offener Mengen offen sind.“)

Beweis

- 1) Sei φ stetig und V offen, $U = \varphi^{-1}(V)$.

Zu zeigen ist

$$\forall x^* \in U \exists \delta > 0 \forall x (d(x, x^*) < \delta \longrightarrow x \in \varphi^{-1}(V)) .$$

Da $\varphi(x^*) \in V$ und V offen ist, existiert $\varepsilon > 0$, so daß

$$\forall z (e(z, \varphi(x^*)) < \varepsilon \longrightarrow z \in V) .$$

Wählen wir ein solches $\varepsilon > 0$ und dazu $\delta > 0$, so daß

$$\forall x (d(x, x^*) < \delta \longrightarrow e(\varphi(x), \varphi(x^*)) < \varepsilon) .$$

Dies zeigt $\{x : d(x, x^*) < \delta\} \subseteq \{x : \varphi(x) \in V\} = \varphi^{-1}(V) = U$.

2) Sei φ unstetig in x^* , d.h.

$$\exists \varepsilon > 0 \forall \delta > 0 \exists x (d(x, x^*) < \delta) \wedge (e(\varphi(x), \varphi(x^*)) \geq \varepsilon) .$$

Betrachte zu einem solchen $\tilde{\varepsilon} > 0$ die offene Kugel

$$V := \{z : e(z, \varphi(x^*)) < \tilde{\varepsilon}\} .$$

Es gilt dann $x^* \in \varphi^{-1}(V)$ und für alle $\delta > 0$

$$\{x : d(x, x^*) < \delta\} \not\subseteq \varphi^{-1}(V) .$$

$\varphi^{-1}(V)$ ist also nicht offen.

b) Folgenstetigkeit

Satz φ stetig in $x^* \iff \forall (x_n)_n (\lim x_n = x^* \implies \lim \varphi(x_n) = \varphi(x^*))$

(„ φ ist genau dann stetig in x^* , wenn für jede gegen x^* konvergierende Folge die Bildfolge gegen $\varphi(x^*)$ konvergiert.“)

Beweis

(a) Die Implikation $\implies$ haben wir bereits bewiesen.

(b) Sei φ unstetig in x^* . Wir konstruieren eine Folge $(x_n)_n$ mit $x_n \rightarrow x^*$ und $\varphi(x_n) \not\rightarrow \varphi(x^*)$.

Wähle $\tilde{\varepsilon} > 0$, so daß

$$\forall \delta > 0 \exists x (d(x, x^*) < \delta) \wedge (e(\varphi(x), \varphi(x^*)) \geq \tilde{\varepsilon}) .$$

$(x_n)_n$ konvergiert gegen x^* , aber $(\varphi(x_n))_n$ konvergiert nicht gegen $\varphi(x^*)$.

Hinweis In der allgemeinen Topologie definiert man die Stetigkeit einer Abbildung als die Eigenschaft, daß volle Urbilder offener Mengen offen sind. Die Folgenstetigkeit impliziert im allgemeinen nicht die Stetigkeit.

c) Halbstetigkeit

Wir betrachten reellwertige Funktionen $f(\cdot)$ auf einem beliebigen metrischen Raum $(D, d(\cdot, \cdot))$. Offenbar gilt der

Satz f stetig auf $D \iff \forall \alpha < \beta \ (\{x : \alpha < f(x) < \beta\} \text{ offen})$.

Definition

f unterhalbstetig auf $D \iff \forall \alpha \ (\{x : \alpha < f(x)\} \text{ offen})$

f oberhalbstetig auf $D \iff (-f)$ unterhalbstetig auf D .

Offenbar gilt

Satz f stetig auf $D \iff f$ unterhalb- und oberhalbstetig.

Den Begriff der Unterhalbstetigkeit erweitert man in der offensichtlichen Weise auf Funktionen mit Werten in $\mathbb{R} \cup \{+\infty\}$. Beispiele werden wir im Kapitel über konvexe Funktionen kennenlernen.

Satz Der Limes einer aufsteigenden Folge unterhalbstetiger Funktionen ist unterhalbstetig.

Beweis

$$f_1 \leq f_2 \leq \dots, \quad f^* = \lim \uparrow f_n, \quad f^*(x) = \sup_n f_n(x).$$

Für jedes $\alpha \in \mathbb{R}$ gilt

$$\{x : f^*(x) > \alpha\} = \bigcup_n \{x : f_n(x) > \alpha\} \text{ offen}.$$

Beispiel

$$f_n(x) = \frac{x^2}{1+x^2} + \frac{x^2}{(1+x^2)^2} + \dots + \frac{x^2}{(1+x^2)^n}.$$

Alle die f_n sind stetig; der aufsteigende Limes f^* ist im Nullpunkt unstetig. f^* ist aber unterhalbstetig.

Satz Die Indikatorfunktion einer offenen Menge ist unterhalbstetig. Man kann sie als aufsteigenden Limes einer Folge stetiger Funktionen darstellen.

Zur Konstruktion der Approximation brauchen wir eine Vorbereitung

Definition Sei $(D, d(\cdot, \cdot))$ ein metrischer Raum und $A \subseteq D$. Wir definieren den Abstand von x zu A

$$\text{dist}(x, A) := \inf\{d(x, z) : z \in A\}.$$

Lemma Die Funktion

$$g(x) := \text{dist}(x, A)$$

ist Lipschitzstetig mit einem Dehnungsfaktor ≤ 1 .

Beweis Zu zeigen ist

$$|g(y) - g(x)| \leq d(x, y) \quad \text{für alle } x, y \in D.$$

Für jedes $z \in A$ gilt

$$g(y) \leq d(y, z) \leq d(y, x) + d(x, z).$$

Daraus ergibt sich

$$g(y) \leq d(y, x) + g(x).$$

Ebenso beweist man $g(x) \leq d(x, y) + g(y)$.

Beweis des Satzes Sei U offen und $A = D \setminus U$.

$$g_n^U(x) := 1 - (1 - n \cdot \text{dist}(x, A))^+.$$

$g_n^U(\cdot)$ ist Lipschitzstetig mit einem Dehnungsfaktor $\leq n$.

$$\lim \uparrow g_n^U = 1_U$$

Wenn nämlich x einen Abstand $\geq \frac{1}{n}$ vom Komplement von U hat, dann gilt $g_n^U(x) = 1$. Jedes $x \in U$ hat einen positiven Abstand von A ; für $x \notin U$ gilt $g_n^U(x) = 0$.

Verallgemeinerung Jede nichtnegative unterhalbstetige Funktion f kann als aufsteigender Limes von stetigen Funktionen gewonnen werden.

Beweisskizze

- (a) Betrachte eine unterhalbstetige Funktion f , welche nur die Werte 0, 1 und 2 annehmen kann.

$$U_2 := \{x : f(x) = 2\} = \{x : f(x) > 1\} \quad \text{offen}$$

$$U_1 := \{x : f(x) > 0\} \quad \text{offen}$$

$$f(x) = \begin{cases} 2 & \text{für } x \in U_2 \\ 1 & \text{für } x \in U_1 \setminus U_2 \\ 0 & \text{sonst} \end{cases}$$

Wir haben offenbar

$$g_n := g_n^{U_1} + g_n^{U_2} \nearrow f.$$

- (b) Wir approximieren f durch einfachere unterhalbstetige Funktionen

$$f^{(m)}(x) = \begin{cases} 0 & \text{falls } f(x) \leq \frac{1}{2^m} \\ \frac{j}{2^m} & \text{falls } f(x) \in \left(\frac{j}{2^m}, \frac{j+1}{2^m}\right] \text{ für } j = 1, 2, \dots \\ +\infty & \text{falls } f(x) = +\infty \end{cases}$$

$f^{(m)}$ ergibt sich also durch Abrundung auf das nächste ganzzahlige Vielfache von $\frac{1}{2^m}$. $f^{(m)}$ ist unterhalbstetig und es gilt

$$f^{(1)} \leq f^{(2)} \leq \dots \lim \uparrow f^{(m)} = f.$$

$$f^{(m)} = \frac{1}{2^m} \left[1_{\{f > \frac{1}{2^m}\}} + 1_{\{f > \frac{2}{2^m}\}} + 1_{\{f > \frac{3}{2^m}\}} + \dots \right]$$

- (c) Jedes $f^{(m)}$ kann man wie in 1) durch stetige Funktionen $g_n^{(m)}$ aufsteigend approximieren. Maxima stetiger Funktionen sind stetig. Die Maxima von immer mehr der $g_n^{(m)}$ liefern eine aufsteigende Approximation.

Den Begriff der Unterhalbstetigkeit kann man ebenso wie den Begriff der Stetigkeit lokalisieren. f ist unterhalbstetig genau dann, wenn f in jedem Punkt x^* unterhalbstetig ist, wobei die Unterhalbstetigkeit in x^* folgendermaßen definiert ist

Definition f unterhalbstetig in x^* : $\Longleftrightarrow$

$$\Longleftrightarrow \quad \forall (x_n)_n \ (\lim x_n = x^* \rightarrow f(x^*) \leq \liminf f(x_n))$$

$$\Longleftrightarrow \quad \forall \varepsilon > 0 \ \exists \delta > 0 \ \forall x \ (d(x, x^*) < \delta \rightarrow f(x) > f(x^*) - \varepsilon)$$

Der Beweis der Äquivalenz ist eine lohnende Übungsaufgabe. Man zeigt auch leicht, daß für zwei in x^* unterhalbstetige Funktionen f und g auch die Funktionen

$$f + g \quad \text{und} \quad f \wedge g$$

in x^* unterhalbstetig sind.

C.6 Der Satz von Picard–Lindelöf

Wir wollen nicht versuchen, allgemein zu sagen, was eine Differentialgleichung ist. Es handelt sich, grob gesagt, um eine Gleichung, in welcher Funktionen und Ableitungen dieser Funktionen vorkommen. Wir beschäftigen uns hier mit Differentialgleichungen der Form

$$y' = f(x, y) .$$

Man nennt sie gewöhnliche Differentialgleichungen erster Ordnung in expliziter Form.

Definition Eine Funktion $y = Y(x)$ heißt Lösung der Differentialgleichung $y' = f(x, y)$ im Intervall $I = (x_-, x_+)$, wenn für alle $x \in (x_-, x_+)$ gilt

$$Y(x) - y_0 = \int_{x_0}^x f(\xi, Y(\xi)) d\xi .$$

Ein solches $Y(\cdot)$ bezeichnet man als eine Lösung durch den Punkt (x_0, y_0) .

Bemerke

- a) Als Lösungskurven von $y' = f(x, y)$ kommen nur totalstetige $g(x)$ in Betracht, für welche $f(x, g(x))$ integrierbar in jedem abgeschlossenen Teilintervall von (x_-, x_+) .
- b) Es sollte klar sein, was es heißt, daß eine Lösung $\tilde{g}$ eine Fortsetzung der Lösung $g(\cdot)$ ist. Wenn $g_1(\cdot)$ und $g_2(\cdot)$ Lösungen auf offenen Intervallen I_1 bzw. I_2 mit nichtleerem Durchschnitt $I_1 \cap I_2 \neq \emptyset$ sind, dann gewinnt man daraus eine Lösung $\tilde{g}(\cdot)$ auf der Vereinigung $I_1 \cup I_2$.

Beispiel 1 Sei $f(x, y) = y^{2/3}$ für alle $y \in \mathbb{R}$. Wir erraten leicht die Lösungen $y = \psi_{(a,b)}(x)$ der Gleichung

$$y' = y^{2/3} .$$

Für jedes Paar $a < b$ erfüllt die Funktion

$$\psi_{(a,b)}(x) = \begin{cases} \frac{1}{27} (x - a)^3 & \text{für } x \leq a \\ 0 & \text{für } x \in [a, b] \\ \frac{1}{27} (x - b)^3 & \text{für } x \geq b \end{cases}$$

die Gleichung. Umgekehrt sieht man leicht, daß jede Lösung $y = Y(x)$ der gegebenen Differentialgleichung ein Stück von einem dieser $\psi_{(a,b)}(\cdot)$ ist.

Beispiel 2 $y' = y^2$.

Für jedes c ist

$$y = \frac{1}{x - c} \quad \text{für } x < c \quad \text{und für } x > c$$

eine Lösung. Eine weitere Lösung ist $y \equiv 0$.

Jede Lösung ist ein Stück einer der angegebenen Lösungen. Jede Lösung, die irgendwo ungleich 0 ist, „explodiert“ in irgendeinem c , wenn man sie fortzusetzen versucht.

Die Bezeichnung t für die unabhängige Variable ist dann beliebt, wenn man sich die Lösungen $y = Y(t)$ als Bahnkurven eines in der Zeit bewegten Körpers vorstellt. Wenn es sich um eine Bewegung im Raum handelt, ist y dreidimensional. Wir betrachten Bewegungen im $\mathbb{R}^d$ beschrieben durch Differentialgleichungen

$$\begin{pmatrix} \dot{y} \\ \dot{y}_1 \\ \dot{y}_2 \\ \vdots \\ \dot{y}_d \end{pmatrix} = \begin{pmatrix} f(t, y) \\ f_1(t, y_1, \dots, y_d) \\ f_2(t, y_1, \dots, y_d) \\ \vdots \\ f_d(t, y_1, \dots, y_d) \end{pmatrix}$$

$f(t, y)$ ist eine Abbildung

$$f: \mathbb{R} \times \mathbb{R}^d \supseteq G \rightarrow \mathbb{R}^d.$$

Beispiel 3

$$\begin{pmatrix} \dot{y}_1 \\ \dot{y}_2 \end{pmatrix} = \begin{pmatrix} -y_2 \\ +y_1 \end{pmatrix}.$$

Lösungen Für c, t^* erhalten wir die Lösung

$$\begin{pmatrix} y_1(t) \\ y_2(t) \end{pmatrix} = c \cdot \begin{pmatrix} \cos(t - t^*) \\ \sin(t - t^*) \end{pmatrix}$$

Abstrakte Beschreibung des Lösungsbegriffs

$y = Y(t)$ ist eine (lokale) Lösung von

$$\dot{y} = f(t, y)$$

durch den Punkt (t_0, y_0) , wenn es ein Intervall $I = (t_-, t_+)$ gibt, so daß

$$y(t) = y_0 + \int_{t_0}^t f(s, y(s)) ds \quad \text{für alle } t \in I.$$

Anders gesagt, wenn die Abbildung

$$\Phi : g(\cdot) \longmapsto y_0 + \int_{t_0}^{\cdot} f(s, g(s)) ds$$

die Funktion $y(\cdot)$ zum Fixpunkt hat.

Für gewisse Typen von Funktionen

$$f : G \longrightarrow \mathbb{R}^d, \quad \text{wobei } G \text{ offen in } \mathbb{R} \times \mathbb{R}^d \text{ ist,}$$

wollen wir nun die (lokale) Existenz und Eindeutigkeit der Lösung durch $(t_0, y_0) \in G$ zeigen.

Definition Sei $G \subseteq \mathbb{R} \times \mathbb{R}^d$ offen.

- a) Die Funktion $f : G \rightarrow \mathbb{R}^d$ heißt eine Lipschitzfunktion (bzgl. der zweiten Koordinate), wenn

$$\begin{aligned} &\exists L \geq 0 \quad \forall t \quad \forall y, \tilde{y} \\ &\left((t, y) \in G, (t, \tilde{y}) \in G \longrightarrow \|f(t, \tilde{y}) - f(t, y)\| \leq L \cdot \|\tilde{y} - y\| \right) \end{aligned}$$

- b) Man sagt, daß $f(\cdot, \cdot)$ lokal eine Lipschitzbedingung erfüllt, wenn zu jedem (x^*, y^*) eine Umgebung $U^* \subseteq G$ existiert, in welcher $f(\cdot, \cdot)$ eine Lipschitzfunktion ist.

Beispiele ($d = 1$)

- 1) $f(t, y) = y^2$ ist eine lokale Lipschitzfunktion in $G = \mathbb{R} \times \mathbb{R}$.
- 2) $f(t, y) = y^{2/3}$ ist nicht Lipschitzfunktion im Gebiet G , wenn G einen Punkt $(t^*, 0)$ enthält.

Satz (Eindeutigkeitssatz)

$$f : \mathbb{R} \times \mathbb{R}^d \supseteq G \longrightarrow \mathbb{R}^d$$

sei stetig und erfülle lokal eine Lipschitzbedingung. Dann stimmen zwei Lösungen von $\dot{y} = f(t, y)$ durch (t^*, y^*) in einer Umgebung von t^* überein.

Beweis Es seien $g(t)$ und $h(t)$ Lösungen durch (t^*, y^*) , also

$$\begin{aligned} g(t) - y^* &= \int_{t^*}^t f(s, g(s)) \, ds \\ h(t) - y^* &= \int_{t^*}^t f(s, h(s)) \, ds . \end{aligned}$$

Subtraktion liefert

$$g(t) - h(t) = \int_{t^*}^t [f(s, g(s)) - f(s, h(s))] \, ds .$$

Wenn t genügend nah an t^* ist, dann gilt wegen

$$\begin{aligned} \|f(s, g(s)) - f(s, h(s))\| &\leq L \cdot \|g(s) - h(s)\| \\ \|g(t) - h(t)\| &\leq L \cdot \int_{t^*}^t \|g(s) - h(s)\| \, ds \end{aligned}$$

Der Integrand ist beschränkt in einer Umgebung von t^* , also gilt für $t \in I$

$$\|g(t) - h(t)\| \leq L \cdot M_I \cdot |t - t^*|$$

wobei $M_I = \sup\{\|g(s) - h(s)\| : s \in I\}$.

Wählen wir I so klein, daß $L \cdot |t - t^*| \leq \frac{1}{2}$, dann gilt

$$\|g(t) - h(t)\| \leq \frac{1}{2} M_I \quad \text{für } t \in I .$$

(Wenn das Supremum M_i kleiner oder gleich dem halben Supremum ist, dann bedeutet das $M_I = 0$.)

Satz (Existenzsatz von Picard–Lindelöf)

$$f : \mathbb{R} \times \mathbb{R}^d \supseteq G \longrightarrow \mathbb{R}^d$$

sei stetig und erfülle eine lokale Lipschitzbedingung. Zu jedem $(t^*, y^*) \in G$ existiert dann in einer genügend kleinen Umgebung I von t^* eine Lösung der Differentialgleichung

$$\dot{y} = f(t, y) .$$

Beweis

1) Wahl von I :

Betrachte das Quadrat $V \subseteq G$

$$V := \{(t, y) : |t - t^*| \leq r, \quad \|y - y^*\| \leq r\}.$$

Wenn r genügend klein ist, gibt es Zahlen M, L , so daß

$$\|f(t, y)\| \leq M, \quad \|f(t, \tilde{y}) - f(t, y)\| \leq L \cdot \|\tilde{y} - y\| \quad \text{für alle } (t, y) \in V.$$

Setze $I = (t^* - \varepsilon, t^* + \varepsilon)$ mit $\varepsilon = \min\{r, \frac{r}{M}\}$.

Wenn $\{(t, g(t)) : t \in I\} \subseteq V$, dann haben wir für

$$G(t) = y^* + \int_{t^*}^t f(s, g(s)) ds, \quad t \in I.$$

$$\|G(t) - y^*\| \leq M \cdot |t - t^*| \leq M \cdot \varepsilon \leq r,$$

also

$$\{(t, G(t)) : t \in I\} \subseteq V.$$

Kurvenstücke über I durch (t^*, y^*) , die in V verlaufen, werden also durch die Abbildung Φ in ebensolche Kurvenstücke abgebildet.

2) Wir definieren rekursiv eine Folge von Kurvenstücken

$$y_0(t) = y^*, \quad y_1(t) = y^* + \int_{t^*}^t f(s, y_0(s)) ds, \quad y_1(\cdot) = \Phi(y_0(\cdot))$$

$$y_{n+1}(t) = y^* + \int_{t^*}^t f(s, y_n(s)) ds, \quad y_{n+1}(\cdot) = \Phi(y_n(\cdot)).$$

Wenn wir zeigen können, daß $(y_n(\cdot))_n$ gleichmäßig konvergiert, dann haben wir im Grenzwert $y^*(\cdot)$ einen Fixpunkt von Φ gefunden.

3) Wir beweisen durch vollständige Induktion

$$\|y_n(t) - y_{n-1}(t)\| \leq M \cdot L^{n-1} \frac{|t - t^*|^n}{n!}.$$

Für $n = 1$ ist das offensichtlich

$$\begin{aligned}
 \|y_{n+1}(t) - y_n(t)\| &\leq \int_{t^*}^t \|f(s, y_n(s)) - f(s, y_{n-1}(s))\| ds \\
 &\leq L \cdot \int_{t^*}^t \|y_n(s) - y_{n-1}(s)\| ds \\
 &\leq L \cdot M \cdot L^{n-1} \cdot \int_{t^*}^t \frac{|s - t^*|^n}{n!} ds \\
 &= M \cdot L^n \cdot \frac{|t - t^*|^{n+1}}{(n+1)!} .
 \end{aligned}$$

Zu jedem $\varepsilon > 0$ existiert ein N , so daß für alle $n > N$

$$\|(y_n - y_N)(t)\| \leq M \cdot \sum_{k=N}^{\infty} \frac{1}{k!} L^{k-1} \cdot |t - t^*|^k \leq \varepsilon \quad \text{für alle } t \in I .$$

Damit ist für

$$\dot{y} = f(t, y)$$

mit Lipschitz-stetigem $f(\cdot, \cdot)$ die lokale Existenz mit Eindeutigkeit der Lösungskurven bewiesen.

(Näheres siehe O. Forster, Analysis 2, S. 96 ff)

C.7 Konvexe Funktionen

Archimedes hat in seiner Schrift über „Kugel und Zylinder“ I. Buch definiert

„2. Nach einer Seite konkav [besser: konvex] nenne ich ein Kurvenstück von der Eigenschaft, daß, wenn man zwei beliebige Punkte des Kurvenstücks geradlinig verbindet, die zwischen diesen beiden Punkten liegenden Kurvenpunkte alle auf dieser Seite der Geraden liegen oder allenfalls auf der Geraden selbst, keinesfalls aber auf der anderen Seite.

4. Nach einer Seite konkav [besser: konvex] nenne ich solche Flächenstücke, die die Eigenschaft haben, daß, wenn man zwei beliebige Punkte der Fläche durch eine Gerade verbindet, alle Punkte dieser Geraden auf derselben Seite der Fläche liegen oder allenfalls auf der Fläche selbst, keinesfalls aber auf der anderen Seite.“

(Archimedes Buch ist 1922 in Ostwalds Klassiker der exakten Wiss. 202/210, Akadem. Verlagsges. Leipzig auf deutsch erschienen.)

In moderner Notation kann man die Definitionen viel knapper fassen. Wir denken nicht an „Kurvenstücke“ oder „Flächenstücke“, sondern an konvexe Funktionen $k(\cdot)$ über einem d -dimensionalen reellen Vektorraum V sowie an konvexe Teilmengen K eines solchen Vektorraums. Konvexe Funktionen dürfen bei uns auch den Wert $+\infty$ annehmen, aber nicht den Wert $-\infty$.

Definition

a) $k(\cdot)$ konvex $\iff$

$$\forall x, y \forall \lambda \in [0, 1] : k((1 - \lambda)x + \lambda y) \leq (1 - \lambda)k(x) + \lambda k(y)$$

b) $K(\cdot)$ konvex $\iff$

$$\forall x, y \in K \forall \lambda \in [0, 1] : (1 - \lambda)x + \lambda y \in K.$$

Bemerke

1) Für festgehaltene x, y beschreibt $(1 - \lambda)x + \lambda y$ die Verbindungsstrecke, wenn $\lambda \in [0, 1]$. Eine Menge ist genau dann konvex, wenn sie mit je zwei Punkten auch die Verbindungsstrecke enthält.

2) Zu einer Menge K betrachten wir die Funktion

$$f_K(x) = \begin{cases} 0 & \text{falls } x \in K \\ +\infty & \text{falls } x \notin K \end{cases}$$

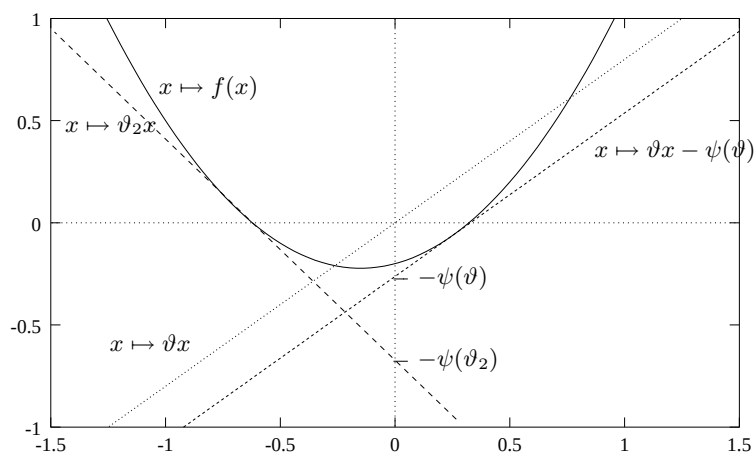
K ist genau dann konvex, wenn f_K eine konvexe Funktion ist.

3) Zu einer Funktion $k(\cdot)$ auf V betrachten wir den „Supergraphen“; das ist die Menge

$$\{(x, z) : z \geq k(x)\} \subseteq V \times \mathbb{R}$$

$k(\cdot)$ ist genau dann konvex, wenn der Supergraph eine konvexe Menge ist.

Beispiel
$$k(x) = \begin{cases} x \cdot \ln x & \text{für } x > 0 \\ 0 & \text{für } x = 0 \\ +\infty & \text{für } x < 0 \end{cases}$$



Beispiele für konvexe Funktionen kann man sehr leicht aufgrund des folgenden Satzes konstruieren.

Satz Das punktweise Supremum konvexer Funktionen ist konvex.

Beweis Sei $\{k_i(\cdot) : i \in I\}$ eine Familie konvexer Funktionen und $k = \sup k_i$. Für alle x, y , $i \in I$ und $0 \leq \lambda \leq 1$ gilt

$$k_i((1-\lambda)x + \lambda y) \leq (1-\lambda)k_i(x) + \lambda k_i(y) \leq (1-\lambda)k(x) + \lambda k(y)$$

also auch $k((1-\lambda)x + \lambda y) \leq (1-\lambda)k(x) + \lambda k(y)$. ■

Bemerkung Wenn die k_i unterhalbstetig sind, ist auch k unterhalbstetig. In der Tat sind die unterhalbstetigen konvexen Funktionen die bei weitem wichtigsten. $k(\cdot)$ ist unterhalbstetig konvex genau dann, wenn der Supergraph abgeschlossen und konvex ist (ohne Beweis).

Bezeichnungen Die affinen Funktionen sind konvex und konkav. Die bekanntere Kennzeichnung der affinen Funktion ist aber die: Eine Funktion $\ell(\cdot)$ auf dem Vektorraum V heißt affin, wenn $\ell(\cdot) - \ell(0)$ eine Linearform ist. Der Gebrauch des Namens „lineare Funktion“ ist nicht einheitlich; manche Autoren nennen die affinen Funktionen lineare Funktionen, andere bezeichnen nur die Linearformen als lineare Funktionen. Die Menge der Linearformen auf V heißt allgemein der Dualraum; er wird mit V^* bezeichnet.

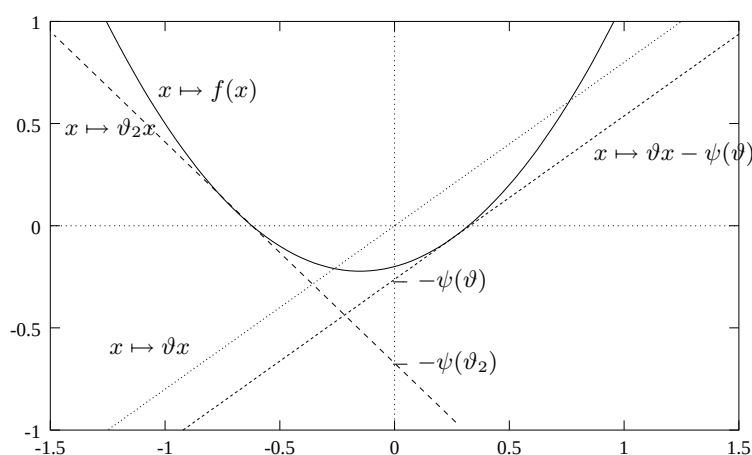
Konstruktion Sei $f(\cdot)$ eine unterhalbstetige Funktion, die mindestens eine affine Minorante besitzt.

Für jedes $\vartheta \in V^*$ definieren wir

$$\psi(\vartheta) = \sup\{\vartheta x - f(x) : x \in V\}.$$

Dieses Supremum kann $+\infty$ sein. Wenn $\psi(\vartheta) < \infty$, dann definieren wir die „**Subtangente zu ϑ** “

$$x \mapsto \vartheta x - \psi(\vartheta)$$



Die Subtangente minorisiert $f(\cdot)$. Außerdem gilt

$$\{x : f(x) - \vartheta(x) + \psi(\vartheta) = 0\} = \{x : f(x) - \vartheta x + \psi(\vartheta) \leq 0\}$$

ist abgeschlossen, wenn $f(\cdot)$ unterhalbstetig ist. Diese Menge von Punkten, in welchen die Subtangente gleich der Funktion ist, kann aber auch leer sein. Wenn $f(\cdot)$ außerdem konvex ist, dann handelt es sich um eine abgeschlossene konvexe Menge.

Das Studium der Subtangenten hat eine ehrwürdige Geschichte. FERMAT (1601–1665) hat sich ausführlich mit ihnen beschäftigt. Er schließt sich an die Definition von Euklid an, der (allerdings nur für den Kreis) definierte „Daß sie den Kreis berühre, sagt man von einer gerade Line, die einen Kreis trifft, ihn aber bei Verlängerung nicht schneidet.“ (Zitiert nach Walter S. 223).

Betrachten wir nun das punktweise Supremum aller möglichen Subtangenten

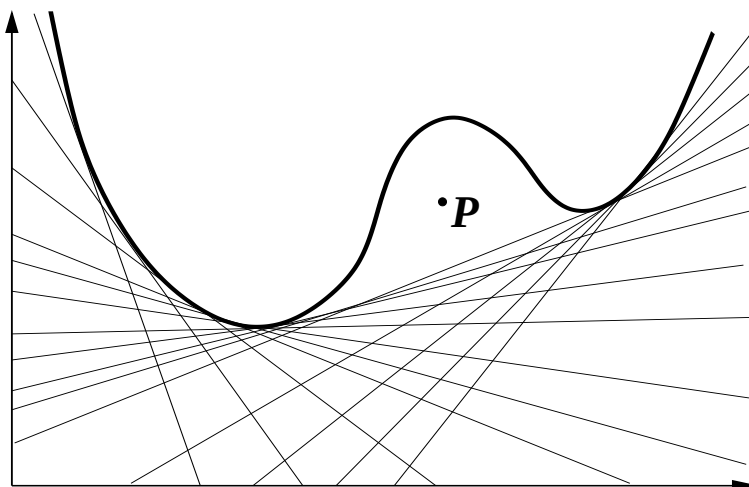
$$\psi^*(x) := \sup\{\vartheta x - \psi(\vartheta) : \vartheta \in V^*\} \quad \text{für } x \in V.$$

Offenbar ist ψ^* unterhalbstetige konvexe Minorante von f .

Theorem Für jede Funktion f (auf einem endlichdimensionalen Vektorraum V), welche mindestens eine affine Minorante besitzt, erhält man die größte unterhalbstetige konvexe Minorante ψ^* durch die Konstruktion

$$\begin{aligned} \psi(\vartheta) &:= \sup\{\vartheta x - f(x) : x \in V\} \quad \text{für } \vartheta \in V^* \\ \psi^*(x) &:= \sup\{\vartheta x - \psi(\vartheta) : \vartheta \in V^*\} \quad \text{für } x \in V. \end{aligned}$$

Bemerkung Jede Funktion, die überhaupt eine unterhalbstetige konvexe Minorante besitzt, besitzt auch eine größte unterhalbstetige konvexe Minorante; das Supremum aller unterhalbstetigen konvexen Minoranten ist ja selbst unterhalbstetig und konvex.



Das Theorem ist offenbar äquivalent mit der folgenden Aussage, die auf den ersten Blick spezieller aussehen mag.

Theorem Sei $k(\cdot)$ unterhalbstetig konvex auf V .

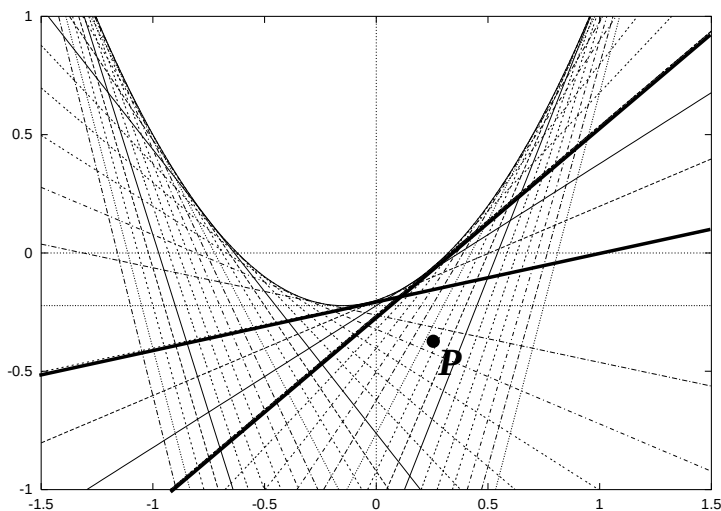
$$k^*(\vartheta) := \sup\{\vartheta x - k(x) : x \in V\} \quad \text{für } \vartheta \in V^*$$

$$(k^*)^*(x) := \sup\{\vartheta x - k^*(\vartheta) : \vartheta \in V^*\} \quad \text{für } x \in V.$$

Dann gilt $(k^*)^* = k$.

$((k^*)^* \leq k$ ist trivialerweise richtig; interessant ist $(k^*)^* \geq k$)

Varianten dieses Theorems (für alle möglichen unendlichdimensionalen Vektorräume) gehören zu den wichtigsten Einsichten der Reellen Analysis überhaupt. Manchmal spricht man von Trennungssätzen für konvexe Mengen; oft formuliert man sie als Varianten des Satzes von Hahn–Banach. Der Name „Trennungssatz“ wird durch das folgende Bild plausibel



Für jeden Punkt P außerhalb des Supergraphen von f , $P = (x, z)$ mit $z < k(x)$, existiert eine Hyperebene, welche P vom Supergraphen trennt.

Die folgenden Aussagen über $k(\cdot)$ sind offenbar äquivalent

- (i) $\forall (x, z) : k(x) > z \exists (\vartheta, \xi) : (\forall y : \vartheta y - \xi \leq k(y)) \wedge (\vartheta x - \xi \geq z)$.
- (ii) $\forall (x, z) : k(x) > z \exists (\vartheta, \xi) : (\xi \geq \sup\{\vartheta y - k(y) : y \in V\}) \wedge (\xi \leq \vartheta x - z)$
- (iii) $\forall (x, z) : k(x) > z \exists \vartheta : \sup\{\vartheta y - k(y) : y \in V\} \leq \vartheta x - z$
- (iv) $\forall (x, z) : k(x) > z \exists \vartheta : \vartheta x - k^*(\vartheta) \geq z$
- (v) $\forall x : \sup\{\vartheta x - k^*(\vartheta) : \vartheta \in V^*\} \geq k(x)$

$$(vi) \quad \forall x : k^*(k^*(x)) \geq k(x)$$

Hinweis zum Beweis des Theorems Man beweist (i) durch vollständige Induktion nach der Dimension von V . Im eindimensionalen Fall (unser Bild) betrachtet man die Geraden durch $P = (x, z)$ mit allen möglichen Steigungen ϑ . Für große ϑ mag die Gerade den Supergraphen rechts von x treffen; für kleine ϑ mag sie links treffen. Aus der Konvexität von $k(\cdot)$ ergibt sich, daß sie für kein ϑ auf beiden Seiten treffen kann. In der Tat gibt es eine Gerade durch P , die den Subgraphen überhaupt nicht trifft. Durch die so gewonnene Gerade ist nun eine Ebene zu legen, die den Supergraphen nicht trifft usw. Wir wollen den Induktionsschritt hier aber nicht durchführen. Trennungssätze für konvexe Mengen gehören nicht zum traditionellen Repertoire der Anfängervorlesung.

Definition („Legendre-Transformation“)

Die Konstruktion, die jeder Funktion f auf V die Funktion $\psi = f^*$ auf V^* zuordnet

$$\psi(\vartheta) = \sup\{\vartheta x - f(x) : x \in V\}$$

heißt die Legendre-Transformation.

$f^*(\cdot)$ heißt Legendre-Transformierte von $f(\cdot)$ oder auch die konvexkonjugierte Funktion zu $f(\cdot)$. (Diejenigen f , für welche überhaupt keine affine Minorante existiert, schließen wir aus. f und f^* dürfen sehr wohl auch den Wert $+\infty$ annehmen; der Wert $-\infty$ kommt jedoch nicht in Betracht.)

Spezialfall 1 Sei f eine Funktion, die nur die Werte 0 und $+\infty$ annehmen kann. $A = \{x : f(x) = 0\}$. Es gilt dann

$$f^*(\vartheta) = \sup\{\vartheta x : x \in A\}$$

f^* heißt die **Stützfunktion** der Menge A . Es gilt

$$f^*(\alpha\vartheta) = \alpha \cdot f^*(\vartheta) \quad \text{für alle } \alpha \geq 0, \quad \vartheta \in V^* \quad (f^* \text{ ist positiv homogen})$$

$f^*(\vartheta)$ ist die kleinste Zahl λ für welche A im Halbraum $H_\lambda = \{x : \vartheta x \leq \lambda\}$ liegt.

Bemerke Eine positiv homogene Funktion $F(\cdot)$ ist genau dann konvex, wenn sie subadditiv ist, d.h. wenn

$$F(x+y) \leq F(x) + F(y) \quad \text{für alle } x, y.$$

Beweis

- 1) Sei $F(\cdot)$ positiv homogen und subadditiv. Dann gilt für alle x, y und alle $\lambda \in [0, 1]$

$$F((1-\lambda)x + \lambda y) \leq F((1-\lambda)x) + F(\lambda y) = (1-\lambda)F(x) + \lambda F(y)$$

$F(\cdot)$ ist also konvex.

- 2) Sei $F(\cdot)$ positiv homogen und konvex. Dann gilt

$$F(x+y) = 2 \cdot F\left(\frac{x}{2} + \frac{y}{2}\right) \leq 2 \cdot \left[\frac{1}{2} F(x) + \frac{1}{2} F(y)\right] = F(x) + F(y)$$

$F(\cdot)$ ist also subadditiv.

Spezialfall 2 Sei F positiv homogen. Dann kann die Legendre-Transformierte $F^*(\cdot)$ nur die Werte 0 und $+\infty$ annehmen.

Beweis Für alle $\vartheta \in V^*$ und alle $\alpha > 0$ gilt

$$\begin{aligned} 0 \leq F^*(\vartheta) &= \sup\{\vartheta x - F(x)\} \\ &= \sup\left\{\alpha \cdot \vartheta \cdot \frac{x}{\alpha} - \alpha \cdot F\left(\frac{x}{\alpha}\right)\right\} \\ &= \alpha \cdot \sup\{\vartheta y - F(y)\} = \alpha \cdot F^*(\vartheta) . \end{aligned}$$

Daraus folgt die Behauptung.

Satz Sei A eine beliebige Menge $\subseteq V$ und K die kleinste abgeschlossene konvexe Obermenge. Sei $F(\cdot)$ die Stützfunktion von A . Dann gilt

$$F^*(x) = \begin{cases} 0 & \text{falls } x \in K \\ +\infty & \text{sonst} . \end{cases}$$

Der Beweis ergibt sich aus dem Theorem; die angegebene Funktion ist die größte unterhalbstetige konvexe Funktion, die auf A verschwindet.

Zwei wichtige Beispiele

Satz Betrachte die p -Norm auf dem $\mathbb{R}^d$. ($p \geq 1$)

$$F(x) = F(x^1, \dots, x^d) = \left(|x^1|^p + \dots + |x^d|^p\right)^{1/p}.$$

Für die Legendre-Transformierte

$$F^*(\vartheta) = F^*(\vartheta^1, \dots, \vartheta^d)$$

gilt dann (mit $\frac{1}{p} + \frac{1}{q} = 1$)

$$F^*(\vartheta) = \begin{cases} 0 & \text{falls } (|\vartheta^1|^q + \dots + |\vartheta^d|^q)^{1/q} \leq 1 \\ \infty & \text{sonst} \end{cases}.$$

Beweis

- 1) $F(\cdot)$ ist offenbar positiv homogen. Minkowskis Ungleichung besagt, daß $F(\cdot)$ subadditiv ist. Somit ist schon einmal klar, daß F^* nur die Werte 0 und $+\infty$ annehmen kann.
- 2) Wir zeigen $F^*(\vartheta) = 0$, wenn $\|\vartheta\|_q \leq 1$.

Aus Hölders Ungleichung

$$|\vartheta x| \leq \|\vartheta\|_q \cdot \|x\|_p$$

ergibt sich in der Tat für $\|\vartheta\|_q \leq 1$

$$\sup\{\vartheta x - F(x)\} = \sup\{|\vartheta x| - \|x\|_p\} \leq \sup\{\|\vartheta\|_q \cdot \|x\|_p - \|x\|_p\} = 0$$

- 3) Hölders Ungleichung ist scharf. Zu jedem ϑ existiert ein $\tilde{x}$ mit $\vartheta \tilde{x} = \|\vartheta\|_q \cdot \|\tilde{x}\|_p$. Es genügt, dies für die $\vartheta = (\vartheta^1, \dots, \vartheta^d)$ mit nichtnegativen Komponenten ϑ^i zu beweisen. Setze

$$\tilde{x}^i = (\vartheta^i)^{q/p} \quad \text{für } i = 1, \dots, d.$$

Für dieses $\tilde{x}$ haben wir

$$\begin{aligned} \|\tilde{x}\|_p &= \left(\sum (\vartheta^i)^q\right)^{1/p}, & \vartheta \tilde{x} &= \sum (\vartheta^i)^q \\ \|\vartheta\|_q \cdot \|\tilde{x}\|_p &= \left(\sum (\vartheta^i)^q\right)^{1/q+1/p} = \vartheta \tilde{x}. \end{aligned}$$

- 4) Wenn $\|\vartheta\|_q > 1$, dann gilt mit den eben konstruierten $\tilde{x}$

$$\sup\{\vartheta x - \|x\|_p\} \geq \|\vartheta\|_q \cdot \|\tilde{x}\|_p > 0.$$

Also $F^*(\vartheta) = +\infty$ für $\|\vartheta\|_q > 1$.

Eine weiterer interessanter Typ von Funktionen, in welchem man die Legendre-Transformierte explizit berechnen kann, sind die nach unten beschränkten quadratischen Funktionen

$$Q(x) = c + b^\top x + \frac{1}{2} x^\top A x .$$

(Wir schreiben die Elemente x, y von V als Spalten und ebenso die Elemente $\vartheta \in V^*$; der Wert von ϑ im Punkte x wird mit $\vartheta^\top x$ bezeichnet.)

Es genügt $F(x) = \frac{1}{2} x^\top A x$ zu studieren; denn

$$\begin{aligned} Q^*(\vartheta) &= \sup \left\{ \vartheta^\top x - c - b^\top x - \frac{1}{2} x^\top A x \right\} \\ &= -c + \sup \left\{ (\vartheta - b)^\top x - \frac{1}{2} x^\top A x \right\} \\ &= -c + F^*(\vartheta - b) \end{aligned}$$

Satz Sei A eine symmetrische $d \times d$ -Matrix. Die quadratische Funktion

$$F(x) = \frac{1}{2} x^\top \cdot A \cdot x$$

ist genau dann konvex, wenn A positiv semidefinit ist, d.h.

$$F(x) \geq 0 \text{ für alle } x .$$

In diesem Fall ist die Legendre-Transformierte $F^*(\vartheta)$ eine quadratische Funktion auf ihrem Endlichkeitsbereich. Dieser Endlichkeitsbereich ist der lineare Unterraum

$$\{\vartheta : F^*(\vartheta) < \infty\} = \{\vartheta : \vartheta = Ay \text{ für ein } y\} .$$

Beweis

- 1) Betrachten wir zunächst den Fall, wo A nichtsingulär ist, wo also A^{-1} existiert und jedes ϑ die Gestalt $\vartheta = Ay$ hat.

Für $\vartheta = Ay$ haben wir für alle x

$$\begin{aligned} \vartheta^\top x - F(x) &= \vartheta^\top x - \frac{1}{2} x^\top A x = y^\top A x - \frac{1}{2} x^\top A x \\ &= \frac{1}{2} y^\top A y - \frac{1}{2} (x + y)^\top A (x + y) \\ F^*(\vartheta) &= \sup \left\{ \vartheta^\top x - F(x) : x \in V \right\} \\ &= \frac{1}{2} y^\top A y = \frac{1}{2} \vartheta^\top \cdot A^{-1} \cdot \vartheta . \end{aligned}$$

2) Auch im allgemeinen Fall haben wir für jedes ϑ der Gestalt $\vartheta = Ay$ den Wert

$$F^*(Ay) = \frac{1}{2} y^\top Ay .$$

Wir zeigen, daß für die übrigen ϑ $F^*(\vartheta) = +\infty$ gilt. Liegt ϑ nicht im „Bild“ $\{Ay : y \in V\}$, dann gibt es ein z mit

$$z^\top Ay = 0 \text{ für alle } y \text{ und } z^\top \vartheta \neq 0 .$$

Es gilt dann für alle α

$$\begin{aligned} \vartheta^\top(\alpha z) - F(\alpha z) &= \alpha \cdot (\vartheta^\top z) - 0 \\ F^*(\vartheta) &\geq \sup \left\{ \vartheta^\top(\alpha z) - F(\alpha z) : \alpha \right\} = +\infty . \end{aligned}$$

Es ist in der Regel schwierig, die Legendre–Transformierte einer vorgegebenen Funktion explizit auszurechnen. Das ist aber hier auch nicht das wesentliche Thema. Es geht uns um das Konstruktionsprinzip.

Anschluß an den „Kalkül“

Der Student kann hier, wenn ihm das Abstraktionsniveau zu hoch geworden ist, Übungen im Umgang mit den elementaren Funktionen anschließen — zur Erholung sozusagen.

Sei $k(\cdot)$ eine konvexe Funktion, die in einem Intervall (a, b) glatt ist. Es kann für manche ϑ passieren, daß die Subtangente

$$\vartheta x - k^*(\vartheta)$$

in einem wohlbestimmten Punkt $x(\vartheta) \in (a, b)$ die gegebene Funktion berührt. Das Supremum

$$k^*(\vartheta) = \sup \{ \vartheta x - k(x) \}$$

wird dann also in $x(\vartheta)$ angenommen und es gilt

$$k^*(\vartheta) = \vartheta \cdot x(\vartheta) - k(x(\vartheta)) .$$

In der Schule lernt man, Maximierungsaufgaben dadurch zu lösen, daß man die Ableitung gleich 0 setzt. In unserem Falle

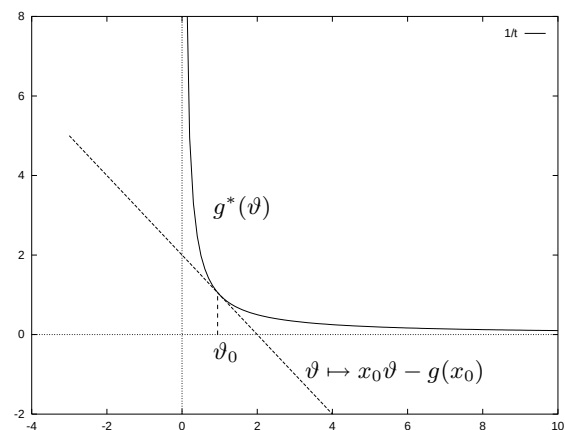
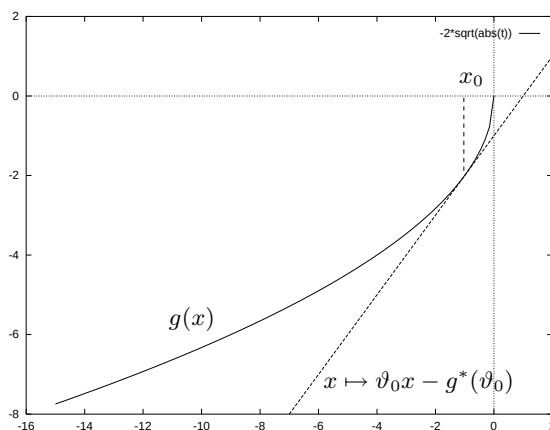
$$0 = (\vartheta x - k(x))' = \vartheta - k'(x) , \quad k'(x(\vartheta)) = \vartheta .$$

In dem in Frage kommenden Intervall ist also $\vartheta \mapsto x(\vartheta)$ die Umkehrfunktion zu $x \mapsto k'(x)$.

Beispiele

$$\begin{aligned}
 1) \quad k(x) &= \begin{cases} \frac{1}{x} & \text{für } x > 0 \\ +\infty & \text{sonst} \end{cases} \\
 k'(x) &= -\frac{1}{x^2} \stackrel{!}{=} \vartheta ; \quad x(\vartheta) = \frac{1}{\sqrt{|\vartheta|}}, \quad \vartheta < 0 \\
 \vartheta \cdot x(\vartheta) - k(x(\vartheta)) &= \vartheta \cdot \frac{1}{\sqrt{|\vartheta|}} - \sqrt{|\vartheta|} = -2 \cdot \sqrt{|\vartheta|} \quad \text{für } \vartheta < 0 \\
 k^*(\vartheta) &= \sup \left\{ \vartheta x - \frac{1}{x} : x > 0 \right\} = \begin{cases} -2 \cdot \sqrt{|\vartheta|} & \text{für } \vartheta \leq 0 \\ +\infty & \text{für } \vartheta > 0. \end{cases}
 \end{aligned}$$

$$1a) \quad g(x) = \begin{cases} -2 \cdot \sqrt{|x|} & \text{für } x \leq 0 \\ +\infty & \text{für } x > 0 \end{cases}$$



$$g^*(\vartheta) = \sup \left\{ \vartheta x + 2 \cdot \sqrt{|x|} : x \leq 0 \right\} = \begin{cases} \frac{1}{\vartheta} & \text{für } \vartheta > 0 \\ +\infty & \text{sonst} \end{cases} .$$

Im Bild sind x_0 und ϑ_0 konjugierte Punkte bzw. konjugierte Steigungen.

$$\begin{aligned}
 2) \quad k(x) &= e^x \quad \text{für } x \in \mathbb{R} \\
 k'(x) &= e^x \stackrel{!}{=} \vartheta \quad \text{für } \vartheta > 0, \quad x(\vartheta) = \ln \vartheta
 \end{aligned}$$

$$\vartheta \cdot x(\vartheta) - k(x(\vartheta)) = \vartheta \ln \vartheta - \vartheta \quad \text{für } \vartheta > 0$$

$$k^*(\vartheta) = \begin{cases} \vartheta \ln \vartheta - \vartheta & \text{für } \vartheta > 0 \\ -1 & \text{für } \vartheta = 0 \\ +\infty & \text{für } \vartheta < 0 \end{cases}$$

$$2a) \quad g(x) = \begin{cases} x \ln x - x + 1 & \text{für } x > 0 \\ 0 & \text{für } x = 0 \\ +\infty & \text{für } x < 0 \end{cases} = \int_1^x \ln y \, dy$$

$$\begin{aligned} g'(x) &= \ln x \stackrel{!}{=} \vartheta, & x(\vartheta) &= e^\vartheta \\ \vartheta \cdot x(\vartheta) - g(x(\vartheta)) &= \vartheta e^\vartheta - (e^\vartheta \cdot \vartheta - e^\vartheta + 1) = e^\vartheta - 1 \\ g^*(\vartheta) &= e^\vartheta - 1 \text{ für alle } \vartheta. \end{aligned}$$

$$3) \quad k(x) = \begin{cases} -\ln x & \text{für } x > 0 \\ +\infty & \text{sonst.} \end{cases}$$

$$\begin{aligned} k'(x) &= -\frac{1}{x} \stackrel{!}{=} \vartheta \text{ für } \vartheta < 0, & x(\vartheta) &= -\frac{1}{\vartheta} \\ \vartheta \cdot x(\vartheta) - k(x(\vartheta)) &= -\vartheta \cdot \frac{1}{\vartheta} + \ln\left(-\frac{1}{\vartheta}\right) \\ &= -1 - \ln(-\vartheta) \text{ für } \vartheta < 0. \end{aligned}$$

$$4) \quad k(x) = \begin{cases} x \ln x + (1-x) \ln(1-x) & \text{für } x \in (0, 1) \\ 0 & \text{für } x = 0, x = 1 \\ +\infty & \text{sonst} \end{cases}$$

$$\begin{aligned} k'(x) &= \ln x - \ln(1-x) \stackrel{!}{=} \vartheta, & \frac{x}{1-x} &\stackrel{!}{=} e^\vartheta \\ x(\vartheta) &= \frac{e^\vartheta}{1+e^\vartheta} \\ \vartheta \cdot x(\vartheta) - k(x(\vartheta)) &= \ln(1+e^\vartheta) = k^*(\vartheta) \text{ für alle } \vartheta. \end{aligned}$$

$$5) \quad k(x) = \cosh x = \frac{1}{2} (e^x + e^{-x})$$

$$k'(x) = \frac{1}{2} (e^x + e^{-x}) \stackrel{!}{=} \vartheta, \quad 2\vartheta e^x = e^{2x} - 1,$$

$$(e^x - \vartheta)^2 = 1 + \vartheta^2$$

$$x(\vartheta) = \ln \left(\vartheta + \sqrt{1 + \vartheta^2} \right)$$

$$\vartheta \cdot x(\vartheta) - k(x(\vartheta)) = \vartheta \cdot \ln \left(\vartheta + \sqrt{1 + \vartheta^2} \right) - \sqrt{1 + \vartheta^2} = k^*(\vartheta).$$

C.8 Normierte Vektorräume; die p -Normen

Definition Eine endlichwertige nichtnegative Funktion $n(\cdot)$ auf einem (reellen oder komplexen) Vektorraum V heißt eine **Norm**, wenn gilt

- (i) $n(x) > 0$ für alle $x \neq 0$
- (ii) $n(\lambda x) = |\lambda|n(x)$ für alle Skalare λ
- (iii) $n(\cdot)$ ist konvex

Bemerke Zu jeder Norm auf V gehört eine Metrik auf V

$$d(x, y) = n(y - x) .$$

Aus (ii) und (iii) folgt nämlich

$$n(x + y) \leq n(x) + n(y) \quad \text{„Subadditivität“}$$

und daraus die Dreiecksungleichung für $d(\cdot, \cdot)$. Man könnte kurz sagen: Jede Norm auf einem Vektorraum entspricht einer mit Translationen und Streckungen verträglichen Metrik auf dem Vektorraum, und umgekehrt.

Definition Ein Vektorraum mit einer ausgezeichneten Norm heißt ein normierter Vektorraum.

Satz Die Vervollständigung eines normierten Vektorraums ist ein vollständiger normierter Vektorraum.

Wir wollen den Satz nicht beweisen, sondern nur klarmachen, was bewiesen werden muß — und in der Tat leicht zu beweisen ist: Die Elemente der Vervollständigung sind zunächst nur Elemente eines metrischen Raums, als Äquivalenzklassen von Cauchy-Folgen abstrakt konstruiert. Dieser vollständige metrische Raum muß nun mit der Struktur eines normierten Vektorraums ausgestattet werden. In der Tat gibt es eine naheliegende Weise, die Elemente zu addieren und mit Skalaren zu multiplizieren. Die Axiome der Vektorrechnung sind erfüllt. Außerdem ist leicht zu sehen, daß die Metrik von einer Norm herrührt, d.h. daß sie mit Translationen und Streckungen im vervollständigten Vektorraum verträglich ist.

Beispiel 1 (Der Raum von $C([0, 1], \|\cdot\|_\infty)$)

Betrachte die Polynome $p: a_0 + a_1 t + \dots + a_n t^n$ mit reellen Koeffizienten (und beliebigem Grad) als Funktionen auf dem Einheitsintervall $[0, 1]$. Die Menge dieser Polynome ist ein Vektorraum V . (Daß sie sogar eine Algebra ist, d.h. daß man Polynome auch multiplizieren kann, interessiert uns hier nicht.) Stattdessen wird V mit der Supremumsnorm aus

$$\begin{aligned} n(p) &= \sup_{t \in [0, 1]} |p(t)| \\ d(p, q) &= \sup_{t \in [0, 1]} |p(t) - q(t)| = \|p - q\|_\infty. \end{aligned}$$

Die Vervollständigung ist (nach dem berühmten Approximationssatz von Weierstraß) der Raum $C([0, 1])$ aller auf $[0, 1]$ stetigen Funktionen, ausgestattet mit der Supremumsnorm.

Hinweis Für die „guten“ (d.h. sowohl transparenten als auch kräftig verallgemeinerbaren) Beweise des Satzes von Weierstraß ist es sehr wohl wichtig, daß die Menge der Polynome eine Algebra ist; die Aussage des Satzes kann man aber auch verstehen, ohne daß man sich um Strukturen kümmert, die über die Vektorraumstruktur hinausgehen.

Beispiel 2 Wir statteten $C([0, 1])$ mit der sog. p -Norm aus

$$\|f\|_p := \left(\int_0^1 |f(t)|^p dt \right)^{1/p} \quad (p \geq 1)$$

Die Vervollständigung heißt der Raum $\mathcal{L}^p([0, 1])$ oder auch $\mathcal{L}^p([0, 1]; dt)$. Um diesen vollständigen Vektorraum zu verstehen, braucht man Lebesguesche Integrationstheorie, insbesondere den Begriff der Nullmenge und den Begriff der Fastüberall-Gleichheit von Funktionen. Man kann nämlich die Elemente von $\mathcal{L}^p([0, 1])$ als Äquivalenzklassen von Funktionen interpretieren. Die Sätze, auf welchen diese grundlegende Tatsache beruht, lauten:

Theorem Für jede $d_p(\cdot, \cdot)$ -Cauchyfolge existiert eine Teilfolge, die fastüberall konvergiert; und die Grenzfunktion ist eine p -integrierbare Funktion, gegen welche die Cauchy-Folge in der p -Norm konvergiert.

(Beweis in Analysis II, Integrationstheorie)

Auf der anderen Seite gilt

Theorem Jede p -integrierbare Funktion läßt sich als Limes (in der p -Norm) einer Cauchy-Folge von stetigen Funktionen gewinnen

Um unsere Überlegungen über p -Normen nicht mit (noch ungewohnten) Begriffen aus der Integrationstheorie zu belasten, untersuchen wir die p -Normen im $\mathbb{R}^d$ (für den $\mathbb{C}^d$ geht alles genauso).

Der zentrale Satz in der Theorie der p -Normen ist die

Höldersche Ungleichung Wenn $p \geq 1$, $q \geq 1$ mit $\frac{1}{p} + \frac{1}{q} = 1$, dann gilt für alle d -Tupel x, y

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

Ausführlich geschrieben:

$$\left| \sum_{i=1}^d x_i \cdot y_i \right| \leq \left(\sum_{i=1}^d |x_i|^p \right)^{1/p} \cdot \left(\sum_{i=1}^d |y_i|^q \right)^{1/q} .$$

Satz

a) Wenn $p, q, r \geq 1$ mit $\frac{1}{p} + \frac{1}{q} = \frac{1}{r}$, dann gilt

$$\|x \cdot y\|_r \leq \|x\|_p \cdot \|y\|_q .$$

b) Zu jedem x existiert ein y so, daß Gleichheit gilt.

Lemma Für $a > 0$, $b > 0$ gilt

$$a^{r/p} \cdot b^{r/q} \leq r \left(\frac{1}{p} a + \frac{1}{q} b \right) .$$

1. Beweis des Lemmas Für $u \geq 0$ und $0 \leq \frac{r}{p} \leq 1$ gilt

$$(1 + u)^{r/p} \leq 1 + \frac{r}{p} u .$$

(Beachte: In Bernoullis Ungleichung hat man statt $\frac{r}{p} \leq 1$ den Exponenten $n \in \mathbb{N}$, und da gilt die umgekehrte Ungleichung.)

Der Beweis für unsere Ungleichung ergibt sich aus der Antitonicität von $\frac{1}{x} \ln(1+x)$ für $x > 0$ (Beweis!)

$$\begin{aligned} \left(\frac{r}{p} u \right)^{-1} \ln \left(1 + \frac{r}{p} u \right) &\geq u^{-1} \ln(1 + u) \\ \ln \left(1 + \frac{r}{p} u \right) &\geq \frac{r}{p} \ln(1 + u) \\ 1 + \frac{r}{p} u &\geq (1 + u)^{r/p} \end{aligned}$$

Für $a \geq b > 0$ und $1 + u = \frac{a}{b} \geq 1$ erhalten wir wegen $\frac{1}{q} = \frac{1}{r} - \frac{1}{p}$, $1 - \frac{r}{p} = \frac{r}{q}$

$$\begin{aligned} \left(\frac{a}{b}\right)^{r/p} &\leq 1 + \frac{r}{p} \left(\frac{a}{b} - 1\right) = \left(1 - \frac{r}{p}\right) + \frac{r}{p} \cdot \frac{a}{b} \\ &= \frac{1}{b} \left(\frac{r}{q} b + \frac{r}{p} a\right) \\ a^{r/p} \cdot b^{r/q} &\leq \frac{r}{q} b + \frac{r}{p} a. \end{aligned}$$

(Diese Ungleichung gilt natürlich auch für beliebige $b \geq a \geq 0$.)

2. Beweis des Lemmas Der Logarithmus ist eine konkave Funktion (Beweis!)

$$\begin{aligned} \ln \left(\frac{r}{p} a + \frac{r}{q} b \right) &\geq \frac{r}{p} \ln a + \frac{r}{q} \ln b \\ \frac{r}{p} a + \frac{r}{q} b &\geq a^{r/p} \cdot b^{r/q}. \end{aligned}$$

(Gleichheit gilt nur für $a = b$.)

Bemerkung Für $p = 2 = q$, $r = 1$ erhalten wir

$$\sqrt{ab} \leq \frac{1}{2} (a + b).$$

Beweis des Satzes

a) zu x, y konstruiere für $i = 1, \dots, d$

$$a_i = \frac{|x_i|^p}{|x_1|^p + \dots + |x_d|^p} \quad b_i = \frac{|y_i|^q}{|y_1|^q + \dots + |y_d|^q}.$$

Bemerkung : $\sum a_i = 1 = \sum b_i$.

Das Lemma ergibt für alle i

$$\left(\frac{r}{p} a_i + \frac{r}{q} b_i \right) \geq a_i^{r/p} \cdot b_i^{r/q} = |x_i|^r \cdot |y_i|^r \cdot (\|x\|_p^{-r} \cdot \|y\|_q^{-r}).$$

Summation über $i = 1, \dots, d$ ergibt

$$\begin{aligned} 1 &= \frac{r}{p} + \frac{r}{q} \geq \left(\sum |x_i \cdot y_i|^r \right) (\|x\|_p^{-r} \cdot \|y\|_q^{-r}) \\ (\|xy\|_r)^r &= \sum |x_i y_i|^r \leq \|x\|_p^r \cdot \|y\|_q^r. \end{aligned}$$

b) Gleichheit gilt genau dann, wenn $a_i = b_i$ für alle i , also

$$\left(\frac{|x_i|}{\|x\|_p} \right)^p = \left(\frac{|y_i|}{\|y\|_q} \right)^q \quad \text{für alle } i.$$

Satz (Beispiel für eine Legendre-Transformierte)

Seien $p, q \geq 1$ mit $\frac{1}{p} + \frac{1}{q} = 1$

$$K := \{y : \|y\|_q \leq 1\}$$

$$f(y) = \begin{cases} 0 & \text{für } y \in K \\ +\infty & \text{für } y \notin K. \end{cases}$$

Die Legendretransformierte ist die p -Norm

$$f^*(x) = \|x\|_p.$$

Beweis Wir schreiben die y als Spalten und die x als Zeilen;

$$xy = \sum x_i y_i$$

$$\begin{aligned} f^*(x) &= \sup\{xy - f(y) : y \text{ bel.}\} \\ &= \sup\{xy : y \in K\} = \|x\|_p; \end{aligned}$$

denn für alle $y \in K$ gilt $xy \leq \|x\|_p \cdot \|y\|_q = \|x\|_p$.

Korollar (Minkowskische Ungleichung)

Für alle $p \geq 1$ und alle d -Tupel x, y gilt

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

Beweis Jede Legendretransformierte ist konvex. Offenbar gilt $\|\lambda x\|_p = |\lambda| \cdot \|x\|_p$. Die Funktion $f^*(x) = \|x\|_p$ ist also eine Norm.

Der Beweis der Hölderschen Ungleichung überträgt sich sofort auch auf den kontinuierlichen Fall.

Sei wieder $p, q, r \geq 1$ mit $\frac{1}{p} + \frac{1}{q} = \frac{1}{r}$; und seien $x(\cdot)$ und $y(\cdot)$ (stetige) Funktionen of $[0, 1]$. Setze

$$a(t) = \left(\frac{|x(t)|}{\|x\|_p} \right)^p; \quad b(t) = \left(\frac{|y(t)|}{\|y\|_q} \right)^q,$$

wobei

$$\|x\|_p = \left(\int_0^1 |x(t)|^p \right)^{1/p}, \quad \|y\|_q = \left(\int_0^1 |y(t)|^q \right)^{1/q}.$$

Das Lemma ergibt die punktweise Ungleichung

$$\frac{r}{p} a(\cdot) + \frac{r}{q} b(\cdot) \geq |x(\cdot)|^r \cdot |y(\cdot)|^r \cdot (\|x\|_p^{-r} \cdot \|y\|_q^{-r}) .$$

Integration liefert

$$1 \geq \int_0^1 |x(s)y(s)|^r ds (\|x\|_p^{-r} \cdot \|y\|_q^{-r}) ,$$

Ziehen der r -ten Wurzel führt auf die Höldersche Ungleichung

$$\|x \cdot y\|_r \leq \|x\|_p \cdot \|y\|_q .$$

Hinweis Die Ungleichung gilt in der Tat für **beliebige Funktionen** $x(\cdot), y(\cdot)$, für die die auftretenden Integrale wohldefiniert und endlich sind.

Die Höldersche Ungleichung begründet die „Dualität“ der Räume $\mathcal{L}^p$ und $\mathcal{L}^q$ ($p, q \geq 1, \frac{1}{p} + \frac{1}{q} = 1$). Wir wollen diese **Dualität** kurz vorstellen, um bei dieser Gelegenheit klarzumachen, was die Lebesguesche Integrationstheorie leisten kann.

- 1) Sei $g(\cdot)$ eine Funktion auf $[0, 1]$, die zunächst als stetig angenommen werden mag. Wir betrachten dazu

$$C([0, 1]) \ni f \longmapsto \lambda_g(f) := \int_0^1 g(t)f(t) dt .$$

Dies ist eine gleichmäßig stetige Funktion auf dem mit der p -Norm ausgestatteten Vektorraum $C([0, 1])$, denn

$$|\lambda_g(f_1) - \lambda_g(f_2)| \leq \|g\|_q \cdot d_p(f_1, f_2) .$$

Man kann $\lambda_g(\cdot)$ also auf $\mathcal{L}^p$ stetig fortsetzen. In der Lebesgueschen Integrationstheorie lernt man, daß man auch die fortgesetzte Funktion $\bar{\lambda}_g(\cdot)$ als Integral schreiben kann

$$\bar{\lambda}_g(f) = \int_0^1 g(t)f(t) dt \quad \text{für } f \in \mathcal{L}^p .$$

- 2) Sei $(g_n)_n$ eine $d_q(\cdot, \cdot)$ -Cauchy-Folge von stetigen Funktionen. Die zugehörigen $\lambda_{g_n}(\cdot)$ konvergieren dann gleichmäßig auf der Einheitskugel des $\mathcal{L}^p$.

$$|\lambda_{g_n}(f) - \lambda_{g_m}(f)| \leq \|g_n - g_m\|_q \quad \text{für alle } f \text{ mit } \|f\|_q \leq 1 .$$

Es existiert also ein gleichmäßiger Limes $\lambda^*(\cdot)$ auf der Einheitskugel; dieser setzt sich zu einer stetigen linearen Funktion auf $\mathcal{L}^p$ fort. In der Lebesgueschen Integrationstheorie lernt man, daß eine geeignete Teilfolge der Folge $(g_n)_n$ fastüberall gegen eine Grenzfunktion g^* konvergiert, und daß die Cauchy-Folge $(g_n)_n$ selbst in der q -Norm gegen g^* konvergiert. Man lernt auch, daß man schreiben kann

$$\lambda^*(f) = \int g^*(s)f(s) ds \quad \text{für alle } f \in \mathcal{L}^p.$$

- 3) Jede stetige lineare Funktion $\lambda(\cdot)$ auf dem $\mathcal{L}^p$ ($p < \infty$) ist von einer q -integrablen Funktion $g(\cdot)$ induziert

$$\lambda(f) = \int g(t)f(t) dt \quad \text{für alle } f \in \mathcal{L}^p.$$

Didaktischer Hinweis

Viele elementare Texte beschäftigen sich mit der Integration (von Funktionen auf dem Einheitsintervall) in einer adhoc-Manier. Man vermeidet den Begriff der Fastüberallgleichheit und die dazugehörige Äquivalenzrelation. Es werden also nicht Äquivalenzklassen integriert sondern mehr oder weniger regelmäßige Funktionen. In der sog. Riemannschen Integrationstheorie benutzt man ein Ausschöpfungsargument mit einfachen Treppenfunktionen; in der Integrationstheorie für Regelfunktionen (z.B. bei BARNER-FLOHR *Analysis I*) arbeitet man gar nur mit den Funktionen, die sich als gleichmäßige Limiten einer Folge einfacher Treppenfunktionen darstellen lassen. Obwohl man bei diesen primitiven Integrationstheorien keinerlei innere Notwendigkeiten aufzeigen kann, und die wirklich starken Sätze (wie der Satz von der majorisierten Konvergenz) nicht zur Verfügung gestellt werden können, wird behauptet, daß der Anfänger hier ein ausreichendes Verständnis dafür gewinnen kann, was Integrationstheorie zu leisten hat. Wir bezweifeln das ganz entschieden; ein wirksamer Integrationsbegriff erfordert die entschlossene Abwendung von allen denjenigen Eigenschaften der Integranden, die irgendwie an Stetigkeit und gleichmäßige Konvergenz erinnern. (Verbindungen zu topologischen Begriffen kann es allenfalls über die σ -Additivität zur Kompaktheit geben (siehe *Analysis II*).)

Unsere Bemerkungen zu den p -Normen zum Lebesguemaß auf dem Einheitsintervall sollten zeigen, daß man ein gründliches Verständnis der Integration braucht, um die Vervollständigungsargumente konkret zu untermauern. Mit halbherziger Integrationstheorie kommt man da auf keinen grünen Zweig.

C.9 Fourier–Reihen aus funktionalanalytischer Sicht

Am Beispiel des Begriffs Theorie der Fourier–Reihe wurde klar, daß der Funktionsbegriff des 18. Jahrhunderts zu eng ist. Dirichlet und Riemann waren die Pioniere bei der Frage nach einem allgemeinen Funktionsbegriff; Dirichlet zunächst im Hinblick auf die mathematische Physik, Riemann dann auch für die Zahlentheorie. Schon Euler hatte eine Fourier–Reihe angegeben, die punktweise gegen eine unstetige Funktion konvergiert, nämlich

$$\sin x + \frac{1}{2} \sin x + \frac{1}{3} \sin 3x + \dots$$

Er wußte damit aber nicht viel anzufangen, wohl hauptsächlich deshalb, weil er keine tragfähige Theorie der Konvergenz der Funktionenfolgen besaß; er arbeitete in einer völlig intuitiven Weise mit divergenten Reihen.

Wir wollen hier die Theorie der Fourier–Reihen vom Standpunkt der Pioniere der Funktionalanalysis aus ein Stück weit entwickeln. Wichtige Beiträge zu dieser Theorie gehen auf F. RIESZ (1880–1956) zurück. Diese Theorie beginnt man wohl am besten mit den trigonometrischen Polynomen.

Einer endlichen Summe

$$\sum c_n e^{inx} \quad \text{oder} \quad \frac{a_0}{2} + \sum (a_k \cos kx + b_k \sin kx)$$

ordnet man eine Norm zu, nämlich

$$\sqrt{\sum |c_n|^2}.$$

Satz Wenn

$$f(x) = \sum c_n e^{inx},$$

dann gilt

$$\frac{1}{2\pi} \int_0^{2\pi} |f(x)|^2 dx = \sum |c_n|^2.$$

Beweis

$$\begin{aligned}
\frac{1}{2\pi} \int_0^{2\pi} e^{-imx} e^{inx} dx &= \begin{cases} 1 & \text{wenn } m = n \\ 0 & \text{wenn } m \neq n \end{cases} \\
\frac{1}{2\pi} \int_0^{2\pi} \left| \sum c_n e^{inx} \right|^2 dx &= \frac{1}{2\pi} \int_0^{2\pi} \left(\sum \bar{c}_n e^{-inx} \right) \left(\sum c_m e^{imx} \right) dx \\
&= \sum_{m,n} \bar{c}_n \cdot c_m \frac{1}{2\pi} \int_0^{2\pi} e^{-inx} e^{imx} dx = \sum |c_n|^2 .
\end{aligned}$$

Definition Wenn f und g trigonometrische Polynome sind, dann definiert man

$$s(g, f) = \frac{1}{2\pi} \int_0^{2\pi} \bar{g}(x) \cdot f(x) dx .$$

Bemerke

$$f = \sum c_n e^{inx} , \quad g = \sum d_m e^{imx} \implies s(g, f) = \sum_n \bar{d}_n c_n .$$

Satz („Inneres Produkt“)

- (i) $s(f, g) = \overline{s(g, f)}$
- (ii) $s(f + g, h) = s(f, h) + s(g, h)$
 $s(f, g + h) = s(f, g) + s(f, h)$
- (iii) $s(f, \alpha g) = \alpha s(f, g)$
 $s(\alpha f, g) = \bar{\alpha} s(f, g)$

Der Beweis ist trivial.

Sprechweise

- a) Wenn V ein komplexer Vektorraum ist und $s(\cdot, \cdot)$ auf $V \times V$ die Eigenschaften (i), (ii), (iii) hat, dann nennt man $s(\cdot, \cdot)$ eine **hermitesche Sequilinearform**.

b) Wenn V ein reeller Vektorraum ist und $b(\cdot, \cdot)$ auf $V \times V$ die folgenden Eigenschaften hat

$$(i) \quad b(f, g) = b(g, f)$$

$$(ii) \quad b(f, g + h) = b(f, g) + b(f, h)$$

$$(iii) \quad b(\alpha f) = \alpha b(f),$$

dann nennt man $b(\cdot, \cdot)$ eine **symmetrische Bilinearform**.

c) Eine hermitesche Sequilinearform auf einem komplexen Vektorraum heißt **nichtnegativ**, wenn

$$s(f, f) \geq 0 \quad \text{für alle } f,$$

sie heißt **positiv**, wenn

$$s(f, f) > 0 \quad \text{für alle } f \neq 0.$$

d) Eine symmetrische Bilinearform auf einem reellen Vektorraum heißt **nichtnegativ**, wenn

$$b(f, f) \geq 0 \quad \text{für alle } f;$$

sie heißt **positiv**, wenn

$$s(f, f) > 0 \quad \text{für alle } f \neq 0.$$

Satz Wenn $s(\cdot, \cdot)$ eine positive Sequilinearform ist, dann ist

$$\|f\| = \sqrt{s(f, f)}$$

eine Norm und für diese Norm gilt

$$\|f + g\|^2 + \|f - g\|^2 = 2(\|f\|^2 + \|g\|^2) \quad (\text{Parallelogrammidentität}).$$

Beweis

1) Nachzuweisen ist

$$(i) \quad \|f\| = 0 \iff f = 0$$

$$(ii) \quad \|\alpha f\| = |\alpha| \cdot \|f\| \quad \text{für alle Skalare } \alpha$$

$$(iii) \quad \|f + g\| \leq \|f\| + \|g\| ;$$

(i) und (ii) sind offensichtlich. Zum Beweis von (iii) bemerken wir

$$\begin{aligned} \|f + g\|^2 &= s(f + g, f + g) = s(f, f) + s(g, g) + s(f, g) + s(g, f) \\ &= \|f\|^2 + \|g\|^2 + 2 \cdot \Re s(f, g) \end{aligned}$$

Es genügt also nachzuweisen

$$|s(f, g)| \leq \|f\| \cdot \|g\| .$$

Den Nachweis dieser Ungleichung führt man mit einem Trick, der die Positivität ausnützt.

Für alle $\lambda \in \mathbb{R}$ gilt

$$\begin{aligned} 0 &\leq s(f + \lambda g, f + \lambda g) = \|f\|^2 + \lambda^2 \|g\|^2 + 2\lambda \cdot \Re(s(f, g)) \\ &= a + 2\lambda b + \lambda^2 c . \end{aligned}$$

Eine quadratische Funktion

$$\lambda \mapsto a + 2\lambda b + \lambda^2 c$$

ist genau dann nichtnegativ für alle $\lambda \in \mathbb{R}$, wenn

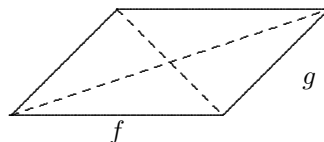
$$a \geq 0, \quad c \geq 0 \quad \text{und} \quad b^2 - ac \leq 0 .$$

In unserem Falle ergibt das

$$|s(f, g)|^2 \leq \|f\|^2 \cdot \|g\|^2 .$$

$$\begin{aligned} 2) \quad \|f + g\|^2 + \|f - g\|^2 &= s(f + g, f + g) + s(f - g, f - g) \\ &= 2 \cdot s(f, f) + 2 \cdot s(g, g) = 2 \cdot \|f\|^2 + 2 \cdot \|g\|^2 \end{aligned}$$

Bemerke Der Name Parallelogrammidentität kommt aus der euklidischen Geometrie



Die Quadrate der Längen der Diagonalen summieren sich zum Doppelten der Summe der Quadrate der Seitenlängen. Wenn die Seiten aufeinander senkrecht stehen, ist das der Satz von Pythagoras. Im allgemeinen Fall ergibt sich die Behauptung aus dem sog. Cosinussatz der euklidischen Geometrie

$$\|f + g\|^2 = \|f\|^2 + \|g\|^2 + 2 \cdot \|f\| \cdot \|g\| \cdot \cos(\angle).$$

Der von f und g eingeschlossene Winkel und der von f und $-g$ eingeschlossene Winkel summieren sich zu π ; der Cosinus des eingeschlossenen Winkels wechselt also einfach das Vorzeichen.

Kehren wir zum Vektorraum der trigonometrischen Polynome zurück. Betrachten wir ihn als einen metrischen Raum $(E, d(\cdot, \cdot))$

$$d(f, g) = \|g - f\| = \left[\frac{1}{2\pi} \int_0^{2\pi} |g(x) - f(x)|^2 dx \right]^{1/2}.$$

Diesen metrischen Raum können wir vervollständigen. Nach dem abstrakten Satz von der Vervollständigung können wir die Punkte von $(\overline{E}, \overline{d}(\cdot, \cdot))$ mit Äquivalenzklassen von Cauchy–Folgen von trigonometrischen Polynomen identifizieren. Im konkreten Fall können wir die Punkte der Vervollständigung aber auch noch auf andere Weise beschreiben.

a) Ein trigonometrisches Polynom $f(\cdot)$ entspricht einer Folge komplexer Zahlen

$$c = (\dots, c_{-1}, c_0, c_1, c_2, \dots),$$

wo nur endlich viele Zahlen ungleich 0 sind.

$$d(c, d) = \left(\sum |d_k - c_k|^2 \right)^{1/2}$$

Die Vervollständigung ergibt die Menge aller quadrat–summierbaren Folgen (vgl. Übungsaufgabe).

b) Ein trigonometrisches Polynom entspricht einer 2π –periodischen Funktion.

Es gilt der

Satz Wenn $(f_n)_n$ eine Cauchy-Folge trigonometrischer Polynome ist, dann existiert eine quadrat-integrierbare 2π -periodische Funktion $\tilde{f}$ mit

$$\lim_{n \rightarrow \infty} \frac{1}{2\pi} \int_0^{2\pi} \|f_n - \tilde{f}\|^2 dx = 0 .$$

Den Beweis können wir hier nicht führen; er braucht die Lebesguesche Integrationstheorie. Hier sollte der Hinweis genügen, daß man $\tilde{f}$ als den Fastüberall-Limes einer geeigneten Teilfolge gewinnt. (Satz von Fischer-Riesz)

Hinweise zur Abrundung des Bildes

- 1) Die Punkte der Vervollständigung sind nach dem obigen Satz Äquivalenzklassen quadratisch integrierbarer Funktionen; die Äquivalenz ist die Gleichheit bis auf eine Lebesgue-Nullmenge. Es zeigt sich, daß jede quadrat-integrierbare Funktion im $\mathcal{L}^2$ -Abstand durch trigonometrische Polynome beliebig gut approximiert werden kann. Die Vervollständigung ist daher gleich der Menge aller (Äquivalenzklassen) quadrat-integrierbarer Funktionen.
- 2) Eine 2π -periodische Funktion $\tilde{f}$, deren Quadrat über $[0, 2\pi]$ integrierbar ist

$$\frac{1}{2\pi} \int_0^{2\pi} |\tilde{f}(x)|^2 dx < \infty$$

ist auch integrierbar über $[0, 2\pi]$

$$\int |\tilde{f}(x)| dx < \infty .$$

Man kann daher zu einer solchen Funktion die Folge der Fourier-Koeffizienten definieren

$$\tilde{c}_n := \frac{1}{2\pi} \int_0^{2\pi} \tilde{f}(x) e^{-inx} dx \quad \text{für } n \in \mathbb{Z} .$$

Die Funktion $\tilde{f}(x)$ ist (bis auf Gleichheit fast überall) durch die Folge der Fourier-Koeffizienten eindeutig bestimmt. Dies ergibt sich aus der Charakterisierung 1) der Punkte der Vervollständigung.

Man schreibt

$$\tilde{f} \sim \sum_{-\infty}^{+\infty} \tilde{c}_n e^{inx}$$

3) Es gilt für die Partialsummen

$$f_N(x) = \sum_{-N}^{+N} \tilde{c}_n e^{inx}$$

$$\|\tilde{f} - f_N\|^2 = \sum_{n: |n| > N} |\tilde{c}_n|^2 \longrightarrow 0 \quad \text{für } N \rightarrow \infty .$$

Diese Faktum kann man so zum Ausdruck bringen:

„Die Folge der Partialsummen der formalen Fourier–Reihe konvergiert für jedes quadrat–integrierbare $\tilde{f}$ im quadratischen Mittel gegen $\tilde{f}$.“

Dies ist ein Satz, den sich jeder Student der Mathematik einprägen sollte.

4) Wer die exakten Beweise wünscht, der möge in irgendeinem modernen Buch über reelle Analysis (z.B. „Real and Abstract Analysis“ von E. Hewitt und K. Stromberg) nachsehen unter dem Stichwort „Riesz–Fischer Theorem“.

Die altmodischen Bücher über reelle Analysis enthalten in der Regel einiges Material zu der (eigentlich überholten) Frage, für welche speziellen quadrat–integralen Funktionen die Folge der Partialsummen der Fourier–Reihe punktweise gegen $\tilde{f}$ konvergiert. Ein Stichwort ist etwa die Dirichletsche Regel (siehe Heuser, Lehrbuch der Analysis Teil 2, S. 138 ff).

Nach dem Satz von L. CARLESON (1966) herrscht für jede quadrat–integrable Funktion Fastüberallkonvergenz. Dieser schwierige Satz wird aber nur in der Spezialliteratur über Fourier Reihen bewiesen.