
Mathematische Methoden für das Lehramt L3

Hendrik van Hees

3. Oktober 2024

Inhaltsverzeichnis

Inhaltsverzeichnis	3
1 Analysis für reelle Funktionen einer Variablen	7
1.1 Allgemeine Hinweise zur Vorlesung	7
1.2 Mengen und reelle Zahlen	8
1.3 Folgen und Grenzwerte	11
1.4 Satz vom Supremum und Infimum	18
1.5 Lineare und quadratische Gleichungen	19
1.6 Funktionen und Stetigkeit	21
1.7 Differentialrechnung für Funktionen einer reellen Veränderlichen	22
1.7.1 Definition der Ableitung einer Funktion	22
1.7.2 Formeln zur Ableitung	23
1.7.3 Trigonometrische Funktionen	26
1.7.4 Kurvendiskussionen, Mittelwertsatz der Differentialrechnung	32
1.8 Die Regeln von de L'Hospital	34
1.9 Integralrechnung	35
1.9.1 Definition des Riemann-Integrals	35
1.9.2 Der Mittelwertsatz der Integralrechnung	38
1.9.3 Der Hauptsatz der Analysis	39
1.9.4 Definition des natürlichen Logarithmus' und der Exponentialfunktion	40
1.9.5 Hyperbelfunktionen	43
1.9.6 Integrationstechniken	45
1.9.7 Funktionenfolgen und -reihen	47
1.9.8 Taylor-Entwicklung und Potenzreihen	50
1.10 Die strikte Definition der trigonometrischen Funktionen	55
2 Lineare Algebra	59
2.1 Geometrische Einführung von Euklidischen Vektoren	59
2.1.1 Definition von Vektoren als Verschiebungen	59
2.1.2 Vektoraddition	60
2.1.3 Länge (Norm) von Vektoren	61
2.1.4 Lineare Unabhängigkeit von Vektoren und Basen	62
2.1.5 Der Vektorraum $\mathbb{R}^3$	64
2.1.6 Basistransformationen	65

2.1.7	Das Skalarprodukt	67
2.1.8	Geometrische Anwendungen des Skalarprodukts	69
2.1.9	Das Skalarprodukt im $\mathbb{R}^3$	70
2.2	Axiomatische Begründung der linearen Algebra und Geometrie	72
2.3	Kartesische Basen und orthogonale Transformationen	75
2.4	Das Vektorprodukt	78
2.5	Das Spatprodukt	80
2.6	Lineare Gleichungssysteme und Determinanten	82
2.6.1	Lineare Gleichungssysteme	82
2.6.2	Determinanten als Volumenform	87
2.6.3	Determinanten von Matrizen	88
2.6.4	Transformationsverhalten des Kreuzprodukts	93
2.7	Drehungen	95
2.7.1	Drehungen in der Ebene	95
2.7.2	Drehungen im Raum	98
2.7.3	Euler-Winkel	98
2.8	Lineare Abbildungen	100
2.9	Eigenwertprobleme	101
2.9.1	Beispiel: Durch Feder verbundene Massen	101
2.9.2	Allgemeine Eigenschaften von Eigenvektoren und Eigenwerten	103
2.9.3	Symmetrische reelle Matrizen und die Hauptachsentransformation	106
2.10	Vektoren und Tensoren	111
3	Vektoranalysis	117
3.1	Kurven in der Ebene	117
3.2	Kegelschnitte	120
3.2.1	Definition der Kegelschnitte	120
3.2.2	Ellipse	121
3.2.3	Hyperbel	123
3.2.4	Parabel	124
3.3	Raumkurven und Frenetsche Formeln	125
3.3.1	Anwendung auf die Bewegung eines Partikels	129
3.4	Skalare Felder, Gradient und Richtungsableitung	132
3.5	Extremwertaufgaben	136
3.6	Vektorfelder, Divergenz und Rotation	138
3.7	Krummlinige Orthogonalkoordinaten	141
3.7.1	Kugelkoordinaten	141
3.7.2	Zylinderkoordinaten	142
3.7.3	Allgemeine krummlinige Orthogonalkoordinaten	142
3.7.4	Die Differentialoperatoren in krummlinigen Orthogonalkoordinaten	143
3.7.5	Polarkoordinaten in der Ebene	147
3.8	Potentialfelder	149
3.9	Wegintegrale und Potentialfelder	151

3.10	Flächenintegrale und der Stokessche Integralsatz	154
3.10.1	Orientierte Flächen im Raum	154
3.10.2	Definition des Flächenintegrals	155
3.10.3	Unabhängigkeit des Flächenintegrals von der Parametrisierung	157
3.10.4	Koordinatenunabhängige Definition der Rotation	159
3.10.5	Der Integralsatz von Stokes	160
3.10.6	Der Greensche Satz in der Ebene	162
3.11	Das Poincaré-Lemma	162
3.12	Volumenintegrale, Divergenz und Gaußscher Integralsatz	165
3.12.1	Definition des Volumenintegrals	166
3.12.2	Die koordinatenunabhängige Definition der Divergenz	167
3.12.3	Der Gaußsche Integralsatz	168
3.12.4	Die Greenschen Integralsätze im Raum	168
3.13	Alternative Herleitung der Differentialoperatoren in krummlinigen Orthogonalkoordinaten	169
3.14	Solenoidfelder und Vektorpotentiale	171
3.15	Die Poisson-Gleichung und Green-Funktionen	173
3.16	Der Helmholtzsche Zerlegungssatz der Vektoranalysis	178
3.16.1	Bestimmung des Potentialfeldanteils	178
3.16.2	Bestimmung des Solenoidfeldanteils	179
3.17	Verallgemeinerte Funktionen (Distributionen)	180
3.18	Transporttheoreme	184
3.18.1	Transporttheorem für Wegintegrale	185
3.18.2	Transporttheorem für Flächenintegrale	186
3.18.3	Reynoldssches Transporttheorem für Volumenintegrale	187
3.19	Allgemein kovariante Formulierung der Vektoranalysis	188
3.19.1	Transformationsverhalten von Vektoren und Tensoren	188
3.19.2	Die Differentialoperatoren grad, div und rot	192
3.19.3	Kovariante Ableitungen	194
3.19.4	Alternierende Formen und Differentialformen	198
3.19.5	Alternierende Differentialformen und der Stokessche Integralsatz	200
3.19.6	Hodge-Dualisierung und Integrale über kontravariante Tensorfeldkomponenten . . .	201
4	Komplexe Zahlen	203
4.1	Definition der komplexen Zahlen	203
4.2	Potenzreihen	205
5	Gewöhnliche Differentialgleichungen	209
5.1	Differentialgleichungen 1. Ordnung	210
5.1.1	Separierbare Differentialgleichungen	210
5.1.2	Homogene Differentialgleichungen 1. Ordnung	212
5.1.3	Exakte Differentialgleichung 1. Ordnung	213
5.1.4	Homogene lineare Differentialgleichung 1. Ordnung	213
5.1.5	Inhomogene lineare Differentialgleichung 1. Ordnung	214

5.2	Lineare Differentialgleichungen 2. Ordnung	214
5.3	Der ungedämpfte harmonische Oszillator	217
5.4	Der gedämpfte harmonische Oszillator	221
5.4.1	Schwingfall ($\omega_0 > \gamma$)	222
5.4.2	Kriechfall ($\omega_0 < \gamma$)	223
5.4.3	Aperiodischer Grenzfall ($\omega_0 = \gamma$)	224
5.5	Der getriebene gedämpfte Oszillator	224
5.5.1	Spezielle Lösung der inhomogenen Gleichung	225
5.5.2	Amplitudenresonanzfrequenz	226
5.5.3	Energieresonanz	227
5.5.4	Lösung des Anfangswertproblems	228
5.5.5	Resonant angetriebener ungedämpfter Oszillator	229
5.5.6	Allgemeine äußere Kräfte und die δ -Distribution	230
5.6	Differentialgleichungen der Fuchsschen Klasse	234
A	Vollständige Induktion	241
A.1	Das Beweisprinzip	241
A.2	Beispiele	242
A.2.1	Der „kleine Gauß“	242
A.2.2	Höhere Summenformeln	243
A.2.3	Geometrische Reihe	243
	Literaturverzeichnis	247

Kapitel 1

Analysis für reelle Funktionen einer Variablen

Nach einigen allgemeinen Hinweisen zur Vorlesung stellen wir in diesem Kapitel einige Grundlagen zusammen, die wir im Folgenden voraussetzen wollen. Dies umfasst den Umgang mit Mengen und reellen Zahlen sowie die Analysis für Funktionen einer reellen Veränderlichen. Im Bedarfsfall sollte dieser Stoff der Schulmathematik auch **selbständig nachgearbeitet** werden [Hef18].

1.1 Allgemeine Hinweise zur Vorlesung

Dies ist das Manuskript zur Vorlesung „Mathematische Methoden für Lehramt L3“. Ziel dieser Vorlesung ist es, die in den Vorlesungen „Theoretische Physik 1–3 für das Lehramt L3“ benötigten mathematischen Methoden zu erarbeiten aber auch vor allem gleich auf konkrete physikalische Probleme anzuwenden. Der Schwerpunkt liegt entsprechend weniger auf formalen Beweisen als vielmehr auf der Vermittlung der Rechen-technik, die sehr wichtig für das Verständnis der theoretischen Physik ist.

Erfahrungsgemäß sind die Vorkenntnisse der Studierenden aus der Schulphysik recht heterogen. Daher beginnen wir die Vorlesung mit einer kurzen Zusammenfassung der „Schulmathematik“, also mit den grundlegenden Begriffen der **reellen Zahlen** und der **Differential- und Integral-Rechnung für Funktionen einer reellen Variablen**, sowie der reellen **Vektoralgebra** und **analytischen Geometrie**.

Bereits die Vorlesung „Theorie 1“ über die Newtonsche Mechanik erfordert die Erweiterung der Methoden der **linearen Algebra** zu solchen der eigentlichen **Vektor-Analysis**. Entsprechend werden wir über die physikalischen Größen Ort, Geschwindigkeit und Beschleunigung die Ableitung von Vektorfunktionen nach äußeren Parametern (hier naturgemäß der Zeit) einführen und geometrisch deuten.

Ebenso werden die wesentlichen Grundbegriffe der **Feldtheorie** eingeführt, wie die Ableitungen **skalarer Felder** und **Vektorfelder** (grad, div und rot), die durch den „Nabla-Operator“ $\vec{\nabla}$ kompakt dargestellt werden. Ebenso besprechen wir auch ausführlich das Rechnen mit Komponenten, den sog. **Ricci-Kalkül**. Dieser Teil wird allerdings erst in der Vorlesung „Theorie 2“ über Elektrodynamik benötigt, sind zum Teil aber auch bereits für die Mechanik von einigem Nutzen.

Den zweiten Schwerpunkt der Vorlesung bilden Techniken zur Lösung von **gewöhnlichen Differentialgleichungen**, also die Integration der einfachsten Typen von Bewegungsgleichungen, wie sie in der klassischen Mechanik auftreten, insbesondere das wichtige Beispiel des **harmonischen Oszillators**. Dazu führen wir auch **komplexe Zahlen** ein und besprechen die wichtigsten **elementaren Funktionen** wie Polynome, die Exponentialfunktion, die trigonometrischen Funktionen und deren Umkehrungen.

Es sei betont, dass in der Vorlesung nicht notwendig alle Inhalte dieses Manuskripts abgearbeitet werden müssen. Der Inhalt der Vorlesung richtet sich nicht zuletzt auch nach den Bedürfnissen der Hörerinnen und Hörer.

Literaturempfehlungen: Zum Auffrischen bzw. Nachholen der Schulphysik empfehlen sich Bücher, die für

Mathematikvorkurse an Universitäten erarbeitet wurden. Ein relativ neues Buch ist [Hef18], zu dem es auch eine offen verfügbare Online-Plattform gibt.

Die Literatur zum Thema „Mathematik für Physiker“ ist nahezu unerschöpflich. Für die Zielsetzung dieser Vorlesung ist besonders [Gro12] (frei als e-book im Netz der Goethe-Universität verfügbar!) zu empfehlen, das den mathematischen Teil des Vorlesungsstoffes weitgehend abdeckt. Ein sehr ausführliches Mathematikbuch für Studierende der Naturwissenschaften ist [AHK⁺18] (frei als e-book im Netz der Goethe-Universität verfügbar!). Eine sehr anschauliche Behandlung der Vektoranalysis bieten [BK88, SH99].

Auch viele Theorie-Lehrbücher bieten einen Überblick über die zur Anwendung kommende Mathematik. Als besonders gelungen empfinde ich die älteren Bücher [Joo89, Sau73, Som92].

Generell orientiere ich mich für die Theorie-Vorlesungen auch an der neuen Lehrbuchreihe über theoretische Physik [BFK⁺18a, BFK⁺18b, BFK⁺18c, BFK⁺18d], wovon im Lehramtsstudiengang naturgemäß nur Teile der Bände 1-3 abgehandelt werden können (frei als e-book im Netz der Goethe-Universität verfügbar!). Hier gibt es allerdings kein eigenständiges Mathe-Kapitel, sondern die benötigten mathematischen Methoden werden im Text in „Mathe-Boxen“ bereitgestellt, wo sie benötigt werden.

Ein Standardlehrbuch mit einem einleitenden Kapitel über die mathematischen Grundlagen ist noch [Nol18] (frei als e-book im Netz der Goethe-Universität verfügbar!).

Es sei auch mit Nachdruck darauf hingewiesen, dass eine ausführliche Beschäftigung mit Übungsaufgaben, auch über das Maß des zur Vorlesung gehörigen Tutoriums hinaus, sehr wichtig ist. Hierzu gibt es auch sehr viele Bücher, z.B. [HH19a, HH19b, HH19c] (frei als e-book im Netz der Goethe-Universität verfügbar!). Auch diese Buchreihe bietet ein umfangreiches Zusatzmaterial online.

1.2 Mengen und reelle Zahlen

Die moderne Mathematik versteht sich als die Lehre über abstrakte Strukturen und basiert (nahezu) vollständig auf der Mengenlehre, die wir hier in einer sehr naiven Auffassung verwenden. Demnach ist eine **Menge** einfach eine Zusammenfassung irgendwelcher (realer oder abstrakter) Gegenstände. Eine Menge M ist demnach dadurch definiert, dass man von beliebigen Gegenstand x sagen kann, ob er zur Menge gehört ($x \in M$: „ x ist in M enthalten“ oder „ x ist Element der Menge M “) oder nicht ($x \notin M$).

Eine Menge M kann zum einen durch einfache Aufzählung der in ihr enthaltenen Elemente definiert werden. Z.B. definiert man als die Menge der natürlichen Zahlen $\mathbb{N} = \{1, 2, 3, \dots\}$ oder $\mathbb{N}_0 = \{0, 1, 2, 3, \dots\}$. Sie kann aber auch durch die Eigenschaft ihrer Elemente definiert sein. Z.B. kann man die Menge aller natürlichen Zahlen, die kleiner als 7 sind zum einen einfach durch Aufzählung $M = \{1, 2, 3, 4, 5, 6\}$ oder durch die betreffende Eigenschaft der Elemente $M = \{n \in \mathbb{N} | n < 7\} = \{n \in \mathbb{N} | n \leq 6\}$, was „alle natürlichen Zahlen n mit der Eigenschaft $n < 7$ bzw. $n \leq 6$ “ bedeutet, definiert werden.

Einige nützliche Notationen sind noch die **Teilmenge** bzw. die **Obermenge**. Man sagt, eine Menge M' ist Teilmenge der Menge M , $M' \subseteq M$, wenn aus $x \in M'$ stets folgt, dass auch $x \in M$ ist. Man sagt in diesem Fall auch, dass M Obermenge von M' ist, $M \supseteq M'$.

Auch das Rechnen mit reellen Zahlen wollen wir als bekannt voraussetzen. Wir deuten nur kurz einige Grundlagen an. Die reellen Zahlen werden im wesentlichen durch die Rechenregeln für Addition und Multiplikation und deren jeweilige Umkehrungen, also Subtraktion und Division definiert. Diese algebraischen Regeln definieren, was die Mathematiker als **Zahlenkörper** bezeichnen.

Zunächst bildet die Menge der **reellen Zahlen** $\mathbb{R}$ zusammen mit der Addition als **Abbildung** zweier reeller Zahlen auf eine reelle Zahl eine **Abelsche Gruppe**, d.h. es gelten die folgenden „Rechenregeln“.

1. Für alle $a, b, c \in \mathbb{R}$ gilt für die Addition stets das **Assoziativgesetz**

$$(a + b) + c = a + (b + c). \quad (1.2.1)$$

2. Es existiert genau eine Zahl $0 \in \mathbb{R}$, so dass für alle $a \in \mathbb{R}$

$$a + 0 = a \quad (1.2.2)$$

gilt. Diese Zahl 0 (Null) ist das **neutrale Element der Addition**.

3. Zu jeder Zahl $a \in \mathbb{R}$ existiert eine Zahl $(-a) \in \mathbb{R}$, so dass

$$a + (-a) = 0. \quad (1.2.3)$$

Man nennt $(-a)$ das inverse Element zu a bzgl. der Addition.

4. Für alle $a, b \in \mathbb{R}$ gilt für die Addition das **Kommutativgesetz**

$$a + b = b + a. \quad (1.2.4)$$

Man schreibt abkürzend auch $a + (-b) = a - b$, d.h. die Subtraktion ist auf die Addition und Bildung des inversen Elements bzgl. der Addition zurückgeführt.

Es gibt noch eine weitere elementare Verknüpfung, die **Multiplikation** (Abbildung zweier reeller Zahlen auf eine reelle Zahl). Auf der Menge $\mathbb{R}^* = \mathbb{R} \setminus \{0\}$ (die reellen Zahlen ohne die Null) bildet auch diese Verknüpfung eine Abelsche Gruppe.

1. Für alle $a, b, c \in \mathbb{R}$ gilt für die Multiplikation stets das **Assoziativgesetz**

$$(ab)c = a(bc). \quad (1.2.5)$$

2. Es existiert genau eine Zahl $1 \in \mathbb{R}$, so dass für alle $a \in \mathbb{R}$

$$a1 = a \quad (1.2.6)$$

gilt. Diese Zahl 1 (Eins) ist das **neutrale Element der Multiplikation**.

3. Zu jeder Zahl $a \in \mathbb{R}^*$ existiert genau eine Zahl $a^{-1} \in \mathbb{R}$, so dass

$$a(a^{-1}) = 1 \quad (1.2.7)$$

ist. Man nennt a^{-1} das inverse Element zu a bzgl. der Multiplikation. Die Null besitzt kein inverses Element bzgl. der Multiplikation („durch Null kann man nicht teilen!“).

4. Für alle $a, b \in \mathbb{R}$ gilt für die Multiplikation das **Kommutativgesetz**

$$ab = ba. \quad (1.2.8)$$

Die Multiplikation mit dem Inversen bezeichnet man auch als Division und definiert für alle $a \in \mathbb{R}$ und alle $b \in \mathbb{R}^*$

$$a(b^{-1}) = \frac{a}{b} = a/b. \quad (1.2.9)$$

Weiter gilt noch für die Verknüpfung der beiden Rechenoperationen das **Distributivgesetz**, d.h. für alle $a, b, c \in \mathbb{R}$ gilt

$$a(b + c) = ab + ac. \quad (1.2.10)$$

Man kann mit diesen Axiomen alle weiteren algebraischen Rechenregeln herleiten.

Die reellen Zahlen sind dadurch aber noch nicht vollständig charakterisiert. Als weiteres Element gibt es eine **Anordnungsrelation**, d.h. man kann von zwei reellen Zahlen sagen, ob die eine kleiner oder größer als die andere ist. Wir definieren für zwei Zahlen $a, b \in \mathbb{R}$ also eine Relation $a \leq b$ (a ist kleiner oder gleich b), die folgende Regeln erfüllt

- Für alle $a, b, c \in \mathbb{R}$ gilt

$$a \leq b \Rightarrow a + c \leq b + c. \quad (1.2.11)$$

- Für alle $a, b \in \mathbb{R}$ und alle $c > 0$ gilt

$$a \leq b \Rightarrow ac \leq bc. \quad (1.2.12)$$

- Seien $a, b \in \mathbb{R}$ mit $a > 0$ und $b > 0$. Dann existiert stets eine natürliche Zahl $n \in \mathbb{N} \subset \mathbb{R}$, so dass

$$b \leq na. \quad (1.2.13)$$

Die bis jetzt gegebenen Gesetze definieren einen **archimedisch angeordneten Zahlenkörper**. Allerdings erfüllt auch die Menge der rationalen Zahlen $\mathbb{Q}$, also die Teilmenge der reellen Zahlen, die sich als „Bruch“ zweier ganzer Zahlen schreiben lassen, diese Axiome. Um $\mathbb{R}$ eindeutig zu charakterisieren, benötigen wir noch die **Vollständigkeit** bzgl. der Grenzwertbildung von Folgen. Darauf gehen wir im folgenden Abschnitt genauer ein.

Für das Folgende benötigen wir auch noch den Begriff des **Betrags** einer reellen Zahl. Dieser ist definiert als

$$|a| = \begin{cases} a & \text{falls } a \geq 0, \\ (-a) & \text{falls } a < 0. \end{cases} \quad (1.2.14)$$

Offensichtlich ist $|a| \geq 0$ für alle $a \in \mathbb{R}$. Es gelten die Rechenregeln

$$|ab| = |a||b|, \quad |a + b| \leq |a| + |b|. \quad (1.2.15)$$

Die letztgenannte Ungleichung heißt **Dreiecksungleichung**.

1.3 Folgen und Grenzwerte

Eine Abbildung $\mathbb{N} \rightarrow \mathbb{R}$, $n \mapsto a_n$ bezeichnet man als **reelle Zahlenfolge**, d.h. jeder natürlichen Zahl n wird eindeutig eine reelle Zahl a_n zugeordnet. Wir bezeichnen eine solche Zahlenfolge abkürzend auch als (a_n) oder genauer $(a_n)_{n \in \mathbb{N}}$.

Eine Folge reeller Zahlen (a_n) konvergiert gegen eine reelle Zahl a genau dann, wenn es zu jedem reellen $\epsilon > 0$ eine natürliche Zahl N gibt, so dass $|a_n - a| < \epsilon$ für alle $n > N$ gilt. Anschaulich bedeutet das, dass man die Abweichung der Folgenglieder mit hinreichend großem Index n von der Zahl a beliebig klein machen kann. Man schreibt dann auch

$$\lim_{n \rightarrow \infty} a_n = a. \quad (1.3.1)$$

Eine Folge, für die eine solche Zahl a existiert, heißt **konvergent** und a der Grenzwert der Folge.

Wir können nun die konvergenten Folgen auch noch charakterisieren, wenn wir den Grenzwert nicht kennen. Zunächst nehmen wir an, die Folge (a_n) sei konvergent mit Grenzwert a . Sei weiter $\epsilon > 0$ eine beliebige positive reelle Zahl. Dann können wir definitionsgemäß eine Zahl N finden, so dass $|a_n - a| < \epsilon/2$ für alle $n > N$ ist. Seien dann $n_1, n_2 > N$ natürliche Zahlen. Dann folgt

$$|a_{n_1} - a_{n_2}| = |(a_{n_1} - a) + (a - a_{n_2})| \leq |a_{n_1} - a| + |a - a_{n_2}| < \epsilon/2 + \epsilon/2 = \epsilon. \quad (1.3.2)$$

Ist also a_n konvergent, kann man zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$ finden, so dass

$$|a_{n_1} - a_{n_2}| < \epsilon \quad \text{für alle } n_1, n_2 > N \quad (1.3.3)$$

gilt. Folgen mit dieser Eigenschaft heißen **Cauchy-Folgen**.

Anschaulich bedeutet das, dass die Folgenglieder (a_n) für hinreichend große N beliebig nahe beieinander liegen, wenn nur ihre Indizes $> N$ sind. Anschaulich liegt es daher nahe, zu denken, dass alle Cauchy-Folgen konvergent sind. Dies ist aber **genau für die reellen Zahlen $\mathbb{R}$** der Fall. Man sagt auch, die reellen Zahlen seien **abgeschlossen bzgl. der Grenzwertbildung**.

Man kann zeigen, dass $\mathbb{R}$ eindeutig als **abgeschlossener Archimedisch angeordneter Zahlkörper** charakterisiert ist.

Man beachte, dass der Körper der rationalen Zahlen $\mathbb{Q} \subset \mathbb{R}$, der alle algebraischen Eigenschaften wie die reellen Zahlen besitzt, *nicht abgeschlossen bzgl. der Grenzwertbildung* sind. Z.B. gibt es rationale Zahlenfolgen, die gegen $\sqrt{2}$ konvergieren, und wir wissen, dass $\sqrt{2}$ keine rationale Zahl ist.

Wir bemerken noch, dass die Grenzwertbildung mit der Addition und Multiplikation „verträglich“ sind, d.h. für konvergente Zahlenfolgen (a_n) und (b_n) mit den Grenzwerten a bzw. b gilt

$$\lim_{n \rightarrow \infty} a_n + b_n = a + b, \quad \lim_{n \rightarrow \infty} a_n b_n = ab. \quad (1.3.4)$$

Ist auch noch $b_n \neq 0$ für alle $n \in \mathbb{N}$ und auch der Grenzwert $b \neq 0$, gilt auch

$$\lim_{n \rightarrow \infty} \frac{a_n}{b_n} = \frac{a}{b}. \quad (1.3.5)$$

Wir wollen als Beispiel für die Anwendung dieser Begriffe einen wichtigen Grenzwertsatz beweisen. Dazu definieren wir zunächst, dass eine Folge **monoton wachsend** ist, wenn für alle $n_1 < n_2$ stets $a_{n_2} \geq a_{n_1}$ ist. Eine Folge heißt **nach oben beschränkt**, wenn es eine Zahl $M \in \mathbb{R}$ gibt, so dass $a_n \leq M$ für alle $n \in \mathbb{N}$ gilt.

Es gilt der folgende Grenzwertsatz: Eine nach oben beschränkte monoton wachsende Folge ist konvergent.

Zum *Beweis* bemerken wir, dass wegen der Monotonie der Folge alle Folgenglieder im Intervall $I_0 = [a_1, K]$ liegen. Dabei bedeuten die eckigen Klammern, dass die Endpunkte in dem Intervall enthalten sein sollen (**abgeschlossenes Intervall**). Jetzt zerlegen wir das Intervall in zwei Hälften $[a_1, (a_1 + K)/2]$ und $[(a_1 + K)/2, K]$. Falls in der rechten Hälfte noch Folgenglieder liegen, bezeichnen wir diese mit I_1 . Andernfalls ist I_1 die linke Hälfte. Da die Folge monoton wachsend ist und wenigstens ein Folgenglied a_{n_1} in I_1 liegt, liegen auch alle Folgenglieder a_n mit $n > n_1$ in diesem Intervall I_1 . Nach Konstruktion gibt es aber keine Folgenglieder, die größer sind als die rechte Grenze dieses Intervalls. So können wir nun beliebig oft fortfahren und so immer kleinere Intervalle I_k bilden, so dass es stets $n_k \in \mathbb{N}$ gibt, so dass $a_n \in I_k$ für alle $n > n_k$ gilt, aber keine Folgenglieder größer als die jeweils rechte Grenze der Intervalle I_k sind. Das Intervall I_k hat offenbar die Länge $L_k = (K - a_1)/2^k$. Da für alle $n > n_k$ die Folgenglieder $a_n \in I_k$ sind, gilt also für alle $n, n' > n_k$, dass $|a_n - a_{n'}| < L_k$ ist. Sei nun $\epsilon > 0$ beliebig vorgegeben. Dann können wir offenbar ein $k \in \mathbb{N}$ finden, so dass $L_k < \epsilon$ ist, denn es ist

$$L_k = \frac{K - a_1}{2^k} < \epsilon \Leftrightarrow 2^k > \frac{K - a_1}{\epsilon}, \quad (1.3.6)$$

und wir können zweifelsohne ein k finden, das diese Bedingung erfüllt. Folglich gibt es zu jedem $\epsilon > 0$ ein $k \in \mathbb{N}$, so dass $|a_n - a_{n'}| < \epsilon$ für alle $n, n' > n_k$. Damit ist aber die Folge nach dem Cauchyschen Konvergenzkriterium konvergent, und das war zu zeigen.

Wichtig ist noch der **Satz von Bolzano-Weierstraß**. Sei (a_n) eine beschränkte Folge, d.h. es gibt Zahlen $m < M \in \mathbb{R}$, so dass $m \leq a_n \leq M$ für alle $n \in \mathbb{N}$. Sei weiter (n_k) eine Folge mit $n_k \in \mathbb{N}$ und $n_k \rightarrow \infty$ für $k \rightarrow \infty$. Dann heißt die Folge $(a'_{n_k}) = (a'_{n_k})$ **Teilfolge** von (a_n) . Der Satz von Bolzano-Weierstraß besagt nun, dass jede beschränkte Folge stets wenigstens eine konvergente Teilfolge enthält.

Zum *Beweis* betrachten wir das Intervall $I_0 = [m, M]$, in dem voraussetzungsgemäß alle (unendlich vielen) Folgenglieder enthalten sind. Jetzt betrachten wir den Mittelpunkt $x_1 = (m + M)/2$. Liegen dann im linken Teilintervall $[m, x_1]$ wieder unendlich viele Folgenglieder, nennen wir dieses Teilintervall I_1 . Andernfalls müssen im rechten Teilintervall unendlich viele Folgenglieder liegen, und wir nennen dieses Teilintervall I_1 . So können wir beliebig oft verfahren. Wir haben dann eine Folge von Intervallen (I_k) , deren Länge $l_k = (M - m)/2^k$ ist. In jedem Intervall I_k liegen unendlich viele Folgenglieder, und wir können daher ein $n_k \in \mathbb{N}$ finden, so dass $a_{n_k} \in I_k$ liegt und dass für $k > k'$ stets $n_k > n_{k'}$ ist. Dann strebt sicher $n_k \rightarrow \infty$, wenn $k \rightarrow \infty$ geht. Es ist also $(a'_{n_k})_k = (a_{n_k})_k$ eine Teilfolge. Wir zeigen nun, dass diese konvergent ist, indem wir nachweisen, dass sie eine Cauchy-Folge ist. Sei dazu $\epsilon > 0$ beliebig vorgegeben. Dann sei k_0 so groß, dass $l_k < \epsilon$ für alle $k > k_0$ ist. Das ist sicher immer möglich, da $l_k = (M - m)/2^k \rightarrow 0$ für $k \rightarrow \infty$ ist. Nun ist für zwei beliebige natürliche Zahlen $k_1, k_2 > k_0$ die Folgenglieder $a'_{n_{k_1}}, a'_{n_{k_2}} \in I_{k_0}$, denn für $k > k_0$ ist aufgrund unserer Konstruktion der Intervalle $I_k \subset I_{k_0}$. Dann ist aber $|a'_{n_{k_1}} - a'_{n_{k_2}}| < \epsilon$. Für jedes $\epsilon > 0$ existiert also ein $k_0 \in \mathbb{N}$, so dass $|a'_{n_{k_1}} - a'_{n_{k_2}}| < \epsilon$ für alle $k_1, k_2 > k_0$, und folglich ist $(a'_{n_k})_k$ konvergent. Da dies eine Teilfolge von $(a_n)_n$ ist, ist damit der Satz von Bolzano-Weierstraß bewiesen.

Eine weitere wichtige Art von Folgen sind die **Reihen**. Sei dazu (a_n) eine beliebige Folge reeller Zahlen. Dann bezeichnet man als **Teilsummenfolge** die durch

$$s_n = a_1 + a_2 + \dots + a_n = \sum_{j=1}^n a_j \quad (1.3.7)$$

definierte Folge (s_n) . Falls diese Teilsummenfolge zu einem Grenzwert s konvergiert, spricht man von einer **konvergenten unendlichen Reihe** und schreibt

$$s = \lim_{n \rightarrow \infty} s_n = \sum_{j=1}^{\infty} a_j. \quad (1.3.8)$$

Eine Reihe heißt **absolut konvergent**, wenn

$$\sum_{j=1}^{\infty} |a_n| \quad (1.3.9)$$

existiert.

Es ist klar, dass eine absolut konvergente unendliche Reihe immer konvergent ist. Um das einzusehen, wenden wir das Cauchysche Konvergenzkriterium auf die Teilsummenfolge an. Demnach ist die Reihe (1.3.8) konvergent, wenn es zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$ gibt, so dass

$$|s_{n_1} - s_{n_2}| < \epsilon \quad \text{für alle } n_1, n_2 > N. \quad (1.3.10)$$

Wir können nun ohne Beschränkung der Allgemeinheit annehmen, dass $n_2 > n_1 > N$ ist. Dann bedeutet (1.3.10)

$$\left| \sum_{j=n_1+1}^{n_2} a_j \right| < \epsilon. \quad (1.3.11)$$

Da nun voraussetzungsgemäß die Reihe absolut konvergent ist, können wir sogar ein $N \in \mathbb{N}$ finden, so dass die Teilsummenfolge von (1.3.9) das Cauchy-Kriterium erfüllt, d.h.

$$\sum_{j=n_1+1}^{n_2} |a_j| < \epsilon \quad \text{für alle } n_2 > n_1 > N. \quad (1.3.12)$$

Wegen der Dreiecksungleichung ist dann aber

$$\left| \sum_{j=n_1+1}^{n_2} a_j \right| \leq \sum_{j=n_1+1}^{n_2} |a_j| < \epsilon \quad \text{für alle } n_2 > n_1 > N. \quad (1.3.13)$$

Das bedeutet aber, dass die Reihe (1.3.8) konvergent ist, wenn sie absolut konvergent ist. *Die Umkehrung gilt allerdings nicht!*

Für absolut konvergente Reihen gilt der **Umordnungssatz**: Sei $P : \mathbb{N} \rightarrow \mathbb{N}$ eine beliebige umkehrbar eindeutige Abbildung und sei die Reihe $\sum a_n$ absolut konvergent mit dem Grenzwert $\sum_{n=1}^{\infty} a_n = a$. Dann gilt auch

$$\sum_{k=1}^{\infty} a_{P(k)} = a. \quad (1.3.14)$$

Beweis: Da die Reihe voraussetzungsgemäß absolut konvergent ist, gibt es zu jedem $\epsilon > 0$ ein $n_0 \in \mathbb{N}$ mit

$$\sum_{k=n_0}^{\infty} |a_k| < \frac{\epsilon}{2}. \quad (1.3.15)$$

Dann ist

$$\left| a - \sum_{k=1}^{n_0-1} a_k \right| \leq \sum_{k=n_0}^{\infty} |a_k| < \frac{\epsilon}{2}. \quad (1.3.16)$$

Da $P : \mathbb{N} \rightarrow \mathbb{N}$ umkehrbar eindeutig ist, gibt es ein $N \in \mathbb{N}$, so dass $\{P(1), P(2), \dots, P(N)\} \supset \{1, 2, \dots, n_0 - 1\}$. Dann gilt für alle $m \geq N$

$$\left| \sum_{k=1}^m a_{P(k)} - a \right| \leq \left| \sum_{k=1}^m a_{P(k)} - \sum_{k=1}^{n_0-1} a_k \right| + \left| \sum_{k=1}^{n_0-1} a_k - a \right| \leq \sum_{k=n_0}^{\infty} |a_k| + \frac{\epsilon}{2} < \epsilon, \quad (1.3.17)$$

wobei wir (1.3.15) und (1.3.16) verwendet haben. Damit ist aber klar, dass tatsächlich (1.3.14) gilt, und das war zu zeigen.

Im Folgenden beweisen wir einige **Konvergenzkriterien** für Reihen, also hinreichende Bedingungen für die Konvergenz von Reihen.

Leibnizsches Konvergenzkriterium für alternierende Reihen: Es sei (a_k) eine monoton fallende Nullfolge mit $a_k \geq 0$. Dann ist die „alternierende Reihe“ $\sum_{n=1}^{\infty} (-1)^{n-1} a_n$ konvergent.

Beweis: Wir betrachten die Partialsummen

$$S_k = \sum_{n=1}^k (-1)^{n-1} a_n. \quad (1.3.18)$$

Dann gilt für die Partialsummenfolge mit geraden Indizes

$$S_{2k+2} - S_{2k} = (-1)^{2k} a_{2k+1} + (-1)^{2k+1} a_{2k+2} = a_{2k+1} - a_{2k+2} \geq 0, \quad (1.3.19)$$

da die Folge $(a_j)_{j \in \mathbb{N}}$ voraussetzungsgemäß monoton fallend ist. Folglich ist die Folge $(S_{2k})_{k \in \mathbb{N}}$ monoton wachsend, d.h. für alle k ist

$$S_2 \leq S_4 \leq \dots \leq S_{2k}. \quad (1.3.20)$$

Für die ungeraden Partialsummenfolglieder ist hingegen

$$S_{2k+3} - S_{2k+1} = (-1)^{2k+1} a_{2k+2} + (-1)^{2k+2} a_{2k+3} = -a_{2k+2} + a_{2k+3} \leq 0 \quad (1.3.21)$$

und damit diese Teilsummenfolge monoton fallend, d.h. für alle $k \in \mathbb{N}$.

$$S_1 \geq S_3 \geq S_5 \geq \dots \geq S_{2k+1}. \quad (1.3.22)$$

Andererseits ist

$$S_{2k+1} - S_{2k} = a_{2k+1} \geq 0 \quad (1.3.23)$$

und folglich für alle k

$$S_{2k} \leq S_{2k+1} \leq S_1. \quad (1.3.24)$$

Da (S_{2k}) somit gemäß (1.3.19) monoton wachsend und wegen (1.3.24) nach oben beschränkt ist, ist diese Folge konvergent. Es sei

$$\lim_{k \rightarrow \infty} S_{2k} = S_1. \quad (1.3.25)$$

Aus (1.3.23) folgt andererseits für alle $k \in \mathbb{N}$

$$S_{2k+1} \geq S_{2k} \geq S_2 \quad (1.3.26)$$

Damit ist also (S_{2k+1}) eine gemäß (1.3.22) monoton fallende und wegen (1.3.26) nach unten beschränkte Folge damit ebenfalls konvergent. Es sei

$$\lim_{k \rightarrow \infty} S_{2k+1} = S_2. \quad (1.3.27)$$

Offenbar gilt $S_1 = S_2$, denn es ist

$$S_{2k+1} - S_{2k} = a_{2k+1} \rightarrow 0 \quad \text{für } k \rightarrow \infty, \quad (1.3.28)$$

da die Folge (a_k) voraussetzungsgemäß eine Nullfolge ist. Demnach konvergieren (S_{2k}) und S_{2k+1} gegen denselben Grenzwert $S = S_1 = S_2$. Das bedeutet, dass es zu $\epsilon > 0$ Zahlen $N_1 \in \mathbb{N}$ und $N_2 \in \mathbb{N}$ gibt, so dass $|S_{2k} - S| < \epsilon$ für alle $2k \geq N_1$ und $|S_{2k+1} - S| < \epsilon$ für alle $2k + 1 \geq N_2$. Setzen wir dann $N = \max(N_1, N_2)$, so ist $|S_k - S| < \epsilon$ für alle $k > N$, und folglich ist (S_k) konvergent, und das war zu zeigen.

Damit können wir auch ein *Beispiel* für eine konvergente aber nicht absolut konvergente Reihe angeben. Betrachten wir die Reihe $\sum_{k=1}^{\infty} (-1)^{k+1}/k = 1 - 1/2 + 1/3 + \dots$, so ist diese nach dem Leibnizschen Konvergenzkriterium konvergent. Dies ist jedoch nicht der Fall für die Reihe $\sum_{k=1}^{\infty} 1/k = 1 + 1/2 + 1/3 + \dots$. Betrachten wir nämlich die Partialsummen dieser letztgenannten Reihe, folgt

$$S_{2k} = 1 + 1/2 + (1/3 + 1/4) + (1/5 + 1/6 + 1/7 + 1/8) + \dots + [1/(2^{k-1} + 1) + \dots + 1/2^k]. \quad (1.3.29)$$

Jede Klammer ist nun aber $\geq K/2^j$ mit $K = 2^j - 2^{j-1} = 2^{j-1}(2 - 1) = 2^{j-1}$, also $\geq 1/2$ und damit

$$S_{2k} \geq 1 + \frac{k}{2} \rightarrow \infty \quad \text{für } k \rightarrow \infty, \quad (1.3.30)$$

und folglich divergiert die Reihe $\sum_{k=1}^{\infty} 1/k \rightarrow \infty$.

Wir zeigen nun auch, dass wir die entsprechende alternierende Reihe so umordnen können, dass sie divergiert. Dazu schreiben wir ihre Glieder in der folgenden Reihenfolge

$$\begin{aligned} &1 - 1/2 + 1/3 - 1/4 \\ &\quad + (1/5 + 1/7) - 1/6 \\ &\quad + (1/9 + 1/11 + 1/13 + 1/15) - 1/8 \\ &\quad + \dots \\ &\quad + [1/(2^j + 1) + 1/(2^j + 3) + \dots + 1/(2^{j+1} - 1)] - 1/(2n + 2) \\ &\quad + \dots \end{aligned} \quad (1.3.31)$$

Für die Klammern gilt

$$1/(2^j + 1) + 1/(2^j + 3) + \dots + 1/(2^{j+1} - 1) \geq (2^{n+1} - 2^n - 1)/2^{n+1} = (2^n - 1)/2^{n+1} \geq 2^{n-1}/2^{n+1} = 1/4. \quad (1.3.32)$$

Damit ist aber jede der Zeilen in (1.3.31) größer als 0, und folglich die Umordnung gegen ∞ divergent.

Ein wichtiges **Kriterium für die absolute Konvergenz** einer Reihe ist das **Vergleichskriterium**. Dazu sei a_j eine beliebige Folge und b_j eine Folge mit $b_j \geq 0$, so dass $|a_j| < b_j$ für alle $j \in \mathbb{N}$ ist. Ist dann die aus den b_j gebildete Reihe konvergent, so ist die aus den a_j gebildete Reihe absolut konvergent.

Beweis: Die Teilsummenfolge der aus den b_j gebildeten Reihe

$$s_n^{(b)} = \sum_{j=1}^n b_j \quad (1.3.33)$$

ist monoton wachsend und voraussetzungsgemäß konvergent. Sei $s^{(b)}$ der Grenzwert. Offenbar gilt $s_n^{(b)} \leq s^{(b)}$ für alle $n \in \mathbb{N}$. Da aber $|a_j| < b_j$ für alle $j \in \mathbb{N}$ ist, gilt für die Teilsummenfolge

$$s_n^{(|a|)} = \sum_{j=1}^n |a_j| \leq \sum_{j=1}^n b_j \leq s^{(b)}. \quad (1.3.34)$$

Folglich ist die monoton wachsende Folge $s_n^{(|a|)}$ nach oben beschränkt und damit nach dem oben bewiesenen Konvergenzkriterium für monoton wachsende Folgen ebenfalls konvergent. Damit ist die aus a_j gebildete Reihe absolut konvergent, was zu beweisen war.

Ein wichtiges Beispiel ist die **geometrische Reihe**. Sei dazu $q \neq 0$. Die **geometrische Folge** ist dann durch $a_j = q^j$ für $j \in \mathbb{N}$ definiert, und die geometrische Reihe wird aus dieser Folge gebildet.

Hier liegt der seltene Fall vor, dass wir die Teilsummenfolge explizit ausrechnen können, denn es gilt

$$s_n = \sum_{j=1}^n q^j = q + q^2 + \cdots + q^n. \quad (1.3.35)$$

Multiplizieren wir diese Gleichung mit q , folgt

$$qs_n = \sum_{j=1}^n q^{j+1} = q^2 + q^3 + \cdots + q^{n+1}. \quad (1.3.36)$$

Subtrahieren wir beide Gleichungen, ergibt sich

$$s_n - qs_n = (1 - q)s_n = q - q^{n+1}. \quad (1.3.37)$$

Falls nun $q \neq 1$ ist, folgt daraus

$$s_n = \frac{q - q^{n+1}}{1 - q}. \quad (1.3.38)$$

Für $q = 1$ ist $s_n = n$. Offensichtlich ist diese Teilsummenfolge genau dann konvergent, wenn $|q| < 1$ ist, denn andernfalls wächst q^{n+1} für $n \rightarrow \infty$ über alle Grenzen.

Für $|q| < 1$ ist $\lim_{n \rightarrow \infty} q^n = 0$ (*Beweis als Übungsaufgabe*) und damit

$$\sum_{j=1}^{\infty} q^j = \frac{q}{1 - q} \quad \text{falls } |q| < 1. \quad (1.3.39)$$

Die geometrische Reihe ist nützlich, um zusammen mit dem oben hergeleiteten Vergleichskriterium für absolute Konvergenz von Reihen einfache Kriterien für die absolute Konvergenz von Reihen zu liefern, das **Quotienten-** und das **Wurzelkriterium**. Gemäß dem Vergleichskriterium ist nämlich die aus der Folge (a_j) gebildete Reihe absolut konvergent, wenn es reelle Zahlen $A > 0$ und $0 < q < 1$ und eine Zahl $N \in \mathbb{N}$ gibt, so dass

$$|a_j| < Aq^j \quad \text{für alle } j > N \quad (1.3.40)$$

gilt.

Angenommen, für die a_j gilt

$$\left| \frac{a_{j+1}}{a_j} \right| < q < 1 \quad \text{für alle } j > N. \quad (1.3.41)$$

Dann folgt

$$|a_{N+2}| < |a_{N+1}|q, \quad |a_{N+3}| < |a_{N+2}|q < |a_{N+1}|q^2, \quad \dots, \quad |a_{N+j}| < |a_{N+1}|q^{j-1} \quad \text{für alle } j \in \mathbb{N}. \quad (1.3.42)$$

Setzt man also $A = |a_{N+1}|q^{-N-1}$, so folgt daraus, dass

$$|a_k| < |a_{N+1}|q^{k-N-1} = Aq^k \quad \text{für alle } k > N. \quad (1.3.43)$$

Demnach ist gemäß dem Vergleichskriterium (1.3.40) die aus der Folge (a_j) gebildete Reihe absolut konvergent.

Offensichtlich ist (1.3.41) insbesondere dann erfüllt, wenn der Grenzwert

$$0 \leq \lim_{j \rightarrow \infty} \left| \frac{a_{j+1}}{a_j} \right| = q' < 1 \quad (1.3.44)$$

ist. Denn dann gibt es zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$, so dass

$$\left| \left| \frac{a_{j+1}}{a_j} \right| - q' \right| < \epsilon \quad \text{für alle } j > N. \quad (1.3.45)$$

Das bedeutet aber, dass für alle $j > N$

$$\left| \frac{a_{j+1}}{a_j} \right| < q' + \epsilon \quad (1.3.46)$$

gilt. Wählen wir nun ϵ so klein, dass $q = q' + \epsilon < 1$, also $0 < \epsilon < 1 - q'$ ist, ist folglich (1.3.41) mit diesem $0 < q < 1$ erfüllt, und die Reihe ist konvergent.

Dies ist das **Quotientenkriterium**: Falls (1.3.44) erfüllt ist, ist die aus (a_j) gebildete Reihe absolut konvergent.

Ebenso ergibt sich das **Wurzelkriterium**. Offensichtlich ist nämlich (1.3.40) genau dann erfüllt, wenn für alle $j > N$

$$|a_j|^{1/j} < A^{1/j} q \quad (1.3.47)$$

gilt. Ist dabei $q < 1$, ist die aus der Folge (a_j) gebildete Reihe absolut konvergent. Dies ist insbesondere der Fall, wenn

$$\lim_{j \rightarrow \infty} |a_j|^{1/j} = q' < 1 \quad (1.3.48)$$

gilt.

Besonders wichtig sind nun die **Potenzreihen**, die es gestatten, außer den rein algebraisch definierbaren **Polynomen** auch allgemeinere Funktionen zu definieren. Dazu betrachtet man Reihen der Form

$$f(x) = \sum_{j=0}^{\infty} a_j x^j. \quad (1.3.49)$$

Die Teilsummen sind Polynome in x . Es stellt sich freilich sofort die Frage, ob die Reihe für irgendein $x \neq 0$ konvergiert und ggf. für welche $x \in \mathbb{R}$ dies der Fall ist.

Das Quotientenkriterium liefert hier eine Antwort:

Die Potenzreihe ist sicher für diejenigen x absolut konvergent, für die

$$\lim_{j \rightarrow \infty} \left| \frac{a_{j+1}x^{j+1}}{a_jx^j} \right| < 1 \quad (1.3.50)$$

ist (vorausgesetzt der Grenzwert existiert). Das kann man aber sofort auch umschreiben:

$$|x| \lim_{j \rightarrow \infty} \left| \frac{a_{j+1}}{a_j} \right| < 1 \Rightarrow |x| < \lim_{j \rightarrow \infty} \left| \frac{a_j}{a_{j+1}} \right| = R. \quad (1.3.51)$$

Falls der rechts stehende Grenzwert R existiert, ist demnach die Potenzreihe (1.3.49) für alle $|x| < R$ absolut konvergent. Falls $R \rightarrow \infty$, ist die Potenzreihe sogar für alle $x \in \mathbb{R}$ absolut konvergent.

1.4 Satz vom Supremum und Infimum

In diesem Abschnitt beleuchten wir nochmals die Vollständigkeitseigenschaften der reellen Zahlen von einem etwas anderen Blickwinkel.

Wir betrachten dazu eine beliebige Teilmenge der reellen Zahlen: $A \subseteq \mathbb{R}$. Diese Menge heißt von oben (unten) beschränkt, wenn es eine Zahl M (m) gibt, so dass für alle $x \in A$ stets $x < M$ ($x > m$) gilt. Dabei müssen die **obere Schranke** M (**untere Schranke** m) selbst nicht notwendig zu A gehören. Man nennt nun die kleinste obere (größte untere) Schranke von A das **Supremum** (**Infimum**) der Menge A (in Formeln $\sup A$ bzw. $\inf A$).

Wir zeigen nun den wichtigen Satz:

Jede nach oben beschränkte Menge $A \subset \mathbb{R}$ besitzt ein Supremum.

Analog besitzt eine nach unten beschränkte Menge A ein Infimum. Der Beweis der letzteren Behauptung verläuft ebenfalls vollständig analog zum Satz vom Supremum.

Beweis: Sei also $A \subseteq \mathbb{R}$ nicht leer und besitze eine obere Schranke M . Sei nun $x_1 \in A$ beliebig gewählt. Ist dann x_1 eine obere Schranke von A , ist diese Zahl offenbar bereits das Supremum dieser Menge, denn es kann dann keine kleinere obere Schranke als x_1 geben. Ist x_1 keine obere Schranke von A , betrachten wir das Intervall $I_1 = [x_1, M]$. Da voraussetzungsgemäß für alle $x \in A$ immer $x \leq M$ gilt und x_1 keine obere Schranke ist, existiert in I_1 noch wenigstens eine Zahl $x \in M$ mit $x_1 < x \leq M$. Betrachten wir nun den Mittelpunkt des Intervalls I_1 , also $x_2 = (x_1 + M)/2$. Falls x_2 obere Schranke von A ist, setzen wir $I_2 = [x_1, x_2]$, andernfalls $I_2 = [x_2, M]$. Dann ist in jedem Fall $A \cap I_2 \neq \emptyset$ und die obere Grenze des Intervalls I_2 ist eine obere Schranke von A . Dieses Verfahren können wir nun beliebig oft wiederholen, und es entsteht eine Folge von Intervallen $I_n = [a_n, b_n]$ mit $0 \leq b_n - a_n = (M - x_1)/2^{n-1}$. Für $n \rightarrow \infty$ wird das Intervall beliebig klein, und a_n und b_n konvergieren demnach zu einem Wert M' . Offenbar ist M' eine obere Schranke von A , denn es gilt für alle $x \in A$ stets $x \leq b_n$ für alle $n \in \mathbb{N}$. Demnach ist aber auch $x \leq \lim_{n \rightarrow \infty} b_n = M'$ und also M' eine obere Schranke. Wir müssen nun noch zeigen, dass M' die kleinste obere Schranke, also $\sup A$ ist. Betrachten wir dazu eine beliebige Zahl $M'' < M'$ und nehmen an, sie sei obere Schranke von A . Es ist klar, dass dann für alle

$n \in \mathbb{N}$ stets $M'' \in I_n$ gelten muss. Da nun aber die linken Grenzen a_n der Intervalle I_n monoton wachsen und gegen M streben, muss es ein $n_0 \in \mathbb{N}$ geben, so dass $a_n > M''$ für alle $n \geq n_0$ gilt. Dann ist aber für diese n immer $M'' \notin I_n$ im Widerspruch zur Annahme, dass M'' obere Schranke von A ist. Folglich muss $M' = \sup A$, also kleinste obere Schranke von A sein.

Wir betonen nochmals, dass wir hier wesentlich die **Vollständigkeit** der reellen Zahlen unter Grenzwertbildungen verwendet haben, denn sie garantiert, dass die eben konstruierte Zahl $M' \in \mathbb{R}$ gilt.

1.5 Lineare und quadratische Gleichungen

Zu den klassischen Aufgaben der Algebra gehört die Auflösung von Gleichungen.

Die einfachste Art von Gleichungen sind die **linearen Gleichungen**. Seien dazu $a, b \in \mathbb{R}$ irgendwelche Zahlen. Dann fragen wir, ob es eine Zahl $x \in \mathbb{R}$ gibt, die die Gleichung

$$ax + b = 0 \tag{1.5.1}$$

erfüllt.

Zunächst können wir die Gleichung vereinfachen, indem wir auf beiden Seiten b subtrahieren:

$$(ax + b) - b = -b \Rightarrow ax + (b - b) = ax = -b. \tag{1.5.2}$$

Diese Gleichung lässt sich nun offenbar genau dann eindeutig nach x auflösen, wenn es ein Inverses von a bzgl. der Multiplikation gibt, d.h. falls $a \neq 0$ ist. Dann folgt

$$x = -\frac{b}{a}. \tag{1.5.3}$$

Falls nun $a = 0$ ist, so ist $ax = 0x = 0$ für alle $x \in \mathbb{R}$. In diesem Fall wird die Gleichung durch alle Zahlen $x \in \mathbb{R}$ gelöst, falls $b = 0$ ist. Falls $b \neq 0$ ist, besitzt die Gleichung keine Lösung.

Kommen wir nun auf die **quadratischen Gleichungen** zu sprechen. Dies sind Gleichungen der Form

$$x^2 + px + q = 0. \tag{1.5.4}$$

Um sie zu lösen, erinnern wir an die **binomische Formel**

$$(a + b)^2 = a^2 + 2ab + b^2, \tag{1.5.5}$$

die man sehr leicht durch Anwendung des Distributivgesetzes nachrechnet (*Übung!*). Setzen wir hierin $a = x$ und $b = p/2$, finden wir

$$\left(x + \frac{p}{2}\right)^2 = x^2 + px + \frac{p^2}{4}. \tag{1.5.6}$$

Verwenden wir dies in (1.5.4), ergibt sich die Gleichung

$$\left(x + \frac{p}{2}\right)^2 + q - \frac{p^2}{4} = 0. \tag{1.5.7}$$

Dies ergibt

$$\left(x + \frac{p}{2}\right)^2 = \frac{p^2}{4} - q. \tag{1.5.8}$$

Wir benötigen nun lediglich eine **Umkehrfunktion** des Quadrierens, die **Quadratwurzel**. Ohne Beweis bemerken wir, dass für $a \geq 0$ die **Quadratwurzelfunktion** wohldefiniert ist, d.h. zu jedem reellen $a \geq 0$ existiert genau eine reelle Zahl $b \geq 0$ mit $b = \sqrt{a}$, die dadurch definiert ist, dass $b^2 = a$ ist. Freilich gilt dann auch $(-b)^2 = b^2 = a$. Für $a \neq 0$ gibt es also zwei Lösungen der Gleichung $b^2 = a$, nämlich $b = +\sqrt{a} \geq 0$ und $b = -\sqrt{a} \leq 0$.

Dies wenden wir nun auf (1.5.8) an. Falls also $p^2/4 - q > 0$ ist, existieren stets zwei Lösungen

$$x_1 + \frac{p}{2} = \sqrt{\frac{p^2}{4} - q} \Rightarrow x_1 = -\frac{p}{2} + \sqrt{\frac{p^2}{4} - q} \quad (1.5.9)$$

und

$$x_2 + \frac{p}{2} = -\sqrt{\frac{p^2}{4} - q} \Rightarrow x_2 = -\frac{p}{2} - \sqrt{\frac{p^2}{4} - q}. \quad (1.5.10)$$

Für $p^2/4 = q$ gibt es nur eine Lösung $x_1 = x_2 = -p/2$. Wir bemerken weiter, dass stets

$$x^2 + px + q = (x - x_1)(x - x_2) \quad (1.5.11)$$

gilt, falls die Gleichung (1.5.4) lösbar ist.

Dies zeigen wir durch Ausmultiplizieren:

$$(x - x_1)(x - x_2) = x^2 - (x_1 + x_2)x + x_1x_2. \quad (1.5.12)$$

Setzt man die Lösungen (1.5.9) und (1.5.10) ein, finden wir durch einfaches Ausrechnen in der Tat

$$-(x_1 + x_2) = p, \quad x_1x_2 = q. \quad (1.5.13)$$

Wie wir gesehen haben, gibt es im Reellen keine Lösung der quadratischen Gleichung (1.5.4), falls $q^2/4 - p < 0$ ist, denn für alle $a \in \mathbb{R}$ ist $a^2 \geq 0$ und $a^2 = 0 \Leftrightarrow a = 0$. Wir werden in Kapitel 4 sehen, wie man die reellen Zahlen erweitern kann, so dass auch in diesem Fall die quadratische Gleichung lösbar ist. Dies führt dann auf den **Körper der komplexen Zahlen**.

1.6 Funktionen und Stetigkeit

Eine **Funktion einer reellen Veränderlichen** ist eine eindeutige Abbildung von einer Teilmenge der reellen Zahlen ($D \subseteq \mathbb{R}$) in die reellen Zahlen. Dies schreiben wir in der Form $f : D \rightarrow \mathbb{R}$. Jeder Zahl $x \in D$ wird eine Zahl $y \in \mathbb{R}$ zugeordnet, nämlich $y = f(x)$. Veranschaulichen können wir uns Funktionen, indem wir (x, y) als Koordinaten eines kartesischen Koordinatensystems auffassen. Die Kurve, die aus den entsprechenden Punkten $(x, f(x))$ besteht, heißt **Graph der Funktion f** .

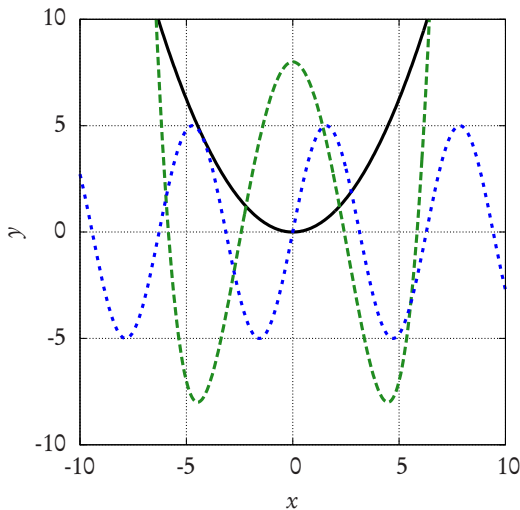


Abbildung 1.1: Beispiele für Graphen von Funktionen: $y = x^2/4$ (schwarz durchgezogen), $y = (0.2x^2 - 1)^2 - 8$ (grün gestrichelt) und $y = 5 \sin x$ (blau gepunktet).

ein $\delta_1 > 0$, so dass $|f(x) - f(x_0)| < \epsilon/2$ für alle $x \in D$ mit $|x - x_0| < \delta_1$. Ebenso gibt es auch ein $\delta_2 > 0$, so dass $|g(x) - g(x_0)| < \epsilon/2$ für alle $x \in D$ mit $|x - x_0| < \delta_2$. Setzen wir nun $\delta = \min(\delta_1, \delta_2)$, dann gilt wegen der Dreiecksungleichung für alle $x < \delta$

$$\begin{aligned} |[f(x) + g(x)] - [f(x_0) + g(x_0)]| &= |[f(x) - f(x_0)] + [g(x) - g(x_0)]| \\ &\leq |f(x) - f(x_0)| + |g(x) - g(x_0)| \\ &< \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon, \end{aligned} \quad (1.6.1)$$

d.h. zu jedem $\epsilon > 0$ können wir ein $\delta > 0$ finden, so dass für alle $x \in D$ mit $|x - x_0| < \delta$ stets

$$|[f(x) + g(x)] - [f(x_0) + g(x_0)]| < \epsilon \quad (1.6.2)$$

gilt. Das bedeutet aber gemäß der Definition, dass die Funktion $f + g$ stetig ist.

Ebenso beweist man, dass mit zwei Funktionen f und g , die stetig in x_0 sind, auch das Produkt fg und (vorausgesetzt $g(x_0) \neq 0$) der Quotient f/g stetig in x_0 sind (*Beweis als Übung!*).

Offensichtlich sind eine konstante Funktion $f(x) = A = \text{const}$ sowie die Funktion $f(x) = x$ in allen $x \in \mathbb{R}$ definiert und stetig. Aus den oben angegebenen Sätzen folgt dann, dass auch jedes Polynom

$$f(x) = a_0 + a_1x + a_2x^2 + \dots + a_nx^n = \sum_{k=0}^n a_k x^k \quad (1.6.3)$$

stetig ist.

In den Beispielen im nebenstehenden Bild sind alle Funktionen auf ganz $\mathbb{R}$ definiert. Außerdem weist ihr Graph keinerlei Unterbrechungen oder ähnliche Anomalien auf. Mathematisch ist eine Voraussetzung dafür, dass eine Funktion stetig ist. Anschaulich bedeutet die **Stetigkeit einer Funktion** an der Stelle x_0 , dass sich der Funktionswert um x_0 nur wenig ändert, wenn man sie in einer kleinen Umgebung um x_0 betrachtet. Formal lässt sich dies so definieren:

Eine Funktion $f : D \rightarrow \mathbb{R}$ heißt stetig an der Stelle $x_0 \in D$, wenn es zu jedem $\epsilon > 0$ ein $\delta > 0$ gibt, so dass für alle $x \in D$ mit $|x - x_0| < \delta$ stets $|f(x) - f(x_0)| < \epsilon$ ist.

Anders formuliert heißt das, dass die Funktionswerte für hinreichend kleine Umgebungen um x_0 beliebig wenig von dem Funktionswert $f(x_0)$ bei x_0 abweichen.

Es ist nun leicht zu zeigen, dass für zwei in x_0 stetige Funktionen f und g auch die Funktion $f + g$ stetig ist. Das sieht man sofort wie folgt ein: Voraussetzungsgemäß gibt es zu jedem $\epsilon > 0$

$$< \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon,$$

Wichtig ist der **Zwischenwertsatz**: Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig. Weiter sei $f(a) < f(b)$. Dann gibt es zu jedem $u \in [f(a), f(b)]$ ein $\xi \in [a, b]$, so dass $f(\xi) = u$ ist (entsprechendes gilt auch falls $f(a) > f(b)$ ist).

Zum Beweis konstruieren wir ξ mittels Intervallschachtelung. Im ersten Schritt betrachten wir also den Mittelpunkt des Ausgangsintervalls $\xi_1 = (a + b)/2$. Falls dann $f(\xi_1) = u$, haben wir ξ schon gefunden. Andernfalls, setzen wir $a_1 = \xi_1$, $b_1 = b$ falls $f(\xi_1) < u$ bzw. $a_1 = a$ und $b_1 = \xi_1$ falls $f(\xi_1) > u$. Wir haben dann wieder ein Intervall $[a_1, b_1]$, für das $f(a_1) \leq u \leq f(b_1)$ ist. Mit diesem Intervall verfahren wir nun genauso. Dadurch entsteht eine Folge von Intervallen $[a_n, b_n]$, für die stets $f(a_n) \leq u \leq f(b_n)$ ist. Die Folge a_n ist monoton wachsend und die Folge b_n monoton fallend. Die Länge des Intervalls ist $|b_n - a_n| = (b - a)/2^n$, und folglich konvergieren beide Folgen gegen denselben Grenzwert ξ . Da weiter f voraussetzungsgemäß im ganzen Intervall $[a, b]$ stetig ist, muss folglich $f(\xi) = u$ sein, und das war zu zeigen.

Weiter gilt der **Satz vom Maximum und Minimum**: Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig. Dann nimmt f an wenigstens einer Stelle $\xi \in [a, b]$ ein Maximum (Minimum) an. Dabei ist das Maximum durch $\sup_{x \in [a, b]} f(x)$ (das Minimum durch $\inf_{x \in [a, b]} f(x)$) definiert.

Wir beweisen den Satz für das Maximum. Der Beweis für das Minimum folgt analog. Da die Funktion im gesamten (abgeschlossenen!) Intervall definiert ist, ist ihr Bildbereich $f([a, b])$ eine beschränkte Menge und besitzt demnach gemäß Abschnitt 1.4 ein Supremum, d.h. es gilt $f(x) \leq M$ für alle $x \in [a, b]$ mit $M = \sup_{x \in [a, b]} f(x)$. Wie wir beim Beweis des Satzes vom Infimum und Supremum in Abschnitt 1.4 gesehen haben, existiert dann eine Zahlenfolge $x_n \in [a, b]$, so dass $\lim_{n \rightarrow \infty} f(x_n) = M$.

1.7 Differentialrechnung für Funktionen einer reellen Veränderlichen

Die Differentialrechnung beschäftigt sich mit **lokalen Eigenschaften** von Funktionen, widmet sich also der Untersuchung des Änderungsverhaltens von Funktionen in kleinen Umgebungen eines beliebigen Punktes im Definitionsbereich der Funktion.

1.7.1 Definition der Ableitung einer Funktion

Es sei $f : D \rightarrow \mathbb{R}$ mit einem offenen Definitionsbereich $D \subseteq \mathbb{R}$, d.h. wir nehmen an, dass mit $x_0 \in D$ auch eine bestimmte Umgebung um diesen Punkt in D enthalten sei, d.h. dass es ein $\delta > 0$ gibt, so dass für alle x mit $|x - x_0| < \delta$ auch $x \in D$ gilt.

In diesem Fall besitzt die Funktion f den **Grenzwert** $A \in \mathbb{R}$ für $x \rightarrow x_0$, wenn es zu jedem $\epsilon > 0$ ein $\delta > 0$ gibt, so dass

$$|f(x) - A| < \epsilon \quad \text{für alle } x \text{ mit } |x - x_0| < \delta. \quad (1.7.1)$$

In dem Fall schreiben wir

$$\lim_{x \rightarrow x_0} f(x) = A. \quad (1.7.2)$$

Wir bemerken, dass unter den obigen Voraussetzungen eine Funktion genau dann stetig in x_0 ist, wenn

$$\lim_{x \rightarrow x_0} f(x) = f(x_0) \quad (1.7.3)$$

gilt.

Die **Ableitung** der Funktion f an der Stelle x_0 ist dann definiert als

$$f'(x_0) = \frac{d}{dx} f(x_0) = \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0}, \quad (1.7.4)$$

vorausgesetzt, dieser Grenzwert existiert. In dem Fall heißt die Funktion f **differenzierbar** an der Stelle x_0 . Sie heißt differenzierbar in einem Bereich $D' \subseteq D$ falls $f'(x)$ in jedem Punkt $x \in D'$ differenzierbar ist.

Wichtige Beispiele sind

- $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = c = \text{const.}$ Dann gilt für jedes $x_0 \in \mathbb{R}$

$$\frac{f(x) - f(x_0)}{x - x_0} = 0 \Rightarrow f'(x_0) = 0. \quad (1.7.5)$$

Die konstante Funktion ist also in ganz $\mathbb{R}$ differenzierbar, und ihre Ableitung verschwindet.

- $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x$. Dann folgt

$$f'(x_0) = \lim_{x \rightarrow x_0} \frac{x - x_0}{x - x_0} = \lim_{x \rightarrow x_0} 1 = 1. \quad (1.7.6)$$

- $f : \mathbb{R} \rightarrow \mathbb{R}$ mit $f(x) = x^2$

$$f'(x_0) = \lim_{x \rightarrow x_0} \frac{x^2 - x_0^2}{x - x_0} = \lim_{x \rightarrow x_0} \frac{(x - x_0)(x + x_0)}{x - x_0} = \lim_{x \rightarrow x_0} (x + x_0) = 2x_0. \quad (1.7.7)$$

Die Ableitung besitzt auch eine **geometrische Bedeutung**. Betrachten wir dazu den Graphen einer Funktion $y = f(x)$. Dann bedeutet $[f(x) - f(x_0)] / (x - x_0)$ die Steigung der Sekante durch die Punkte $[x, f(x)]$ und $[x_0, f(x_0)]$. Die Ableitung ergibt, falls sie existiert, entsprechend die Steigung der **Tangente** an die durch den Funktionsgraphen gegebenen Kurve am Punkt $[x_0, f(x_0)]$.

1.7.2 Formeln zur Ableitung

Im Folgenden leiten wir einige allgemeine Formeln zur Ableitung von Funktionen her. Dabei nehmen wir an, dass alle beteiligten Funktionen im Punkt x_0 ihres Definitionsbereichs differenzierbar sind. Dann gilt

$$\frac{d}{dx} (f + g)(x_0) = f'(x_0) + g'(x_0). \quad (1.7.8)$$

Zum Beweis müssen wir nur den entsprechenden Limes betrachten

$$\begin{aligned} \frac{d}{dx} (f + g)(x_0) &= \lim_{x \rightarrow x_0} \frac{f(x) + g(x) - [f(x_0) + g(x_0)]}{x - x_0} \\ &= \lim_{x \rightarrow x_0} \left[\frac{f(x) - f(x_0)}{x - x_0} + \frac{g(x) - g(x_0)}{x - x_0} \right] \\ &= f'(x_0) + g'(x_0). \end{aligned} \quad (1.7.9)$$

Mit zwei Funktionen besitzt also auch deren Summe an der Stelle x_0 eine Ableitung, und diese berechnet sich als die Summe der Ableitungen. Ebenso einfach zeigt man, dass für eine beliebige Konstante $a \in \mathbb{R}$

$$\frac{d}{dx}(af)(x_0) = af'(x_0) \quad (1.7.10)$$

gilt. Zusammen mit (1.7.8) bedeutet dies, dass die Ableitung eine **lineare Operation** ist, d.h. für beliebigen Funktionen f, g , die bei x_0 differenzierbar sind und $a, b \in \mathbb{R}$ gilt

$$\frac{d}{dx}[af(x_0) + bf(x_0)] = a \frac{d}{dx}f(x_0) + b \frac{d}{dx}f(x_0) = af'(x_0) + bf'(x_0). \quad (1.7.11)$$

Weiter gilt die **Produktregel**

$$\frac{d}{dx}(fg)(x_0) = f'(x_0)g(x_0) + f(x_0)g'(x_0). \quad (1.7.12)$$

Zum Beweis verwenden wir wieder die Definition der Ableitung:

$$\begin{aligned} & \frac{f(x_0 + \Delta x)g(x_0 + \Delta x) - f(x_0)g(x_0)}{\Delta x} \\ &= \frac{[f(x_0 + \Delta x) - f(x_0)]g(x_0 + \Delta x) + f(x_0)[g(x_0 + \Delta x) - g(x_0)]}{\Delta x}. \end{aligned} \quad (1.7.13)$$

Nun ist eine Funktion, die in x_0 differenzierbar ist, dort auch stetig, und damit ergibt sich aus (1.7.13) im Limes $\Delta x \rightarrow 0$ die Produktregel (1.7.12).

Um zu zeigen, dass eine in x_0 differenzierbare Funktion f auch stetig ist, gehen wir auf die Definition des Limes zurück. Da f in x_0 differenzierbar ist, gibt es demnach zu jedem ϵ_0 ein $\delta > 0$, so dass

$$\left| \frac{f(x_0 + \Delta x) - f(x_0)}{\Delta x} - f'(x_0) \right| < \epsilon \quad \text{falls} \quad |\Delta x| < \delta. \quad (1.7.14)$$

Wir können nun diese Ungleichung mit $|\Delta x|$ multiplizieren:

$$|f(x_0 + \Delta x) - f(x_0) - \Delta x f'(x_0)| < \epsilon |\Delta x|. \quad (1.7.15)$$

Für $\Delta x \rightarrow 0$ strebt die rechte Seite der Ungleichung sowie $\Delta x f'(x_0)$ unter dem Betrag auf der linken Seite gegen 0. Damit ist aber

$$\lim_{\Delta x \rightarrow 0} |f(x_0 + \Delta x) - f(x_0)| = 0 \Leftrightarrow \lim_{\Delta x \rightarrow 0} f(x_0 + \Delta x) = f(x_0), \quad (1.7.16)$$

d.h. f ist an der Stelle x_0 stetig, und das war zu zeigen.

Weiter können wir nun auch die **Kettenregel** beweisen. Seien f und g Funktionen, wobei g bei x_0 und f bei $g(x_0)$ differenzierbar seien. Dann gilt

$$\frac{d}{dx}f[g(x_0)] = f'[g(x_0)]g'(x_0). \quad (1.7.17)$$

Zum Beweis verwenden wir wieder die Definition der Ableitung

$$\frac{f[g(x_0 + \Delta x)] - f[g(x_0)]}{\Delta x} = \frac{f[g(x_0 + \Delta x)] - f[g(x_0)]}{g(x_0 + \Delta x) - g(x_0)} \frac{g(x_0 + \Delta x) - g(x_0)}{\Delta x}. \quad (1.7.18)$$

1.7. Differentialrechnung für Funktionen einer reellen Veränderlichen

Da nach der gerade durchgeführten Überlegung g an der Stelle x_0 stetig ist, folgt daraus für $\Delta x \rightarrow 0$ die Kettenregel (1.7.17).

Mit der Produktregel können wir nun leicht mittels **vollständiger Induktion** die Formel

$$\frac{d}{dx} x^n = nx^{n-1}, \quad n \in \mathbb{N}_0 \quad (1.7.19)$$

beweisen, wobei $x \in \mathbb{R}$ beliebig sein darf, d.h. Potenzfunktionen (und damit auch Polynome) sind überall differenzierbar. Zum Beweis bemerken wir, dass die Formel sicher für $n = 0$ gilt. Denn dann ist $f(x) = x^0 = 1 = \text{const}$, und die Ableitung verschwindet gemäß (1.7.5). Nehmen wir nun an, die Formel gilt für $n = k$. Dann folgt mit der Produktregel

$$\frac{d}{dx} x^{k+1} = \frac{d}{dx} (x x^k) = \frac{dx}{dx} x^k + x \frac{d}{dx} x^k = x^k + x k x^{k-1} = (k+1)x^k, \quad (1.7.20)$$

und das ist genau (1.7.19) für $n = k+1$, und damit ist die Formel nach dem Beweisprinzip der vollständigen Induktion bewiesen.

Weiter ist die Funktion $f(x) = 1/x$ für $x \neq 0$ ebenfalls differenzierbar, denn es gilt

$$\frac{1}{\Delta x} \left(\frac{1}{x + \Delta x} - \frac{1}{x} \right) = \frac{1}{\Delta x} \left(\frac{x - (x + \Delta x)}{x(x + \Delta x)} \right) = -\frac{1}{x(x + \Delta x)}. \quad (1.7.21)$$

Mit $\Delta x \rightarrow 0$ erhalten wir damit für alle $x \neq 0$

$$\frac{d}{dx} \left(\frac{1}{x} \right) = -\frac{1}{x^2}. \quad (1.7.22)$$

Wenden wir weiter die Produkt- und die Kettenregel auf die Funktion $f/g = f \cdot 1/g$ an, wobei f und g in x_0 differenzierbar und $g(x_0) \neq 0$ sei, erhalten wir die **Quotientenregel**

$$\frac{d}{dx} \left(\frac{f}{g} \right) (x_0) = \frac{d}{dx} \left(f \frac{1}{g} \right) (x_0) = \frac{f'(x_0)}{g(x_0)} + f(x_0) \frac{d}{dx} \left(\frac{1}{g} \right) (x_0) = \frac{f'(x_0)}{g(x_0)} - f(x_0) \frac{g'(x_0)}{g^2(x_0)}. \quad (1.7.23)$$

Nochmals zusammengefasst ergibt sich also

$$\frac{d}{dx} \left(\frac{f}{g} \right) (x_0) = \frac{f'(x_0)g(x_0) - f(x_0)g'(x_0)}{g^2(x_0)}. \quad (1.7.24)$$

Damit können wir nun auch für $n \in \mathbb{N}$ die Funktionen $1/x^n$ für alle $x \neq 0$ ableiten:

$$\frac{d}{dx} \frac{1}{x^n} = \frac{0 - nx^{n-1}}{x^{2n}} = -nx^{-n-1}. \quad (1.7.25)$$

Es gilt also die Ableitungsregel

$$\frac{d}{dx} x^k = kx^{k-1} \quad \text{für alle } k \in \mathbb{Z}. \quad (1.7.26)$$

Schließlich betrachten wir noch eine Funktion $f: I \rightarrow \mathbb{R}$, die in einem Intervall I strikt monoton ist, d.h. es gilt für alle $x_1 < x_2 \in I$ stets $f(x_1) < f(x_2)$ (**strikt monoton wachsende**) oder $f(x_1) > f(x_2)$ (**strikt monoton**

fallende) Funktion. Es ist dann klar, dass in diesem Intervall die Zuordnung des Funktionswertes $y = f(x)$ zum Bildpunkt x umkehrbar eindeutig ist, d.h. zu jedem Wert y im Wertebereich der Funktion existiert genau ein $x \in I$ mit $f(x) = y$, d.h. diese Gleichung ist eindeutig nach x auflösbar.

Dies ermöglicht die Definition der **Umkehrfunktion** der Funktion $f, f^{-1} : f(I) \rightarrow I$. Dabei ist $f(I) = \{y \mid \text{es gibt ein } x \in I \text{ mit } f(x) = y\}$ die **Wertemenge** der Funktion f .

Wir nehmen nun an, die Funktion f sei differenzierbar im Intervall I . Aufgrund der Definition der Umkehrfunktion gilt für alle $x \in I$

$$f^{-1}[f(x)] = x. \quad (1.7.27)$$

Da die rechte Seite differenzierbar und auch f differenzierbar ist, ist demnach auch f^{-1} differenzierbar und nach der Kettenregel gilt

$$(f^{-1})'[f(x)]f'(x) = 1 \Rightarrow (f^{-1})'[f(x)] = \frac{1}{f'(x)}. \quad (1.7.28)$$

Mit Hilfe dieser Formel können wir die Ableitung der Umkehrfunktionen monotoner Funktionen finden, wenn wir die Ableitung der Funktion selbst kennen. Das erkennt man einfacher, indem man $y = f(x)$ und $x = f^{-1}(y)$ schreibt.

Dann ist gemäß (1.7.28)

$$f^{-1}'(y) = \frac{1}{f'(x)} = \frac{1}{f'[f^{-1}(y)]}. \quad (1.7.29)$$

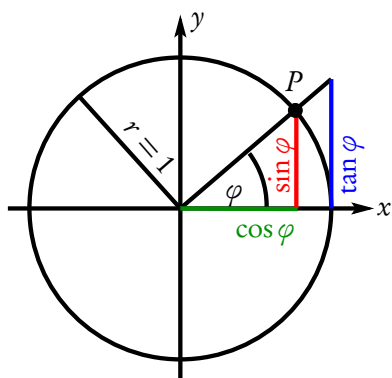
Betrachten wir als Beispiel $f(x) = x^2$. Setzen wir dann $y = f(x) = x^2$, folgt $x = \sqrt{y}$ (für $y > 0$). Mit (1.7.29) erhalten wir daraus

$$\frac{d}{dy} \sqrt{y} = \frac{1}{f'(x)} = \frac{1}{2x} = \frac{1}{2\sqrt{y}}. \quad (1.7.30)$$

Nennen wir jetzt y wieder x folgt die Ableitungsformel für die Wurzelfunktion

$$\frac{d}{dx} \sqrt{x} = \sqrt{x}' = \frac{1}{2\sqrt{x}}. \quad (1.7.31)$$

1.7.3 Trigonometrische Funktionen



Die trigonometrischen Funktionen werden gewöhnlich **geometrisch am Einheitskreis** entsprechend der nebenstehenden Skizze definiert. Durchläuft P den Einheitskreis, dann sind dessen Koordinaten in dem rechtwinkligen Koordinatensystem definitionsgemäß durch $(\cos \varphi, \sin \varphi)$ gegeben. Dabei ist der Winkel φ als die Länge des Kreisbogens a definiert. Für den Vollkreis ergibt sich 2π , für den gestreckten Winkel π und für den rechten Winkel $\pi/2$. Diese Definition des Winkels bezeichnet man als **Bogenmaß**. Der Zusammenhang zum **Gradmaß** ist durch die Beziehung $360^\circ = 2\pi$, d.h. $180^\circ = \pi$ gegeben. Man liest aus der Ähnlichkeit der beiden eingezeichneten rechtwinkligen Dreiecke die Beziehung

$$\tan \varphi = \frac{\sin \varphi}{\cos \varphi} \quad (1.7.32)$$

ab.

Aus der Definition wird unmittelbar klar, dass $\sin$ und $\cos$ **periodische Funktionen** mit der Periode 2π sind, d.h. es gilt für alle $\varphi \in \mathbb{R}$

$$\sin(\varphi + 2\pi) = \sin \varphi, \quad \cos(\varphi + 2\pi) = \cos \varphi. \quad (1.7.33)$$

Außerdem macht man sich schnell klar, dass

$$\sin(\varphi + \pi) = -\sin \varphi, \quad \cos(\varphi + \pi) = -\cos \varphi \quad (1.7.34)$$

ist. Daraus folgt, dass der $\tan$ eine **periodische Funktion** der Periode π ist,

$$\tan(\varphi + \pi) = \frac{\sin(\varphi + \pi)}{\cos(\varphi + \pi)} = \frac{-\sin \varphi}{-\cos \varphi} = \frac{\sin \varphi}{\cos \varphi} = \tan \varphi. \quad (1.7.35)$$

Aus der Definition am Einheitskreis und dem Satz des Pythagoras folgt, dass für alle $\varphi \in \mathbb{R}$

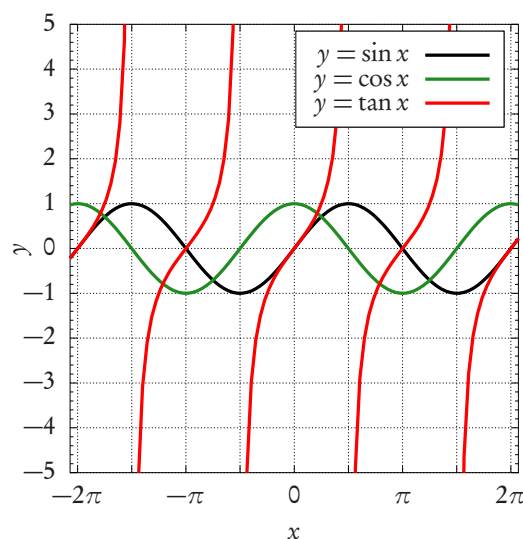
$$\sin^2 \varphi + \cos^2 \varphi = 1 \quad (1.7.36)$$

gilt.

Einige Werte für spezielle Winkel lassen sich leicht direkt aus der geometrischen Definition ablesen. Für $\varphi = \pi/4 = 45^\circ$ liegt der Punkt P auf der Winkelhalbierenden der Koordinatenachsen. Demnach gilt $\sin(\pi/4) = \cos(\pi/4)$, und wegen (1.7.36) folgt daraus wiederum $\cos^2(\pi/4) + \sin^2(\pi/4) = 2\sin^2(\pi/4) = 1$. Da offenbar $\sin(\pi/4) > 0$ ist, folgt daraus eindeutig $\sin(\pi/4) = \cos(\pi/4) = 1/\sqrt{2} = \sqrt{2}/2$. Ferner gilt $\tan(\pi/4) = \sin(\pi/4)/\cos(\pi/4) = 1$.

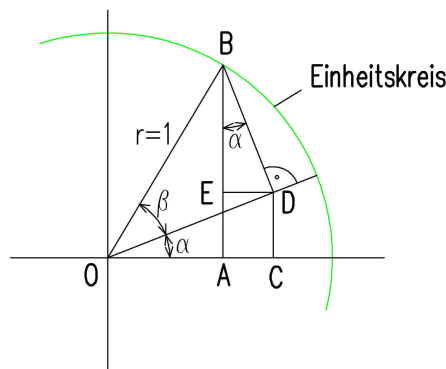
Wählt man hingegen den Punkt P für $\varphi = \pi/6 = 30^\circ$ und denkt sich das rechtwinklige Dreieck an der x -Achse gespiegelt, ergibt sich insgesamt ein gleichseitiges Dreieck, und es folgt $\sin(\pi/6) = 1/2$ und daraus $\cos(\pi/6) = \sqrt{1 - \sin^2(\pi/6)} = \sqrt{3}/2$ und folglich $\tan(\pi/6) = 1/\sqrt{3} = \sqrt{3}/3$.

Die Funktionsgraphen der trigonometrischen Funktionen $\sin$ (schwarz), $\cos$ (grün) und $\tan$ (rot) ergeben sich insgesamt zu



Als nächstes leiten wir die **Additionstheoreme** her. Betrachten wir dazu die folgende Skizze¹:

¹ von Petflo2000 - Eigenes Werk, Gemeinfrei, <https://commons.wikimedia.org/w/index.php?curid=8326298>



Aus dem Dreieck OAB liest man ab

$$\sin(\alpha + \beta) = |AB|, \quad (1.7.37)$$

$$\cos(\alpha + \beta) = |OA|, \quad (1.7.38)$$

aus dem Dreieck OCD

$$\sin \alpha = \frac{|CD|}{|OD|}, \quad (1.7.39)$$

$$\cos \alpha = \frac{|OC|}{|OD|}, \quad (1.7.40)$$

aus dem Dreieck ODB

$$\sin \beta = |BD|, \quad (1.7.41)$$

$$\cos \beta = |OD| \quad (1.7.42)$$

und dem Dreieck EDB

$$\sin \alpha = \frac{|DE|}{|BD|}, \quad (1.7.43)$$

$$\cos \alpha = \frac{|BE|}{|BD|}. \quad (1.7.44)$$

Nun gilt offenbar unter Anwendung der Beziehungen (1.7.37-1.7.44)

$$\begin{aligned} \sin(\alpha + \beta) &= |AB| = |AE| + |EB| = |CD| + |EB| = |OD| \sin \alpha + |BD| \cos \alpha \\ &= \sin \alpha \cos \beta + \cos \alpha \sin \beta, \end{aligned} \quad (1.7.45)$$

$$\begin{aligned} \cos(\alpha + \beta) &= |OA| = |OC| - |AC| = |OC| - |DE| = |OD| \cos \alpha - |BD| \sin \alpha \\ &= \cos \alpha \cos \beta - \sin \alpha \sin \beta. \end{aligned} \quad (1.7.46)$$

Für den $\tan$ ergibt sich daraus

$$\tan(\alpha + \beta) = \frac{\sin(\alpha + \beta)}{\cos(\alpha + \beta)} = \frac{\sin \alpha \cos \beta + \cos \alpha \sin \beta}{\cos \alpha \cos \beta - \sin \alpha \sin \beta} = \frac{\tan \alpha + \tan \beta}{1 - \tan \alpha \tan \beta}. \quad (1.7.47)$$

Dabei haben wir im letzten Schritt den Bruch durch $\cos \alpha \cos \beta$ gekürzt.

Für $\alpha = \beta$ ergeben sich speziell die oft nützlichen **Doppelwinkelformeln**

$$\sin(2\alpha) = 2 \sin \alpha \cos \alpha, \quad \cos(2\alpha) = \cos^2 \alpha - \sin^2 \alpha, \quad \tan(2\alpha) = \frac{2 \tan \alpha}{1 - \tan^2 \alpha}. \quad (1.7.48)$$

Die Doppelwinkelformel für den $\cos$ können wir mittels (1.7.36) noch weiter umformen zu

$$\cos(2\alpha) = 1 - 2\sin^2 \alpha \Rightarrow \sin^2 \alpha = \frac{1}{2}[1 - \cos(2\alpha)] \quad (1.7.49)$$

oder auch

$$\cos(2\alpha) = 2\cos^2 \alpha - 1 \Rightarrow \cos^2 \alpha = \frac{1}{2}[1 + \cos(2\alpha)]. \quad (1.7.50)$$

Um die Ableitungen der trigonometrischen Funktionen zu berechnen, wollen wir zwei Wege angeben. Die erste geometrische Methode ist der anschaulichere Weg, benötigt aber die Vorwegnahme der Vektorrechnung. Ausgangspunkt ist die Definition von $\sin$ und $\cos$ am Einheitskreis, wonach $(x(\varphi), y(\varphi))$ als Koordinaten in einem kartesischen Koordinatensystem betrachtet die Punkte auf dem Einheitskreis beschreiben, wobei $\varphi \in [0, 2\pi)$ alle Punkte des Kreises genau einmal erfasst. Nach dem Satz des Pythagoras gilt also

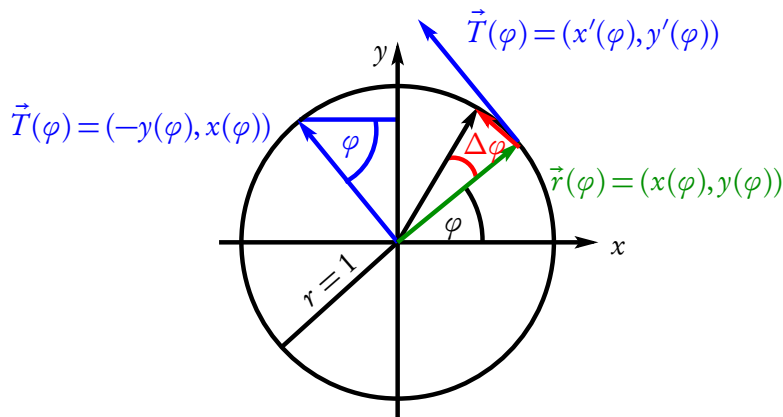
$$x^2(\varphi) + y^2(\varphi) = 1. \quad (1.7.51)$$

Leiten wir diese Funktion nach φ ab finden wir unter Verwendung der Kettenregel auf der linken Seite

$$\frac{d}{d\varphi}[x^2(\varphi) + y^2(\varphi)] = 2x(\varphi)x'(\varphi) + 2y(\varphi)y'(\varphi) = \frac{d}{d\varphi}1 = 0. \quad (1.7.52)$$

Es gilt also

$$x(\varphi)x'(\varphi) + y(\varphi)y'(\varphi) = 0. \quad (1.7.53)$$



Dies können wir nun geometrisch interpretieren, wenn wir wie in der obigen Skizze (x, y) als die Komponenten eines Vektors betrachten, also als die gerichtete Strecke vom Ursprung $(0, 0)$ des Koordinatensystems zum Punkt (x, y) , den wir als den entsprechenden Pfeil eingezeichnet haben. Betrachtet man dann zunächst den Differenzenquotienten $(\Delta x, \Delta y)/\Delta\varphi$, sieht man, dass die Richtung dieses Vektors der Richtung der Sekante entspricht, die die Punkte $(x(\varphi), y(\varphi))$ und $(x(\varphi + \Delta\varphi), y(\varphi + \Delta\varphi))$ verbindet. Lässt man dann $\Delta\varphi \rightarrow 0$ gehen, ergibt sich die Ableitung des Vektors, also $(x'(\varphi), y'(\varphi))$, und dieser Vektor weist offenbar in die Richtung der Tangente an den Kreis im Punkt $(x(\varphi), y(\varphi))$. Anschaulich ist klar, dass die Tangente in einem beliebigen Punkt auf dem Kreis immer senkrecht auf dem entsprechenden Radius steht. Genau das ist aber auch die Aussage von (1.7.53). Ein Einheitsvektor, der senkrecht auf dem Einheitsvektor $\vec{n}_1 = (x, y)$ steht, ist nämlich durch $\vec{n}_2 = (-y, x)$ gegeben, wie ebenfalls aus der Skizze zu entnehmen ist, wenn man den Tangentenvektor (blau) so verschiebt, dass sein Anfangspunkt im Koordinatenursprung zu liegen kommt. Schließlich ist $d\varphi \sqrt{[x'(\varphi)]^2 + [y'(\varphi)]^2}$ die Länge eines infinitesimalen Kreisbogenstückchens, das $(x(\varphi + d\varphi), y(\varphi + d\varphi))$ mit $(x(\varphi), y(\varphi))$ verbindet. Diese Länge ist aber gemäß der Definition des Winkels φ im Bogenmaß gerade $d\varphi$. Demnach ist der Tangentenvektor $\vec{T}(\varphi) = (x'(\varphi), y'(\varphi))$ tatsächlich ein Vektor der Länge 1. Auch die

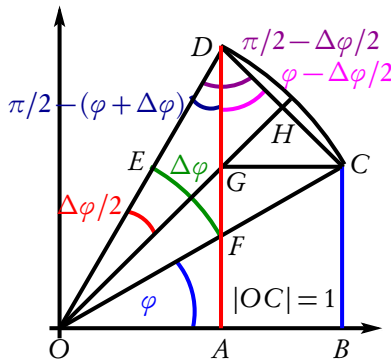
Richtung wird anhand der Grenzwertbildung von den Differenzenquotienten zur Ableitung, also vom Richtungsvektor der Sekanten zum Richtungsvektor der Tangenten klar. Demnach ist eindeutig

$$(\cos'(\varphi), \sin'(\varphi)) = (-\sin \varphi, \cos \varphi), \quad (1.7.54)$$

d.h. die Ableitungen für $\cos$ und $\sin$ sind durch

$$\cos'(x) = -\sin x, \quad \sin'(x) = \cos x \quad (1.7.55)$$

gegeben, wobei wir für das Argument jetzt x statt φ geschrieben haben.



Die zweite Berechnung der Ableitung des Sinus' verwendet ebenfalls die geometrische Definition der trigonometrischen Funktionen als Seitenverhältnisse eines rechtwinkligen Dreiecks, um die Ableitung des Sinus' als Grenzwert des entsprechenden Differenzenquotienten zu finden. Dazu folgen wir [Sch84] und betrachten die nebenstehende Abbildung. Die Herleitung beruht auf der Ungleichung

$$|OF|\Delta\varphi < |CD| < \Delta\varphi. \quad (1.7.56)$$

Dabei haben wir benutzt, dass die Länge des grünen Bogens $|OF|\Delta\varphi$ und die des schwarzen Kreisbogens $|OC|\Delta\varphi = \Delta\varphi$ ist, da definitionsgemäß $|OC| = 1$ ist. Die eingezeichneten Winkel $\angle(ODC)$ und $\angle(ODA)$ folgen jeweils aus dem Winkelsummensatz, d.h. dass für ein rechtwinkliges Dreieck die Summe der beiden nicht-rechtwinkligen Winkel $\pi/2$ ist. Daraus folgt dann $\angle(GDC) = \angle(ODC) - \angle(ODA) = \varphi - \Delta\varphi/2$, wie in der Zeichnung angegeben. Am rechtwinkligen Dreieck OAF lesen wir nun ab, dass $|OA|/|OF| = \cos \varphi$ und am rechtwinkligen Dreieck OAD , dass $|OA|/|OD| = |OA| = \cos(\varphi + \Delta\varphi)$ ist. Damit folgt

$$|OF| = \frac{|OA|}{\cos \varphi} = \frac{\cos(\varphi + \Delta\varphi)}{\cos \varphi}. \quad (1.7.57)$$

Weiter ist $|DG| = |AD| - |AG| = |AD| - |BC| = \sin(\varphi + \Delta\varphi) - \sin \varphi$ und damit folgt aus $|DG|/|DC| = \cos(\varphi + \Delta\varphi/2)$, dass

$$|DC| = \frac{|DG|}{\cos(\varphi + \Delta\varphi/2)} = \frac{\sin(\varphi + \Delta\varphi) - \sin \varphi}{\cos(\varphi + \Delta\varphi/2)}. \quad (1.7.58)$$

Damit ergibt sich aus (1.7.56) die Ungleichung

$$\frac{\cos(\varphi + \Delta\varphi)}{\cos \varphi} \Delta\varphi < \frac{\sin(\varphi + \Delta\varphi) - \sin \varphi}{\cos(\varphi + \Delta\varphi/2)} < \Delta\varphi. \quad (1.7.59)$$

multiplizieren wir diese Ungleichung mit $\cos(\varphi + \Delta\varphi/2)/\Delta\varphi$ folgt

$$\frac{\cos(\varphi + \Delta\varphi)}{\cos \varphi} \cos(\varphi + \Delta\varphi/2) < \frac{\sin(\varphi + \Delta\varphi) - \sin \varphi}{\Delta\varphi} < \cos(\varphi + \Delta\varphi/2). \quad (1.7.60)$$

Für $\Delta\varphi \rightarrow 0$ ergibt sich daraus, dass

$$\sin' \varphi = \lim_{\Delta\varphi \rightarrow 0} \frac{\sin(\varphi + \Delta\varphi) - \sin \varphi}{\Delta\varphi} = \cos \varphi \quad (1.7.61)$$

ist. Dabei haben wir aufgrund der geometrischen Definition der cos-Funktion angenommen, dass diese stetig ist.

Weiter folgt aus der Definition der Winkelfunktionen am Einheitskreis und dem Winkelsummensatz für rechtwinklige Dreiecke, dass $\cos \varphi = \sin(\pi/2 - \varphi)$ ist. Mit der Kettenregel und der eben hergeleiteten Ableitungsformel für sin (1.7.61) folgt weiter

$$\cos' \varphi = \frac{d}{d\varphi} \sin(\pi/2 - \varphi) = -\sin'(\pi/2 - \varphi) = -\cos(\pi/2 - \varphi) = -\sin \varphi. \quad (1.7.62)$$

Dabei haben wir im letzten Schritt verwendet, dass $\cos(\pi/2 - \varphi) = \sin \varphi$ ist.

Aus den Ableitungsformeln für sin und cos folgen noch die wichtigen Grenzwerte

$$\sin' 0 = \cos 0 = 1 = \lim_{x \rightarrow 0} \frac{\sin x}{x}, \quad (1.7.63)$$

$$\cos' 0 = \sin 0 = 0 = \lim_{x \rightarrow 0} \frac{\cos x - 1}{x}. \quad (1.7.64)$$

Mit (1.7.61) und (1.7.62) können wir mit Hilfe der Quotientenregel auch den Tangens ableiten. Für $x \neq (2n+1)\pi/2, n \in \mathbb{Z}$ folgt

$$\frac{d}{dx} \tan x = \frac{d}{dx} \frac{\sin x}{\cos x} = \frac{\cos^2 x + \sin^2 x}{\cos^2 x} = \frac{1}{\cos^2 x} = 1 + \tan^2 x. \quad (1.7.65)$$

Die trigonometrischen Funktionen sind periodisch² und haben daher keine eindeutig definierten **Umkehrfunktionen**. Gewöhnlich definiert man die Umkehrfunktion auf einem eingeschränkten Intervall für den Winkel. Beginnen wir mit dem Kosinus: Dieser ist im Intervall $[0, \pi]$ streng monoton fallend. Entsprechend definieren wir als Umkehrfunktion $\arccos: [-1, 1] \rightarrow [0, \pi]$ durch $\cos(\arccos x) = x$. Wegen der strikten Monotonie des cos im Intervall $[0, \pi]$ ist durch die Einschränkung $\arccos x \in [0, \pi]$ der Wert der arccos-Funktion eindeutig bestimmt. Mit dem Satz von der Ableitung der Umkehrfunktion folgt aus $y = \arccos x$:

$$\frac{d}{dx} \arccos x = \frac{1}{\cos' y} = -\frac{1}{\sin y}. \quad (1.7.66)$$

Nun ist für $y \in [0, \pi]$ der Sinus positiv, so dass dort wegen (1.7.36) $\sin y = +\sqrt{1 - \cos^2 y}$ ist. Für $x \in (-1, 1)$, also $y \in (0, \pi)$ ist $\sin y \neq 0$, und folglich gilt

$$\frac{d}{dx} \arccos x = -\frac{1}{\sqrt{1 - \cos^2 y}} = -\frac{1}{\sqrt{1 - x^2}}. \quad (1.7.67)$$

Demnach haben wir die Ableitungsregel

$$\arccos' x = -\frac{1}{\sqrt{1 - x^2}}, \quad x \in (-1, 1). \quad (1.7.68)$$

Die Ableitung divergiert offensichtlich für $x \rightarrow \pm 1$, d.h. der Funktionsgraph des arccos besitzt für $x \rightarrow \pm 1$ senkrechte Tangenten.

²sin und cos haben die Periode 2π und tan die Periode π .

Die Umkehrfunktion des Sinus definieren wir im Winkelbereich $x \in [-\pi/2, \pi/2]$, wo der Sinus strikt monoton wachsend ist. Dann ist $\arcsin : [-1, 1] \rightarrow [-\pi/2, \pi/2]$ durch $\sin(\arcsin x) = x$ eindeutig definiert. Analog wie beim Kosinus zeigt man durch Ableiten der Umkehrfunktion (*Übung!*)

$$\arcsin' x = +\frac{1}{\sqrt{1-x^2}}, \quad x \in (-\pi/2, \pi/2). \quad (1.7.69)$$

Der Tangens ist im offenen Intervall $(-\pi/2, \pi/2)$ strikt monoton wachsend und nimmt beliebige reelle Werte an, denn $\lim_{x \rightarrow \pm\pi/2} \tan x = \pm\infty$. Entsprechend definieren wir die Umkehrfunktion $\arctan : \mathbb{R} \rightarrow (-\pi/2, \pi/2)$ eindeutig durch $\tan(\arctan x) = x$. Über die Ableitung der Umkehrfunktion finden wir mit (1.7.65) (*Übung*)

$$\arctan' x = \frac{1}{1+x^2}. \quad (1.7.70)$$

1.7.4 Kurvendiskussionen, Mittelwertsatz der Differentialrechnung

Wir besprechen als wichtige Anwendung der Differentialrechnung die Untersuchung der **lokalen Eigenschaften** von Funktionen einer reellen Veränderlichen.

Eine Funktion $f : (a, b) \rightarrow \mathbb{R}$ besitzt in $\xi \in (a, b)$ ein **lokales Maximum (lokales Minimum)**, wenn es ein $\epsilon > 0$ gibt, so dass für alle $x \in (a, b)$ mit $|x - \xi| < \epsilon$ stets $f(x) \leq f(\xi)$ ($f(x) \geq f(\xi)$) gilt. Besitzt f in ξ ein lokales Maximum oder Minimum, sagt man, dass f in ξ ein **lokales Extremum** aufweist. Es gilt der **Satz**: Sei $f : (a, b) \rightarrow \mathbb{R}$ stetig und wenigstens einmal differenzierbar. Besitzt dann f in $\xi \in (a, b)$ ein lokales Extremum, so gilt $f'(\xi) = 0$.

Zum Beweis nehmen wir an, f besitze ein lokales Maximum bei $x = \xi$. Da f voraussetzungsgemäß differenzierbar in $x = \xi$ ist, gilt

$$f'(\xi) = \lim_{h \rightarrow 0^+} \frac{f(\xi + h) - f(\xi)}{h} \leq 0, \quad (1.7.71)$$

denn voraussetzungsgemäß nimmt bei Einschränkung auf eine hinreichend kleine Umgebung f an der Stelle ξ den größten Wert an. Ebenso argumentiert man aber, dass

$$f'(\xi) = \lim_{h \rightarrow 0^-} \frac{f(\xi + h) - f(\xi)}{h} \geq 0, \quad (1.7.72)$$

Nun können aber (1.7.71) und (1.7.72) nur beide gelten, wenn $f'(\xi) = 0$. Analog zeigt man die Behauptung, falls die Funktion ein Minimum in (a, b) besitzt.

Bemerkung: Die Bedingung $f'(\xi) = 0$ ist nur *notwendig aber nicht hinreichend* für das Vorliegen eines Extremums. Die Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = x^3$ ist z.B. strikt monoton wachsend, besitzt also sicher kein lokales Extremum. Andererseits gilt $f'(x) = 3x^2$ und folglich $f'(0) = 0$. Die Ableitung von f verschwindet bei $x = 0$, besitzt dort aber weder ein lokales Minimum noch ein lokales Maximum.

Satz von Rolle (Michel Rolle, 1652-1719): Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig auf (a, b) differenzierbar. Gilt dann $f(a) = f(b)$, existiert mindestens eine Stelle $\xi \in (a, b)$ mit $f'(\xi) = 0$.

Beweis: Ist $f(x) = \text{const}$, ist die Behauptung trivial, denn dann ist $f'(x) = 0$ für alle $x \in (a, b)$. Falls dies nicht der Fall ist, gibt es wenigstens eine Stelle x_0 mit $f(x_0) > f(a) = f(b)$ oder $f(x_0) < f(a) = f(b)$. Da die Funktion f in dem abgeschlossenen Intervall $[a, b]$ voraussetzungsgemäß stetig ist, nimmt sie nach dem Satz vom Maximum und Minimum (vgl. Abschnitt 1.6) irgendwo in diesem Intervall das Infimum und das Supremum ihres Bildbereiches an. Im ersten Fall muss offenbar das Supremum und im zweiten Fall das Infimum an wenigstens einer Stelle ξ im offenen Intervall (a, b) angenommen werden, d.h. es liegt allemal ein Extremum in (a, b) vor, und nach dem eben bewiesenen Satz ist $f'(\xi) = 0$.

Schließlich gilt der **Mittelwertsatz der Differentialrechnung**: Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und auf (a, b) differenzierbar. Dann gibt es eine Stelle $\xi \in (a, b)$ mit

$$\frac{f(b) - f(a)}{b - a} = f'(\xi). \quad (1.7.73)$$

Zum **Beweis** wenden wir den Satz von Rolle auf die Funktion

$$g(x) = f(x) - \frac{f(b) - f(a)}{b - a}(x - a) \quad (1.7.74)$$

an. In der Tat erfüllt diese Funktion die Voraussetzungen des Satzes von Rolle. Da f auf $[a, b]$ stetig und auf (a, b) differenzierbar ist, gilt dies sicher auch für g . Außerdem ist $g(a) = g(b) = f(a)$. Folglich gibt es eine Stelle $\xi \in (a, b)$ mit $g'(\xi) = 0$. Nun ist aber

$$g'(x) = f'(x) - \frac{f(b) - f(a)}{b - a} \Rightarrow 0 = g'(\xi) = f'(\xi) - \frac{f(b) - f(a)}{b - a}, \quad (1.7.75)$$

und das war zu zeigen.

Eine weitere Verallgemeinerung ist

Seien f und g auf einem abgeschlossenen Intervall $[a, b]$ stetig und im offenen Intervall (a, b) differenzierbar. Dann existiert $\xi \in (a, b)$, so dass

$$[f(b) - f(a)]g'(\xi) = [g(b) - g(a)]f'(\xi). \quad (1.7.76)$$

Falls $g'(x) \neq 0$ für $x \in (a, b)$, gilt für mindestens ein $\xi \in (a, b)$

$$\frac{f(b) - f(a)}{g(b) - g(a)} = \frac{f'(\xi)}{g'(\xi)}. \quad (1.7.77)$$

Der **Beweis** folgt sofort durch Anwendung des Satzes von Rolle auf die Funktion

$$F(x) = [f(b) - f(a)]g(x) - [g(b) - g(a)]f(x). \quad (1.7.78)$$

Satz von der Monotonie: Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und auf (a, b) differenzierbar. Dann ist f genau dann strikt monoton wachsend (monoton fallend) falls $f'(x) > 0$ ($f'(x) < 0$).

Beweis: Sei f strikt monoton wachsend und $x \in (a, b)$. Dann gibt es ein $\epsilon > 0$, so dass $x + h \in [a, b]$ für alle $h \in (-\epsilon, \epsilon)$. Für diese h gilt dann wegen des strikt monotonen Wachstums von f

$$\frac{f(x + h) - f(x)}{h} > 0. \quad (1.7.79)$$

Da f in (a, b) differenzierbar ist, erhält man durch Grenzwertbildung $h \rightarrow 0$, dass $f'(x) > 0$.

Sei nun umgekehrt $f'(x) > 0$. Wir nehmen nun an, dass f nicht strikt monoton wachsend ist. Dann gibt es zwei Stellen $x_1 < x_2 \in (a, b)$ mit $f(x_1) \geq f(x_2)$. Nach dem oben bewiesenen Mittelwertsatz der Differentialrechnung gibt es dann eine Stelle $\xi \in (a, b)$ mit

$$f'(\xi) = \frac{f(x_2) - f(x_1)}{x_2 - x_1} \leq 0, \quad (1.7.80)$$

und das steht im Widerspruch zur Voraussetzung, dass $f'(x) > 0$ im ganzen offenen Intervall (a, b) gilt. Damit muss also f strikt monoton wachsend sein.

Der Beweis für strikt monoton fallende Funktionen erfolgt analog.

Schließlich geben wir noch ein *hinreichendes Kriterium* für ein lokales Extremum an:

Sei $f : [a, b]$ stetig und auf (a, b) zweimal differenzierbar. Ist dann $f'(\xi) = 0$ für ein $\xi \in (a, b)$ und $f''(\xi) < 0$ ($f''(\xi) > 0$), so besitzt f bei $x = \xi$ ein lokales Maximum (Minimum).

Wir beweisen die Behauptung für den Fall $f''(\xi) < 0$. Es gilt

$$f''(\xi) = \lim_{h \rightarrow 0} \frac{f'(\xi + h) - f'(\xi)}{h} = \lim_{h \rightarrow 0} \frac{f'(\xi + h)}{h} < 0. \quad (1.7.81)$$

Dies bedeutet, dass es ein $\epsilon > 0$ gibt, so dass $f'(\xi + h)/h < 0$ für $h \in (-\epsilon, \epsilon) \setminus \{0\}$. Damit ist aber $f'(x) < 0$ für $x \in (\xi, \xi + \epsilon]$ und $f'(x) > 0$ für $x \in [\xi - \epsilon, \xi)$. Damit ist f zumindest in einer kleinen Umgebung von ξ für $x < \xi$ strikt monoton wachsend und für $x > \xi$ strikt monoton fallend. Damit muss f bei ξ ein lokales Maximum besitzen.

Der Beweis, dass für $f''(\xi) > 0$ und $f'(\xi) = 0$ ein Minimum vorliegt, folgt analog.

Um zu zeigen, dass das obige Kriterium *hinreichend aber nicht notwendig* ist, betrachten wir die Funktion $f(x) = x^4$. Sie besitzt offenbar bei $x = 0$ ein lokales (sogar ein absolutes) Minimum. In der Tat ist $f'(x) = 4x^3$ und damit $f'(0) = 0$, wie es der Satz von lokalen Extrema differenzierbarer Funktionen sagt. Allerdings gilt weiter $f''(x) = 12x^2$ und folglich $f''(0) = 0$, d.h. obwohl ein Minimum vorliegt ist $f''(0)$ *nicht* > 0 .

1.8 Die Regeln von de L'Hospital

Eine der Regeln von de L'Hospital Guillaume François Antoine, Marquis de L'Hospital (1661–1704) erlaubt die Berechnung von Grenzwerten von Funktionen vom Typ " $0/0$ " oder " ∞/∞ ":

Seien f und g auf dem offenen Intervall (a, b) differenzierbar und $g'(x) \neq 0$ für $x \in (a, b)$. Ferner sei eine der beiden Bedingungen

$$(1) \lim_{x \rightarrow a+0^+} f(x) = \lim_{x \rightarrow a+0^+} g(x) = 0,$$

$$(2) \lim_{x \rightarrow a+0^+} f(x) = \pm\infty, \lim_{x \rightarrow a+0^+} g(x) = \pm\infty$$

erfüllt. Dann gilt

$$\lim_{x \rightarrow a+0^+} \frac{f(x)}{g(x)} = \lim_{x \rightarrow a+0^+} \frac{f'(x)}{g'(x)}, \quad (1.8.1)$$

falls der letztgenannte Grenzwert existiert oder $\pm\infty$ ist.

Beweis: Sei zunächst die Voraussetzung (1) erfüllt: Offenbar können wir $f(a) = g(a) = 0$ setzen. Dann sind die Voraussetzungen des verallgemeinerten Mittelwertsatzes der Differentialrechnung in der Form (1.7.77) erfüllt, d.h. es gibt für $x \in (a, b)$ eine Zahl $\xi \in (a, x)$, so dass

$$\frac{f(x)}{g(x)} = \frac{f(x) - f(a)}{g(x) - g(a)} = \frac{f'(\xi)}{g'(\xi)}. \quad (1.8.2)$$

Da mit $x \rightarrow a + 0^+$ auch $\xi \rightarrow a + 0^+$ gilt, ist für diesen Fall die Behauptung bewiesen.

Ist die Bedingung (2) erfüllt, gibt es offenbar ein $x_0 \in (a, b)$ so dass $f(x) \neq 0$ und $g(x) \neq 0$ für alle $x \in (a, x_0]$. Einerseits gibt es dann gemäß dem verallgemeinerten Mittelwertsatzes der Integralrechnung für $x \in (a, x_0)$ ein ξ zwischen x und x_0 mit

$$\frac{f(x) - f(x_0)}{g(x) - g(x_0)} = \frac{f'(\xi)}{g'(\xi)}. \quad (1.8.3)$$

Andererseits können wir

$$\frac{f(x) - f(x_0)}{g(x) - g(x_0)} = \frac{f(x)}{g(x)} \frac{1 - f(x_0)/f(x)}{1 - g(x_0)/g(x)}. \quad (1.8.4)$$

Setzt man (1.8.3) und (1.8.4), folgt

$$\frac{f(x)}{g(x)} = \frac{f'(\xi)}{g'(\xi)} \frac{1 - g(x_0)/g(x)}{1 - f(x_0)/f(x)}. \quad (1.8.5)$$

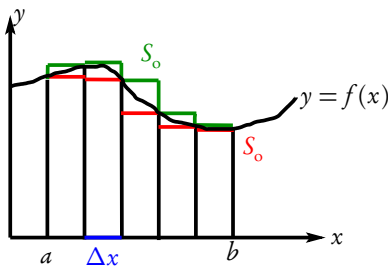
Für $x \rightarrow a + 0^+$ folgt dann die Behauptung.

1.9 Integralrechnung

Die Integralrechnung ist, wie wir sehen werden, in gewissem Sinne die Umkehrung der Differentialrechnung. Sie beschäftigt sich generell mit dem geometrischen Inhaltsbegriff (Bogenlänge einer Kurve, Flächeninhalte von Flächen, Volumen von Körpern) und ist daher eng mit der **Maßtheorie** verknüpft. In dieser Vorlesung behandeln wir nur das **Riemann-Integral**. Für an der modernen Lebesgue-Integrationstheorie Interessierte sei auf [Wei80] verwiesen, wo man eine Darstellung findet, die nicht die allgemeine Maßtheorie voraussetzt.

1.9.1 Definition des Riemann-Integrals

Das Ziel der Integralrechnung für reellwertige Funktionen einer reellen Veränderlichen ist, anschaulich formuliert, die Berechnung des Flächeninhaltes der Fläche unter dem Graphen der Funktion $f : [a, b] \rightarrow \mathbb{R}$, wobei $[a, b] \subset \mathbb{R}$ ein endliches Intervall ist. Zunächst setzen wir nur voraus, dass die Funktion f beschränkt sei, d.h. es gibt Zahlen m und M , so dass $m \leq f(x) \leq M$ für alle $x \in [a, b]$. Die grundlegende Idee, den Flächeninhalt zu bestimmen, ist, das Intervall $[a, b]$ in n Teilintervalle $[x_{j-1}, x_j]$ mit $j \in \{0, 1, \dots, n\}$ zu zerlegen, wobei $a = x_0 < x_1 < \dots < x_n = b$ ist und die Funktion durch eine entsprechende **Treppenfunktion** zu approximieren. Dabei ist eine Treppenfunktion eine Funktion, die auf jedem Teilintervall $[x_{j-1}, x_j]$ konstant ist.



Nun kann man für jedes Teilintervall das Infimum bzw. das Supremum der Funktionswerte in diesem Intervall verwenden. Nach dem Satz vom Infimum und Supremum existieren nämlich beide Größen, weil ja voraussetzungsgemäß die Funktion f beschränkt ist.

Dann bezeichnen wir als **Unter- und Obersumme** der Funktion f bzgl. der Zerlegung $[x_{j-1}, x_j]$ des Intervalls $[a, b]$ die Größen

$$\begin{aligned} S_u &= \sum_{j=1}^n (x_j - x_{j-1}) \inf\{f([x_{j-1}, x_j])\}, \\ S_o &= \sum_{j=1}^n (x_j - x_{j-1}) \sup\{f([x_{j-1}, x_j])\}. \end{aligned} \quad (1.9.1)$$

In der nebenstehenden Skizze entspricht der Untersumme die Fläche unter der grün eingezeichneten und der Obersumme die Fläche unter der rot eingezeichneten Treppenfunktion.

Es ist anschaulich klar, dass der Flächeninhalt der Fläche unter der Kurve für alle Zerlegungen des Intervalls $[a, b]$ stets zwischen diesen beiden Werten liegen muss und dass wir den Flächeninhalt erhalten, indem wir die Unterteilung immer feiner machen, d.h. wir lassen die Anzahl der Intervalle $n \rightarrow \infty$ gehen, wobei zugleich $\max_{j \in \{1, \dots, n\}} (x_j - x_{j-1}) \rightarrow 0$ gehen soll.

Der Flächeninhalt existiert dann sicher, wenn S_u und S_o gegen denselben Wert streben. Wir sagen dann, die Funktion f sei über das Intervall $[a, b]$ **Riemannintegrierbar**, und wir definieren das Integral als den entsprechenden Grenzwert

$$\int_a^b dx f(x) = \lim_{n \rightarrow \infty} S_u = \lim_{n \rightarrow \infty} S_o. \quad (1.9.2)$$

Da bei einer Verfeinerung der Zerlegung von $[a, b]$ die Untersummen stets größer und die Obersummen stets kleiner werden (*warum?*), kann man das Integral auch als $\sup S_u = \inf S_o$ definieren, wobei Supremum bzw. Infimum über alle möglichen Zerlegungen zu nehmen ist.

Wir bemerken sogleich, dass wenn $f, g : [a, b] \rightarrow \mathbb{R}$ Riemann-integrierbar sind, auch die Linearkombination $\lambda_1 f + \lambda_2 g$ für beliebige $\lambda_1, \lambda_2 \in \mathbb{R}$ Riemann-integrierbar sind und dass dann

$$\int_a^b dx [\lambda_1 f(x) + \lambda_2 g(x)] = \lambda_1 \int_a^b dx f(x) + \lambda_2 \int_a^b dx g(x) \quad (1.9.3)$$

gilt. Der einfache Beweis sei dem Leser zur *Übung* überlassen.

Nehmen wir weiter an $f : [a, c] \rightarrow \mathbb{R}$ sei über die Intervalle $[a, b]$ und $[b, c]$ Riemann-integrierbar. Man zeigt dann leicht (*Übung*), dass f dann auch über $[a, c]$ Riemann-integrierbar ist und dass dann

$$\int_a^c dx f(x) = \int_a^b dx f(x) + \int_b^c dx f(x) \quad (1.9.4)$$

ist.

Definition: Sei $f : [a, b] \rightarrow \mathbb{R}$ eine beliebige Funktion. Dann definieren wir

$$\begin{aligned} f_+(x) &= \begin{cases} f(x) & \text{falls } f(x) > 0, \\ 0 & \text{falls } f(x) \leq 0, \end{cases} \\ f_-(x) &= \begin{cases} -f(x) & \text{falls } f(x) < 0, \\ 0 & \text{falls } f(x) \geq 0. \end{cases} \end{aligned} \quad (1.9.5)$$

Offenbar gilt dann $f = f_+ - f_-$ und $|f| = f_+ + f_-$.

Satz: Sei $f : [a, b] \rightarrow \mathbb{R}$ eine Riemann-integrierbare Funktion. Dann sind auch f_+ , f_- , $|f|$ und f^2 Riemann-integrierbar.

Beweis: Nach Voraussetzung ist $\sup S_u = \inf S_o$, wobei Supremum und Infimum der Ober- bzw. Untersumme stets über alle Zerlegungen des Intervalls $[a, b]$ zu nehmen sind. Das bedeutet aber, dass es zu jedem $\epsilon > 0$ Treppenfunktionen φ und ψ gibt, so dass $\varphi \leq f \leq \psi$ ist, so dass

$$0 \leq \int_a^b dx [\psi(x) - \varphi(x)] < \epsilon \quad (1.9.6)$$

ist. Wegen $\varphi \leq f \leq \psi$ folgt aus (1.9.5) sofort, dass $\varphi_+ \leq f_+ \leq \psi_+$ und $\psi_- \leq f_- \leq \varphi_-$. Nun sind $\varphi_{\pm}$ und $\psi_{\pm}$ wieder Treppenfunktionen, und es gilt

$$\begin{aligned} 0 &\leq \int_a^b dx [\psi_+(x) - \varphi_+(x)] \leq \int_a^b dx [\psi(x) - \varphi(x)] < \epsilon, \\ 0 &\leq \int_a^b dx [\varphi_-(x) - \psi_-(x)] \leq \int_a^b dx [\psi(x) - \varphi(x)] < \epsilon. \end{aligned} \quad (1.9.7)$$

Zu jedem $\epsilon > 0$ gibt es also Zerlegungen des Intervalls, so dass Ober- und Untersumme von $f_{\pm}$ um weniger als ϵ voneinander abweichen. Folglich stimmen das Supremum der Unter- und das Infimum der Obersummen überein, und die Funktionen $f_{\pm}$ sind folglich Riemann-integrierbar. Da mit ist auch $|f| = f_+ + f_-$ Riemann-integrierbar.

Weiter ist $f^2 = |f|^2$. Da $|f|$ Riemann-integrierbar ist, gibt es zu jedem $\epsilon > 0$ Treppenfunktionen $\tilde{\psi}, \tilde{\varphi} : [a, b] \rightarrow \mathbb{R}$ mit $\tilde{\varphi} \leq |f| \leq \tilde{\psi}$, so dass

$$\int_a^b dx [\psi(x) - \varphi(x)] < \frac{\epsilon}{2M} \quad \text{mit} \quad M = \sup_{x \in [a, b]} |f(x)| - \inf_{x \in [a, b]} |f(x)|. \quad (1.9.8)$$

Dabei gehen wir davon aus, dass f nicht auf $[a, b]$ konstant ist. In dem Fall wäre auch f^2 konstant und somit auch f^2 integrierbar, und die Behauptung des Satzes also erfüllt. Da nun die Funktion $g(x) = x^2$ für $x \geq 0$ monoton wachsend ist, gilt auch

$$\tilde{\varphi}^2 \leq |f|^2 \leq \tilde{\psi}^2, \quad (1.9.9)$$

und $\tilde{\varphi}$ und $\tilde{\psi}$ sind Treppenfunktionen. Nach dem Mittelwertsatz der Differentialrechnung, angewandt auf $g(x) = x^2$, gilt für beliebige Zahlen $\tilde{\psi}^2 - \tilde{\varphi}^2 = g'(\xi)(\psi - \varphi) = 2\xi(\psi - \varphi) \leq 2M(\psi - \varphi)$ mit $\varphi \leq \xi \leq \psi$. Damit ist also

$$0 \leq \int_a^b dx [\tilde{\psi}^2(x) - \tilde{\varphi}^2(x)] \leq 2M \int_a^b dx [\tilde{\psi}(x) - \tilde{\varphi}(x)] \leq \epsilon. \quad (1.9.10)$$

Folglich gibt es zu jedem ϵ Treppenfunktionen $\tilde{\varphi}^2$ und $\tilde{\psi}^2$ mit $\tilde{\varphi}^2 \leq f^2 \leq \tilde{\psi}^2$, die (1.9.10) erfüllen.

Satz: Seien weiter $f, g : [a, b] \rightarrow \mathbb{R}$ Riemann-integrierbare Funktionen. Dann ist auch $f g$ Riemann-integrierbar.

Beweis: Nach den obigen Sätzen sind $f + g$ und $f - g$ sowie $(f + g)^2$ und $(f - g)^2$ Riemann-integrierbar. Damit ist aber auch

$$f g = \frac{1}{4}[(f + g)^2 - (f - g)^2] \quad (1.9.11)$$

Riemann-integrierbar.

Satz: Stetige Funktionen $f : [a, b] \rightarrow \mathbb{R}$ sind Riemann-integrierbar.

Beweis: Dazu zeigen wir zuerst, dass in diesem Fall f sogar **gleichmäßig** stetig ist, d.h. zu jedem $\epsilon > 0$ existiert ein $\delta > 0$, so dass für alle $x, x' \in [a, b]$ mit $|x - x'| < \delta$ stets $|f(x) - f(x')| < \epsilon$ ist. Zum Beweis nehmen wir an, dass f zwar stetig aber nicht gleichmäßig stetig ist. Dann gibt es zu einem $\epsilon > 0$ zu jedem $n \in \mathbb{N}$ Zahlen $x_n, x'_n \in [a, b]$ mit $|x_n - x'_n| < 1/n$, so dass $|f(x_n) - f(x'_n)| \geq \epsilon$. Die Zahlenfolgen $(x_n)_n$ und $(x'_n)_n$ sind beschränkt und besitzen daher konvergente Teilfolgen $(x_{n_k})_k$ bzw. $(x'_{n_k})_k$. Wegen der Voraussetzung, dass $|x_n - x'_n| < 1/n$ gilt $\xi := \lim_{k \rightarrow \infty} x_{n_k} = \lim_{k \rightarrow \infty} x'_{n_k}$. Folglich existiert auch $\lim_{k \rightarrow \infty} f(x_{n_k}) = \lim_{k \rightarrow \infty} f(x'_{n_k}) = f(\xi)$, weil f voraussetzungsgemäß stetig auf dem Intervall $[a, b]$ ist. Andererseits soll aber gemäß unserer obigen Annahme $|f(x_{n_k}) - f(x'_{n_k})| \geq \epsilon$ für alle $k \in \mathbb{N}$ gelten. Dann wären aber die Folgen $(f(x_{n_k}))_k$ und $(f(x'_{n_k}))_k$ wiederum nicht konvergent im Widerspruch zu unserer gerade durchgeführten Argumentation. Folglich muss also f auf $[a, b]$ gleichmäßig stetig sein.

Nun können wir leicht zeigen, dass eine auf $[a, b]$ stetige Funktion Riemann-integrierbar ist. Sei dazu $\epsilon > 0$ beliebig gewählt. Dann gibt es ein $\delta > 0$, so dass $|f(x) - f(x')| < \epsilon/(b-a)$ ist, wenn nur $|x - x'| < \delta$. Es sei nun $I_j = (x_{j-1}, x_j)$ eine so feine Zerlegung des Intervalls $[a, b]$, so dass $\max_j (x_j - x_{j-1}) < \delta$ ist. Weiter gibt es nach dem Satz vom Minimum und Maximum (s. Abschnitt 1.6) $\xi_j \in I_j$ und $\xi'_j \in I_j$, so dass $\inf_{x \in I_j} f(x) = f(\xi_j)$ und $\sup_{x \in I_j} f(x) = f(\xi'_j)$. Dann ist sicher $|\xi_j - \xi'_j| < \delta$ und folglich

$$0 \leq S_o - S_u = \sum_{j=1}^n (x_j - x_{j-1}) [f(\xi'_j) - f(\xi_j)] < \frac{\epsilon}{b-a} \sum_{j=1}^n (x_j - x_{j-1}) = \epsilon. \quad (1.9.12)$$

Folglich kann man durch beliebige Verfeinerung der Zerlegung Ober- und Untersumme beliebig aneinander annähern. Folglich existieren das Infimum der Ober- und das Supremum der Untersummen, und beide sind gleich. Folglich ist f über das Intervall $[a, b]$ Riemann-integrierbar.

1.9.2 Der Mittelwertsatz der Integralrechnung

Mittelwertsatz der Integralrechnung: Sei $f : [a, b] \rightarrow \mathbb{R}$ stetig und $\varphi : [a, b] \rightarrow \mathbb{R}$ mit $\varphi \geq 0$ integrierbar. Dann gibt es eine Zahl $\xi \in [a, b]$, so dass

$$\int_a^b dx f(x) \varphi(x) = f(\xi) \int_a^b dx \varphi(x). \quad (1.9.13)$$

Beweis: Nach den obigen Sätzen zur Riemann-Integrierbarkeit ist mit den obigen Voraussetzungen $f \varphi$ Riemann-integrierbar. Da f auf dem abgeschlossenen Intervall stetig ist, existieren $M = \sup_{x \in [a, b]} f(x)$ und

1.9. Integralrechnung

$m = \inf_{x \in [a, b]} f(x)$. Da weiter $\varphi \geq 0$ ist, folgt $m\varphi(x) \leq f(x)\varphi(x) \leq M\varphi(x)$, und damit

$$m \int_a^b dx \varphi(x) \leq \int_a^b f(x)\varphi(x) \leq M \int_a^b dx \varphi(x). \quad (1.9.14)$$

Da auch $\int_a^b dx \varphi(x) \geq 0$ ist, gibt es demnach eine Zahl μ mit $m \leq \mu \leq M$, so dass

$$\int_a^b dx f(x)\varphi(x) = \mu \int_a^b dx \varphi(x). \quad (1.9.15)$$

Da $f[a, b] \rightarrow \mathbb{R}$ stetig ist, nimmt f die Werte m und M tatsächlich für irgendwelche Werte $\xi_1, \xi_2 \in [a, b]$ tatsächlich an: $f(\xi_1) = m$, $f(\xi_2) = M$. Wegen des Zwischenwertsatzes für stetige Funktionen (vgl. Abschnitt 1.6) gibt es eine Zahl ξ zwischen ξ_1 und ξ_2 mit $f(\xi) = \mu$. Wegen (1.9.15) ist damit der Zwischenwertsatz der Integralrechnung bewiesen.

1.9.3 Der Hauptsatz der Analysis

Nach dem vorigen Abschnitt können wir für eine auf $[a, b]$ stetige Funktion f die **Integralfunktion**

$$F(x) = \int_a^x dx' f(x') \quad (1.9.16)$$

definieren.

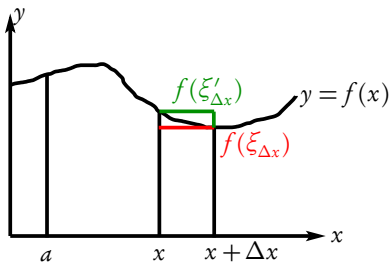
Wir wollen nun den **Hauptsatz der Analysis** beweisen, demzufolge dann F eine auf $[a, b]$ differenzierbare Funktion ist und

$$F'(x) = f(x) \quad (1.9.17)$$

gilt.

Beweis: Offenbar ist für $\Delta x > 0$

$$F(x + \Delta x) - F(x) = \int_x^{x+\Delta x} dx' f(x'). \quad (1.9.18)$$



Sei nun $m = \inf_{x' \in [x, x+\Delta x]} f(x')$ und $M = \sup_{x' \in [x, x+\Delta x]} f(x')$. Beide Werte existieren, und es gibt nach dem Satz von Bolzano-Weierstraß Zahlen $\xi_{\Delta x}, \xi'_{\Delta x} \in [x, x + \Delta x]$, so dass $m = f(\xi_{\Delta x})$ und $M = f(\xi'_{\Delta x})$. Demnach ist

$$\int_x^{x+\Delta x} dx f(\xi_{\Delta x}) \leq F(x + \Delta x) - F(x) \leq \int_x^{x+\Delta x} dx f(\xi'_{\Delta x}) \quad (1.9.19)$$

oder nach Ausführung der Integrale

$$\Delta x f(\xi_{\Delta x}) \leq F(x + \Delta x) - F(x) \leq \Delta x f(\xi'_{\Delta x}). \quad (1.9.20)$$

Dividieren wir durch Δx , erhalten wir

$$f(\xi_{\Delta x}) \leq \frac{F(x + \Delta x) - F(x)}{\Delta x} \leq f(\xi'_{\Delta x}). \quad (1.9.21)$$

Für $\Delta x \rightarrow 0$ gehen nun $\xi_{\Delta x}, \xi'_{\Delta x} \rightarrow x$ und wegen der Stetigkeit von f entsprechen $f(\xi_{\Delta x}), f(\xi'_{\Delta x}) \rightarrow f(x)$, womit (1.9.17) bewiesen ist. Analog gehen wir für $\Delta x < 0$ vor. Wir definieren dabei ganz allgemein für Integrale

$$\int_a^b dx f(x) = - \int_b^a f(x) \quad \text{falls } a < b. \quad (1.9.22)$$

Für $\Delta x < 0$ müssen wir dann nur die obige Betrachtung für das Integral

$$\int_x^{x+\Delta x} dx f(x) = - \int_{x+\Delta x}^x dx f(x) \quad (1.9.23)$$

anstellen. Dann ergibt sich auch für den Grenzwert $\Delta x \rightarrow 0$ für $\Delta x < 0$, dass $F'(x) = f(x)$ ist. Damit ist F differenzierbar und besitzt die Ableitung $F' = f$ für alle $x \in [a, b]$.

Wir bemerken noch, dass demnach für jede **Stammfunktion** $\tilde{F} : [a, b] \rightarrow \mathbb{R}$ von f , gilt

$$\int_a^b dx f(x) = \tilde{F}(b) - \tilde{F}(a). \quad (1.9.24)$$

Dabei heißt $\tilde{F}$ Stammfunktion zu f , falls

$$\tilde{F}'(x) = f(x) \quad (1.9.25)$$

gilt.

Zum **Beweis** zeigen wir, dass $\tilde{F}(x) = F(x) + \tilde{F}(a)$ ist. Es ist nämlich $\tilde{F}'(x) = F'(x) = f(x)$ und also $\tilde{F}'(x) - F'(x) = 0$. Nun muss eine auf $[a, b]$ differenzierbare Funktion, deren Ableitung überall im Intervall $[a, b]$ verschwindet, konstant sein. Da $F(a) = 0$ ist, gilt also die Behauptung.

Man bezeichnet eine beliebige Stammfunktion von f auch als **unbestimmtes Integral** und schreibt kurz

$$\tilde{F}(x) = \int dx f(x) + C, \quad C = \text{const.} \quad (1.9.26)$$

1.9.4 Definition des natürlichen Logarithmus' und der Exponentialfunktion

Mit dem Riemann-Integral lässt sich nun auch der **natürliche Logarithmus** sehr einfach definieren, und zwar durch

$$\ln x = \int_1^x dx' \frac{1}{x'}. \quad (1.9.27)$$

Der Definitionsbereich ist $x > 0$, denn der Integrand divergiert bei $x' = 0$. Da $1/x' > 0$ für $x' > 0$ ist, ist $\ln x$ im gesamten Definitionsbereich $x > 0$ streng monoton wachsend.

Aus dem Hauptsatz der Differential- und Integralrechnung folgt sofort, dass

$$\frac{d}{dx} \ln x = \frac{1}{x} \quad (1.9.28)$$

ist. Außerdem ist $\ln 1 = 0$, d.h. $\ln x < 0$ für $0 < x < 1$ und $\ln x > 0$ für $x > 1$.

1.9. Integralrechnung

Weiter gilt für $0 < a_1 < a_2$

$$\int_{a_1}^{a_2} dx \frac{1}{x} = \ln a_2 - \ln a_1. \quad (1.9.29)$$

Nun substituieren wir in dem Integral $y = x/a_1$, d.h. $dx = a_1 dy$. Damit wird

$$\ln a_2 - \ln a_1 = \int_1^{a_2/a_1} dy \frac{a_1}{a_1 y} = \int_1^{a_2/a_1} dy \frac{1}{y} = \ln(a_2/a_1), \quad (1.9.30)$$

Betrachten wir die Funktion $f(x) = \ln(ax)$ für $a > 0$ und $x > 0$. Nach der Kettenregel ist $f'(x) = (ax)' \ln'(ax) = a/(ax) = 1/x$. Nach dem Fundamentalsatz der Differential- und Integralrechnung ist also $f(x) = \ln x + C$. Nun ist $f(1) = C = \ln a$, d.h. es gilt für $a, x > 0$

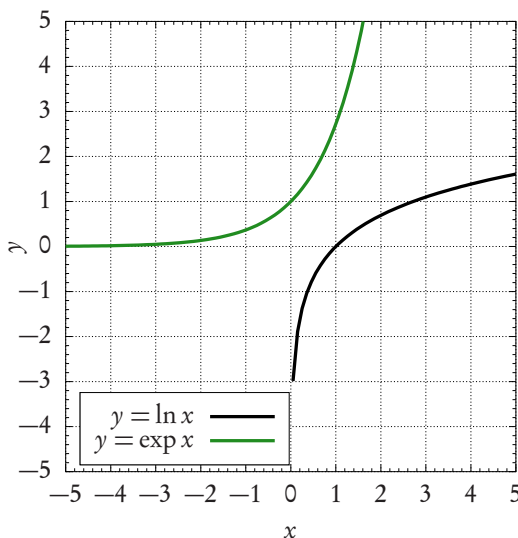
$$\ln(ax) = \ln a + \ln x. \quad (1.9.31)$$

Daraus folgt für $n \in \mathbb{N}$ und $x > 0$

$$\ln(x^n) = n \ln x, \quad (1.9.32)$$

wie man aus (1.9.31) sehr leicht durch vollständige Induktion beweisen kann (*Übung!*). Mit (1.9.30) folgt daraus wiederum

$$\ln(1/x^n) = \ln 1 - \ln x^n = -n \ln x. \quad (1.9.33)$$



Da $\ln x$ streng monoton wachsend ist, gibt es genau eine Zahl $e > 1$, so dass $\ln e = 1$ ist. Die Zahl $e \approx 2,781828$ heißt **Eulersche Zahl**. Daraus folgt

$$\ln(e^n) = n \ln e = n. \quad (1.9.34)$$

Für $n \rightarrow \infty$ bzw. $n \rightarrow -\infty$ ergibt sich daraus

$$\lim_{x \rightarrow \infty} \ln x = \infty, \quad \lim_{x \rightarrow 0^+} \ln x = -\infty, \quad (1.9.35)$$

da $e^n \rightarrow 0$ für $n \rightarrow -\infty$ ist. Der Wertebereich des Logarithmus' ist also ganz $\mathbb{R}$.

Da $\ln$ eine strikt monoton wachsende Funktion ist, ist deren Umkehrfunktion eindeutig auf ganz $\mathbb{R}$ definiert. Diese Funktion heißt **Exponentialfunktion** $\exp x$. Offensichtlich ist $\exp 0 = 1$ und $\exp x > 0$ für alle $x \in \mathbb{R}$.

Mit Hilfe der Regel von der Ableitung der Umkehrfunktion ergibt sich die Ableitung von $\exp x$ wie folgt: Wir setzen $y = \exp x$. Dann ist $x = \ln y$ und die Ableitung nach x lautet nach der Kettenregel

$$\frac{y'(x)}{y(x)} = 1 \Rightarrow y'(x) = \exp'(x) = y(x) = \exp x, \quad (1.9.36)$$

d.h. die Ableitung der Exponentialfunktion ist wieder die Exponentialfunktion selbst.

Betrachten wir nun die Funktion $f(x) = \exp(x+a)/\exp x$. Mit der Ketten- und Quotientenregel folgt für die Ableitung

$$f'(x) = \frac{\exp x \exp'(x+a) - \exp'(x) \exp(x+a)}{\exp^2 x} = 0 \Rightarrow f(x) = \text{const.} \quad (1.9.37)$$

Wegen $f(0) = \exp a / \exp 0 = \exp a$ ist $f(x) = \exp a = \text{const}$ und folglich

$$\exp(a+x) = \exp a \exp x \quad \text{für alle } x, a \in \mathbb{R}. \quad (1.9.38)$$

Setzen wir nun $f(x) = \exp x \exp(-x)$ ergibt sich mit der Produktregel $f'(x) = 0$ (*nachprüfen!*) und damit $f(x) = \text{const}$. Für $x = 0$ folgt $f(x) = f(0) = 1$ und damit

$$\exp x \exp(-x) = 1 \Rightarrow \exp(-x) = \frac{1}{\exp x} \quad \text{für alle } x \in \mathbb{R}. \quad (1.9.39)$$

Wie oben gezeigt, gilt $\ln(e^n) = n$ für $n \in \mathbb{Z}$. Draus folgt

$$\exp n = e^n \quad \text{für alle } n \in \mathbb{Z}. \quad (1.9.40)$$

Weiter gilt für $n \in \mathbb{Z}$ definitionsgemäß $(e^{1/n})^n = e$. Nimmt man davon den Logarithmus, folgt einerseits $\ln(e^{1/n})^n = \ln e = 1$ und andererseits $\ln(e^{1/n})^n = n \ln(e^{1/n})$. Folglich ist $\ln(e^{1/n}) = 1/n$. Für ein beliebiges $m \in \mathbb{Z}$ ergibt sich damit $\ln(e^{m/n}) = \ln[(e^{1/n})^m] = m \ln(e^{1/n}) = m/n$.

Mit der Exponentialfunktion folgt

$$\exp[\ln(e^{m/n})] = e^{m/n} = \exp(m/n). \quad (1.9.41)$$

Für alle $x \in \mathbb{Q}$ ist also $\exp x = e^x$. Es liegt nahe, entsprechend zu **definieren**, dass $e^x = \exp x$ für alle $x \in \mathbb{R}$ gilt.

Wir definieren weiter für jedes $a > 0$

$$a^x = \exp(x \ln a) = e^{x \ln a}. \quad (1.9.42)$$

Dann gilt ganz allgemein

$$\ln(a^x) = x \ln a. \quad (1.9.43)$$

Für $a \neq 1$ ist offenbar $f(x) = a^x$ eine echt monotone Funktion, und zwar echt monoton wachsend für $a > 1$ und echt monoton fallend für $0 < a < 1$ (*warum?*). Damit gibt es also auch eine Umkehrfunktion zu $f(x) = a^x$, die wir als Logarithmus zur Basis a bezeichnen. Definitionsgemäß ist also

$$\log_a(a^x) := x. \quad (1.9.44)$$

mit (1.9.43) folgt aber auch

$$\log_a(a^x) = x = \frac{\ln(a^x)}{\ln a}. \quad (1.9.45)$$

Setzen wir $y = a^x$, folgt für alle $y > 0$

$$\log_a y = \frac{\ln y}{\ln a}, \quad (1.9.46)$$

d.h. alle Logarithmen zu einer beliebigen Basis $a > 0$ mit $a \neq 1$ sind gemäß (1.9.46) durch den natürlichen Logarithmus berechenbar.

1.9.5 Hyperbelfunktionen

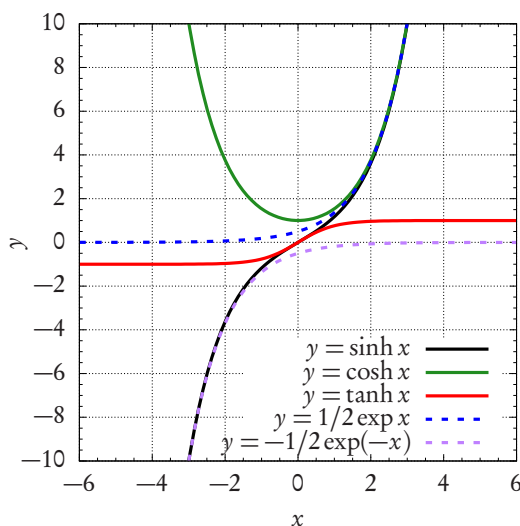
Nützlich sind oft auch noch die sog. **Hyperbelfunktionen**. Sie werden mit Hilfe der Exponentialfunktion durch

$$\begin{aligned} \sinh x &= \frac{\exp x - \exp(-x)}{2}, \\ \cosh x &= \frac{\exp x + \exp(-x)}{2}, \\ \tanh x &= \frac{\sinh x}{\cosh x} = \frac{\exp x - \exp(-x)}{\exp x + \exp(-x)} \end{aligned} \quad (1.9.47)$$

für $x \in \mathbb{R}$ definiert. Sie heißen sinus hyperbolicus, cosinus hyperbolicus bzw. tangens hyperbolicus. Durch direktes Nachrechnen (*Übung*) zeigt man, dass für alle $x \in \mathbb{R}$

$$\cosh^2 x - \sinh^2 x = 1 \quad (1.9.48)$$

gilt.



Wie aus dem nebenstehenden Plot zu ersehen und auch direkt aus (1.9.47) ersichtlich, gilt für große $|x|$

$$\begin{aligned} \sinh x &\cong \pm \frac{1}{2} \exp(\pm x) \quad \text{für } x \rightarrow \pm\infty, \\ \cosh x &\cong \frac{1}{2} \exp(\pm x) \quad \text{für } x \rightarrow \pm\infty, \\ \tanh x &\cong \pm 1 \quad \text{für } x \rightarrow \pm\infty. \end{aligned} \quad (1.9.49)$$

Aus der Ableitung der Exponentialfunktion $\exp' x = \exp x$ folgt sofort (*nachrechnen!*)

$$\sinh' x = \cosh x, \quad \cosh' x = \sinh x. \quad (1.9.50)$$

Für die Ableitung des $\tanh$ verwendet man wieder die Quotientenregel, was zu

$$\begin{aligned} \tanh' x &= \left(\frac{\sinh x}{\cosh x} \right)' = \frac{\cosh^2 x - \sinh^2 x}{\cosh^2 x} \\ &= 1 - \tanh^2 x = \frac{1}{\cosh^2 x}. \end{aligned} \quad (1.9.51)$$

führt.

Der $\sinh$ ist wegen $\sinh' x = \cosh x > 0$ streng monoton wachsend und daher überall eindeutig umkehrbar. Die entsprechende Umkehrfunktion $\operatorname{arsinh} : \mathbb{R} \rightarrow \mathbb{R}$ (area sinus hyperbolicus) ist demnach eindeutig durch $\sinh(\operatorname{arsinh} x) = x$ definiert.

Wir erhalten die Ableitung durch Anwenden der Ableitungsregel für Umkehrfunktionen. Setzen wir also $y = \operatorname{arsinh} x$. Dann ist $x = \sinh y$. Leiten wir diese Gleichung nach x ab folgt unter Verwendung der Kettenregel

$$\frac{dx}{dx} = 1 = y'(x) \cosh y = y'(x) \sqrt{1 + \sinh^2 y}. \quad (1.9.52)$$

Dabei haben wir im letzten Schritt (1.9.48) verwendet. Wegen $\cosh x > 0$ ist hierbei eindeutig die positive Wurzel zu verwenden. Nun ist aber $\sinh y = x$ und damit

$$y'(x) = \operatorname{arsinh}' x = \frac{1}{\sqrt{x^2 + 1}}. \quad (1.9.53)$$

Der $\cosh$ ist eine gerade Funktion und für $x > 0$ strikt monoton wachsend (für $x < 0$ strikt monoton fallend). Es gilt stets $\cosh x \geq 1$, denn die Ableitung $\cosh'(x) = \sinh x$ verschwindet für $x = 0$, und da $\cosh''(0) = \cosh(0) = 1 > 0$, besitzt die $\cosh$ -Funktion bei $x = 0$ ein Minimum.

Daher wird die Umkehrfunktion, der area cosinus hyperbolicus, eindeutig durch $\operatorname{arcosh} : \mathbb{R}_{\geq 1} \rightarrow \mathbb{R}_{\geq 0}$ via $\cosh(\operatorname{arcosh} x) = x$ definiert. Die Ableitung der Umkehrfunktion liefert (*nachrechnen*)

$$\operatorname{arcosh}' x = \frac{1}{\sqrt{x^2 - 1}}. \quad (1.9.54)$$

Der $\tanh$ ist wegen $\tanh' x = 1/\cosh^2 x > 0$ überall strikt monoton wachsend und folglich eindeutig umkehrbar. Der Wertebereich des $\tanh$ ist $(-1, 1)$, denn offensichtlich gilt $\lim_{x \rightarrow \pm\infty} \tanh x = \pm 1$. Damit ist durch $\operatorname{artanh} : (-1, 1) \rightarrow \mathbb{R}$ via $\tanh(\operatorname{artanh} x) = x$ der area tangens hyperbolicus eindeutig definiert, und mit (1.9.51) erhält man über die Ableitung der Umkehrfunktion (*nachrechnen!*)

$$\operatorname{artanh}' x = \frac{1}{1 - x^2}. \quad (1.9.55)$$

Die Kurve $\vec{r}(t) = [\cosh t, \sinh t]$, $t \in \mathbb{R}$ beschreibt den Ast einer **Hyperbel** (vgl. Abschn. 3.2), woher der Name für diese Funktionen stammt. Dies ist in Analogie zu den trigonometrischen Funktionen, für die ja $\vec{r}(t) = [\cos t, \sin t]$ für $t \in [0, 2\pi)$ einen Einheitskreis beschreibt. Die Erklärung, woher der Name „Area-Funktionen“ für die Umkehrfunktionen der Hyperbelfunktionen herrührt, findet sich in Abschnitt 3.10.2. Wir können die Area-Funktionen auch mit Hilfe des Logarithmus' ausdrücken. Betrachten wir als Beispiel arsinh . Setzen wir $y = \sinh x$, so ist definitionsgemäß $x = \operatorname{arsinh} y$. Wir können aber auch die Definition (1.9.47) für den $\sinh$ verwenden. Setzen wir dazu $z = \exp x$, folgt

$$y = \sinh x = \frac{\exp x - \exp(-x)}{2} = \frac{1}{2} \left(z - \frac{1}{z} \right). \quad (1.9.56)$$

Diese Gleichung lösen wir nun nach z auf. Multiplikation mit $2z$ und eine einfache Umformung liefert die quadratische Gleichung

$$z^2 - 2yz - 1 = 0 \Rightarrow (z - y)^2 - 1 - y^2 = 0 \Rightarrow (z - y) = \sqrt{y^2 + 1}. \quad (1.9.57)$$

Dabei haben wir beim Wurzelziehen verwendet, dass

$$z - y = \exp x - \sinh x = \exp x - \frac{1}{2}[\exp x - \exp(-x)] = \frac{1}{2}[\exp x + \exp(-x)] = \cosh x > 0 \quad (1.9.58)$$

ist. Damit wird aber eindeutig

$$z = \exp x = \sqrt{y^2 + 1} + y \Rightarrow x = \operatorname{arsinh} y = \ln(\sqrt{y^2 + 1} + y). \quad (1.9.59)$$

Es ist also

$$\operatorname{arsinh} x = \ln(\sqrt{x^2 + 1} + x) \quad \text{für } x \in \mathbb{R}. \quad (1.9.60)$$

Analog zeigt man auch, dass (*Übung!*)

$$\begin{aligned} \operatorname{arcosh} x &= \ln(\sqrt{x^2 - 1} + x) \quad \text{für } x \geq 1, \\ \operatorname{artanh} x &= \frac{1}{2} \ln\left(\frac{1+x}{1-x}\right) \quad \text{für } x \in (-1, 1). \end{aligned} \quad (1.9.61)$$

1.9.6 Integrationstechniken

Mit dem Hauptsatz der Differential- und Integralrechnung können wir nun sehr viele Integrale berechnen, indem wir eine beliebige Stammfunktion zu der zu integrierenden Funktion suchen. Dafür gibt es einige Rechentechniken, die wir hier kurz zusammenfassen. Zunächst kann man sich aus jeder Ableitungsregel eine Integrationsregel herleiten. Z.B. folgt aus der Ableitung von Potenzfunktionen $(x^n)' = nx^{n-1}$ die Regel

$$\int dx x^n = \frac{x^{n+1}}{n+1} + C. \quad (1.9.62)$$

Dies gilt für alle $n \in \mathbb{Z} \neq \{-1\}$.

Die Ausnahme ist $n = -1$, aber auch hier kennen wir eine Stammfunktion:

$$\int dx \frac{1}{x} = \ln x + C. \quad (1.9.63)$$

Für $n < 0$ müssen wir bei der Anwendung des Hauptsatzes aufpassen: Die entsprechenden Funktionen $f(x) = x^n$ besitzen nämlich dann bei $x = 0$ eine Singularität und sind insbesondere unstetig. Sie können daher *nur* über Intervalle $[a, b]$ integriert werden für die $0 \notin [a, b]$ ist, und nur dann liefert der Hauptsatz der Analysis den korrekten Wert für das Integral!

Weitere Grundintegrale sind

$$\begin{aligned}
 \int dx \sin x &= -\cos x + C, & \int dx \cos x &= \sin x + C, \\
 \int dx \frac{1}{\cos^2 x} &= \int dx(1 + \tan^2 x) = \tan x + C, & \int dx \frac{1}{\sin^2 x} &= \int dx(1 + \cot^2 x) = -\cot x + C, \\
 \int dx \exp x &= \exp x + C, & \int dx \sinh x &= \cosh x + C, & \int dx \cosh x &= \sinh x + C, \\
 \int dx \frac{1}{\cosh^2 x} &= \int dx(1 - \tanh^2 x) = \tanh x + C, \\
 \int dx \frac{1}{\sinh^2 x} &= \int dx(\coth^2 x - 1) = -\coth x + C. & (1.9.64) \\
 \int dx \frac{1}{\sqrt{1-x^2}} &= \arcsin x + C = -\arccos x + C', & \int dx \frac{1}{1+x^2} &= \arctan x + C = -\operatorname{arccot} x + C', \\
 \int dx \frac{1}{1-x^2} &= \operatorname{artanh} x + C \quad \text{falls } x < 1, & \int dx \frac{1}{1-x^2} \operatorname{arcoth} x + C & \quad \text{falls } x > 1, \\
 \int dx \frac{1}{\sqrt{1+x^2}} &= \operatorname{arsinh} x + C, & \int dx \frac{1}{\sqrt{x^2-1}} &= \operatorname{arcosh} x + C \quad \text{falls } x > 1.
 \end{aligned}$$

Aus allgemeinen Regeln der Differentialrechnung erhalten wir entsprechend Rechenregeln für unbestimmte Integrale.

Aus der Produktregel der Integralrechnung folgt die **Regel von der partiellen Integration**. Aus $(uv)' = u'v + v'u$ folgt

$$\int dx u'(x)v(x) = u(x)v(x) - \int dx u(x)v'(x) + C. \quad (1.9.65)$$

Als Anwendungsbeispiel betrachten wir das Integral von $x \exp x$. Setzen wir $u'(x) = \exp x$ und $v(x) = x$, gilt $u(x) = \exp x$ und $v'(x) = 1$. Damit ist

$$\int dx x \exp x = x \exp x - \int dx \exp x = (x-1)\exp x + C. \quad (1.9.66)$$

In der Tat findet man durch Ableiten, dass tatsächlich $[(x-1)\exp x]' = \exp x + (x-1)\exp x = x \exp x$ gilt. Ein weiteres Beispiel ist $\int dx \ln x$. Hier setzen wir $u'(x) = 1$ und $v(x) = \ln x$. Dann ist $u(x) = x$ und $v'(x) = 1/x$. Damit liefert (1.9.65)

$$\int dx \ln x = x \ln x - \int dx 1 = x(\ln x - 1) + C. \quad (1.9.67)$$

In der Tat folgt aus der Produktregel $[x(\ln x - 1)]' = (\ln x - 1) + x/x = \ln x$.

Die Kettenregel der Differentialrechnung liefert die **Substitutionsregel für Integrale**. Sei dazu F eine Stammfunktion von f .

$$\frac{d}{du} F[x(u)] = \frac{d}{dx} F[x(u)] \frac{dx(u)}{du} = f[x(u)] \frac{dx(u)}{du} \Rightarrow \int du \frac{dx(u)}{du} f[x(u)] = F[u(x)] + C. \quad (1.9.68)$$

1.9. Integralrechnung

Diese Regel lässt sich leichter merken, wenn man formal $dx = du \, dx/du$ schreibt. Dann wird (1.9.69) zu

$$\int du \frac{dx}{du} f[x(u)] = \left[\int dx f(x) \right]_{x=x(u)}. \quad (1.9.69)$$

Als Beispiel betrachten wir die Aufgabe, die Fläche eines Halbkreises um den Ursprung mit Radius r zu berechnen. Für den Halbkreis gilt $x^2 + y^2 = r^2$ bzw. $y(x) = \sqrt{r^2 - x^2}$ mit $x \in [-r, r]$. Damit ist die Fläche

$$A = \int_{-r}^r dx \sqrt{r^2 - x^2}. \quad (1.9.70)$$

Um das Integral zu berechnen, substituieren wir nun $x = r \cos u$. Das Intervall $[-r, r]$ wird dann umkehrbar eindeutig auf das Intervall $u \in [0, \pi]$ abgebildet. Weiter ist $dx = -du \, r \sin u$. Für $x = -r$ ist offenbar eindeutig $u = \pi$ und $x = r$ für $u = 0$. Damit folgt

$$A = \int_{\pi}^0 du (-r \sin u) \sqrt{r^2 - r^2 \cos^2 u} = r^2 \int_0^{\pi} du \sin^2 u. \quad (1.9.71)$$

Um schließlich dieses Integral zu berechnen, bemerken wir, dass

$$\cos(2u) = \cos^2 u - \sin^2 u = 1 - 2 \sin^2 u \Rightarrow \sin^2 u = \frac{1}{2} [1 - \cos(2u)]. \quad (1.9.72)$$

Dies in (1.9.71) eingesetzt liefert, wie aus der Geometrie bekannt,

$$A = \frac{r^2}{2} \int_0^{\pi} du [1 - \cos(2u)] = \frac{r^2}{2} \left[u - \frac{1}{2} \sin(2u) \right]_{u=0}^{u=\pi} = \frac{\pi r^2}{2}. \quad (1.9.73)$$

1.9.7 Funktionenfolgen und -reihen

In diesem Abschnitt untersuchen wir Funktionen, die sich als Grenzwerte von Folgen oder Reihen ergeben, deren Glieder selbst Funktionen sind. Insbesondere interessiert uns die Frage, ob Folgen oder Reihen stetiger Funktionen selbst gegen stetige Funktionen konvergieren und ob man ggf. Operationen wie die Differentiation und Integration mit der Grenzwertbildung vertauschen darf. Dabei ist es wichtig zu beachten, in welchem Sinne die Funktionenfolgen und -reihen konvergent sein sollen. Dazu definieren wir zwei Begriffe:

Es sei $(f_n(x))_n$ eine Folge von Funktionen $f_n : D \rightarrow \mathbb{R}$ mit einem beliebigen Definitionsbereich $D \subseteq \mathbb{R}$. Wir sagen, die Funktionenfolge konvergiere **punktweise** im Definitionsbereich D gegen eine Funktion $f : D \rightarrow \mathbb{R}$, wenn sie für jedes $x \in D$ konvergiert. D.h. zu jedem x existiert zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$, so dass $|f_n(x) - f(x)| < \epsilon$ für alle $n > N$.

Davon zu unterscheiden ist die stärkere Forderung nach **gleichmäßiger Konvergenz** gegen die Funktion f . Dazu verlangen wir, dass zu jedem ϵ ein $N \in \mathbb{N}$ existiert, so dass für alle $x \in D$ stets $|f_n(x) - f(x)| < \epsilon$ für alle $n > N$ gilt.

Im Unterschied zur punktweisen Konvergenz, muss also zu jedem ϵ das N unabhängig von $x \in D$ gewählt werden können, d.h. die Forderung nach gleichmäßiger Konvergenz ist stärker als die nach punktwiser.

Man kann die gleichmäßige Konvergenz auch noch anders beschreiben. Dies macht die Arbeit mit diesem Begriff erheblich einfacher.

Dazu führen wir einen Abstands begriff, eine **Norm**, für Funktionen ein, und zwar die **Supremumsnorm**. Sei also $f : D \rightarrow \mathbb{R}$ eine Funktion mit beliebigem Definitionsbereich $D \subseteq \mathbb{R}$. Dann ist die **Supremumsnorm** durch

$$\|f\| = \sup_{x \in D} |f(x)|. \quad (1.9.74)$$

Aus der Definition geht unmittelbar hervor, dass die folgenden Regeln für eine Norm gelten (*Beweis als Übung!*):

$$\|f\| \geq 0, \quad \|f\| = 0 \Leftrightarrow f(x) = 0 \quad \text{für alle } x \in D, \quad (1.9.75)$$

$$\|\lambda f\| = |\lambda| \|f\|, \quad (1.9.76)$$

$$\|f + g\| \leq \|f\| + \|g\|. \quad (1.9.77)$$

Eine Funktionenfolge ist offenbar genau dann gleichmäßig konvergent, wenn sie im Sinne dieser Norm konvergiert. Nehmen wir nämlich an, die Funktionenfolge $(f_n)_n$ konvergiert gleichmäßig gegen f . Dann gibt es voraussetzungsgemäß zu jedem $\epsilon \in \mathbb{R}$ ein $N \in \mathbb{N}$, so dass für alle x stets $|f_n(x) - f(x)| < \epsilon/2$ gilt, wenn nur $n > N$ ist. Dann ist aber auch $\|f_n - f\| = \sup_{x \in D} |f_n(x) - f(x)| \leq \epsilon/2 < \epsilon$. Kurz: Falls $(f_n)_n$ gleichmäßig konvergiert, gibt es zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$, so dass $\|f_n - f\| < \epsilon$ für alle $n > N$ gilt. Das bedeutet aber, dass die Funktionenfolge im Sinne der Supremumsnorm gegen f konvergiert.

Ist umgekehrt im Sinne die Funktionenfolge (f_n) im Sinne der Supremumsnorm konvergent gegen f , so gibt es zu jedem $\epsilon > 0$ ein $n \in \mathbb{N}$, so dass $\|f_n - f\| = \sup_{x \in D} |f_n(x) - f(x)| < \epsilon$ für alle $n > N$. Dann gilt aber für alle $x \in D$ offenbar $|f_n(x) - f(x)| < \epsilon$, d.h. die Funktionenfolge ist gleichmäßig konvergent.

Satz von der Stetigkeit stetiger Funktionenfolgen: Seien alle Glieder der $f_n : D \rightarrow \mathbb{R}$ stetig und gleichmäßig gegen f konvergente. Dann ist auch f stetig.

Beweis: Wegen der gleichmäßigen Konvergenz der Funktionenfolge $(f_n)_n$ gegen f , gibt es zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$, so dass für alle $x \in D$ für alle $n > N$ stets $|f_n(x) - f(x)| < \epsilon/3$ gilt. Sei nun $n > N$ festgehalten. Da weiter die f_n stetig in $x \in D$ sind, gibt es zu jedem $\epsilon > 0$ ein $\delta > 0$, so dass für alle $x' \in D$ mit $|x - x'| < \delta$ stets $|f_n(x) - f_n(x')| < \epsilon/3$ ist. Damit folgt nun aber, dass

$$|f(x) - f(x')| \leq |f(x) - f_n(x)| + |f_n(x) - f_n(x')| + |f_n(x') - f(x')|. \quad (1.9.78)$$

Wegen der Wahl von $n > N$ sind die beiden letzten Terme beide $< \epsilon/3$, während der mittlere für $|x - x'| < \delta$ ebenfalls $< \epsilon/3$ ist. Wir haben also zu einem beliebigen $\epsilon > 0$ ein $\delta > 0$ gefunden, so dass für alle $x' \in D$ mit $|x - x'| < \delta$ stets $|f(x) - f(x')| < \epsilon$ gilt. Damit ist f definitionsgemäß stetig, und das war zu zeigen.

Oft ist es wichtig zu wissen, ob man bestimmte Operationen an Funktionenfolgen mit der Limesbildung vertauschen darf. Hier interessieren uns Integration und Differentiation.

Integration von Funktionenfolgen: Seien $f_n : [a, b] \rightarrow \mathbb{R}$ stetige Funktionen und die Funktionenfolge $(f_n)_n$ auf $[a, b]$ gleichmäßig konvergent gegen eine Funktion $f : [a, b] \rightarrow \mathbb{R}$. Dann gilt

$$\int_a^b dx f(x) = \int_a^b dx \lim_{n \rightarrow \infty} f_n(x) = \lim_{n \rightarrow \infty} \int_a^b dx f_n(x). \quad (1.9.79)$$

Kurz gesagt darf man bei gleichmäßiger Konvergenz von Funktionenfolgen die Grenzwertbildung und die Integration vertauschen.

Beweis: Nach dem vorigen Satz ist f stetig auf $[a, b]$ und folglich Riemann-integrierbar. Wegen der gleichmäßigen Konvergenz der Funktionenfolge $(f_n)_n$ gegen f gilt

$$\left| \int_a^b dx f(x) - \int_a^b dx f_n(x) \right| \leq \int_a^b dx |f(x) - f_n(x)| \leq (b-a) \|f - f_n\| \xrightarrow{n \rightarrow \infty} 0, \quad (1.9.80)$$

d.h. der Limes der Integrale existiert. Folglich gilt (1.9.79) und das war zu zeigen.

Differentiation von Funktionenfolgen: Seien $f_n : [a, b] \rightarrow \mathbb{R}$ stetig differenzierbare Funktionen und die Funktionenfolge $(f_n)_n$ punktweise gegen eine Funktion $f : [a, b] \rightarrow \mathbb{R}$ konvergent. Die Folge $(f'_n)_n$ der Ableitungen konvergiere gleichmäßig. Dann ist f differenzierbar, und es gilt

$$f'(x) = \lim_{n \rightarrow \infty} f'_n(x). \quad (1.9.81)$$

MaW. darf man unter den angegebenen Voraussetzung Grenzwertbildung und Differentiation vertauschen.

Beweis: Es sei $\tilde{f}(x) = \lim_{n \rightarrow \infty} f'_n(x)$ für $x \in [a, b]$. Da die f'_n allesamt stetig sind und die Folge der Ableitungen voraussetzungsgemäß gleichmäßig konvergiert, ist $\tilde{f}$ nach dem oben bewiesenen Satz stetig und besitzt eine Stammfunktion, denn nach dem Satz über die Vertauschbarkeit der Integration mit der Grenzwertbildung für gleichmäßig konvergente Folgen stetiger Funktionen gilt

$$\int_a^x dx' \tilde{f}(x') = \lim_{n \rightarrow \infty} \int_a^x dx' f'_n(x) = \lim_{n \rightarrow \infty} [f_n(x) - f_n(a)]. \quad (1.9.82)$$

Wegen der vorausgesetzten punktwisen Konvergenz von $(f_n)_n$ gegen f gilt also

$$\int_a^x dx' \tilde{f}(x') = f(x) - f(a). \quad (1.9.83)$$

Ableiten dieser Gleichung nach x liefert nach dem Hauptsatz der Differential- und Integralrechnung schließlich $f'(x) = \tilde{f}(x)$, und das war zu zeigen.

Schließlich definieren wir eine **Funktionenreihe** $\sum_{k=1}^{\infty} f_k(x)$ als gleichmäßig konvergent gegen f , wenn ihre **Partialsummenfolge** $(S_n(x))_n$,

$$S_n(x) = \sum_{k=1}^n f_k(x), \quad (1.9.84)$$

gleichmäßig (also im Sinne der Supremumsnorm) gegen f konvergiert.

Weierstraßsches Konvergenzkriterium: Seien $f_n : D \rightarrow \mathbb{R}$ Funktionen und

$$\sum_{k=1}^{\infty} \|f_k\| \quad \text{konvergent}, \quad (1.9.85)$$

so ist die Funktionenreihe $\sum f_n$ absolut und gleichmäßig konvergent.

Beweis: Zunächst zeigen wir, dass die Reihe punktweise absolut konvergiert. Sei also $x \in D$. Dann gilt $|f_k(x)| \leq \|f_k\|$. Da voraussetzungsgemäß (1.9.85) gilt, ist also $\sum_k |f_k(x)|$ konvergent. Nun definieren wir im Sinne dieser *punktweisen* Konvergenz

$$f(x) = \sum_{k=1}^{\infty} f_k(x) \quad (1.9.86)$$

und zeigen, dass die Funktionenreihe sogar gleichmäßig gegen f konvergiert. Seien dazu die Partialsummen durch (1.9.84) definiert. Dann gilt

$$|S_n(x) - f(x)| = \left| \sum_{k=n+1}^{\infty} f_k(x) \right| \leq \sum_{k=n+1}^{\infty} |f_k(x)| \leq \sum_{k=n+1}^{\infty} \|f_k\|. \quad (1.9.87)$$

Wegen (1.9.85) können wir nun zu jedem $\epsilon > 0$ ein $N \in \mathbb{N}$ finden, so dass die letztere Summe $< \epsilon$ für alle $n > N$ ist. Damit gilt für alle $x \in D$

$$|S_n(x) - f(x)| < \epsilon \quad \text{für } n > N, \quad (1.9.88)$$

und folglich ist die Funktionenreihe tatsächlich nicht nur punktweise sondern sogar gleichmäßig gegen f konvergent, und das war zu zeigen.

1.9.8 Taylor-Entwicklung und Potenzreihen

Sei $f : D \rightarrow \mathbb{R}$ mit offenem Definitionsbereich $D \subseteq \mathbb{R}$ eine Funktion, die in einem abgeschlossenen Intervall $I \subset D$, mindestens $(n+1)$ -mal ($n \in \mathbb{N}_0$) stetig differenzierbar ist und sei $a \in D$. Wir suchen dann eine Näherung von f durch ein **Polynom** n -ten Grades für Argumente „in der Nähe“ von a . Sei $x \in I$. Dann gibt es Koeffizienten $c_k \in \mathbb{R}$, so dass

$$f(x) = \sum_{k=0}^n c_k (x-a)^k + R_n(x, a) \quad (1.9.89)$$

ist, wobei das Restglied $R_n(x, a) = \mathcal{O}[(x-a)^{n+1}]$ ist, d.h.

$$\lim_{n \rightarrow \infty} \frac{R_n(x, a)}{(x-a)^{n+1}} = \text{const.} \quad (1.9.90)$$

Um die Koeffizienten c_k in (1.9.89) zu bestimmen, leiten wir diese Gleichung k -mal ($k \in \{0, 1, \dots, n\}$) ab und setzen danach $x = a$:

$$\left. \frac{d^k}{dx^k} f(x) \right|_{x=a} =: f^{(k)}(a) = k! c_k, \quad (1.9.91)$$

denn wegen (1.9.90) verschwindet die k -te Ableitung des Restglieds für $x = a$.

Damit wird

$$f(x) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (x-a)^k + R_n(x, a). \quad (1.9.92)$$

Dies ist die **Taylorsche Formel**.

Zum **Beweis** bemerken wir, dass wegen der Stetigkeit von $f'(x)$

$$f(x) = f(a) + \int_a^x dx' f'(x') \quad (1.9.93)$$

gilt. Ist $n \geq 2$, können wir eine partielle Integration in (1.9.93) vornehmen. Dazu setzen wir

$$u'(x') = 1, \quad v(x') = f'(x'), \quad u(x') = x' - x,$$

so dass

$$f(x) = f(a) + (x-a)f'(a) - \int_a^x dx' (x' - x)f''(x') \quad (1.9.94)$$

resultiert. Dies wiederholen wir noch weitere $n - 1$ -mal. Dann entsteht

$$f(x) = \sum_{k=0}^{\infty} \frac{f^{(k)}(a)}{k!} (x-a)^k + \underbrace{(-1)^n \int_a^x \frac{(x'-x)^n}{n!} f^{(n+1)}(x') dx'}_{R_n(x,a)}. \quad (1.9.95)$$

Da voraussetzungsgemäß $f^{(n+1)}$ im Intervall (a, x) bzw. (x, a) stetig ist und die Funktion $x' \mapsto (x' - x)^n$ dort entweder stets negativ oder positiv ist, gibt es nach dem Mittelwertsatz der Integralrechnung wenigstens ein $\xi \in (a, x)$ bzw. $\xi \in (x, a)$, so dass

$$R_n(x, a) = (-1)^n \frac{f^{(n+1)}(\xi)}{n!} \int_a^x dx' (x' - x)^n = \frac{f^{(n+1)}(\xi)}{(n+1)!} (x-a)^{n+1}. \quad (1.9.96)$$

Dies ist die **Restgliedformel von Lagrange**.

Man kann freilich den Mittelwertsatz der Integralrechnung auch direkt auf die Funktion $x' \mapsto (x' - x)^n f^{(n+1)}(x')$ anwenden.

Dann resultiert die **Restgliedformel von Cauchy**:

$$R_n(x, a) = (-1)^n \frac{f^{(n+1)}(\xi')}{n!} (\xi' - x)^n \int_a^x dx' = \frac{f^{(n+1)}(\xi')}{n!} (x - \xi')^n (x - a). \quad (1.9.97)$$

Dabei ist $\xi' \in (a, x)$ bzw. $\xi' \in (x, a)$. Aus der Lagrangeschen Restgliedformel (1.9.96) folgt in der Tat (1.9.90), womit die Taylorsche Formel bewiesen ist.

Falls nun f im Intervall (a, x) bzw. (x, a) sogar beliebig oft stetig differenzierbar ist und ein $S > 0$ existiert, so dass $|f^{(n+1)}(\xi)| < S$ für alle $\xi \in (a, x)$ bzw. $\xi \in (x, a)$ ist, folgt aus (1.9.96)³

$$|R_n(x, a)| < S \frac{|x-a|^{n+1}}{(n+1)!} \rightarrow 0 \quad \text{für } n \rightarrow \infty. \quad (1.9.98)$$

In diesem Fall besitzt f die **Taylor-Entwicklung**

$$f(x) = \sum_{k=0}^{\infty} \frac{f^{(k)}(a)}{k!} (x-a)^k. \quad (1.9.99)$$

Als **Beispiel** betrachten wir die Taylor-Entwicklung von $x \mapsto \exp(x)$ um $a = 0$. Da

$$f^{(k+1)}(x) = \exp(x), \quad (1.9.100)$$

diese Funktion die Bedingungen für die Taylor-Entwicklung für jedes $x \in \mathbb{R}$ erfüllt, gilt

$$\exp(x) = \sum_{k=0}^{\infty} \frac{x^k}{k!}. \quad (1.9.101)$$

³Um zu zeigen, dass für beliebige $y > 0$ $\lim_{n \rightarrow \infty} y^n/n! = 0$ gilt, bemerken wir, dass für $n > 2y$ und $N \in \mathbb{N}$ die Abschätzung $0 \leq y^{n+N}/(n+N)! \leq y^n/n!(1/2)^N$ gilt. Der letzte Ausdruck strebt aber für $N \rightarrow \infty$ gegen 0.

Setzen wir in (1.9.99) $x = a + y$, können wir symbolisch

$$f(a + y) = \sum_{k=0}^{\infty} \frac{f^{(k)}(a)}{k!} y^k = \exp\left(y \frac{d}{da}\right) f(a) \quad (1.9.102)$$

schreiben. Dabei ist die Exponentialfunktion des **Differentialoperators** durch die formale Reihe (1.9.101)

$$\exp\left(y \frac{d}{da}\right) = \sum_{k=0}^{\infty} \frac{y^k}{k!} \frac{d^k}{da^k} \quad (1.9.103)$$

definiert. Wendet man diesen Operator auf f an, ergibt sich in der Tat die Taylor-Reihe von f in der in (1.9.102) angegebenen Form.

Für die Exponentialfunktion selbst gilt $f^{(k)}(a) = \exp a$ für alle $k \in \mathbb{N}$. Daraus folgt mit (1.9.102)

$$\exp(a + y) = \exp a \sum_{k=0}^{\infty} \frac{y^k}{k!} = \exp a \exp y. \quad (1.9.104)$$

Wichtig sind noch die Taylor-Entwicklungen der trigonometrischen Funktionen. Betrachten wir zuerst den Sinus:

$$f(x) = \sin x, \quad f'(x) = \cos x, \quad f''(x) = -\sin x, \quad \dots \quad (1.9.105)$$

Es gilt also

$$f^{(2k)}(0) = 0, \quad f^{(2k+1)}(0) = (-1)^k, \quad k \in \{0, 1, 2, \dots\}. \quad (1.9.106)$$

Die Taylor-Reihe für den Sinus ist also

$$\sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \dots = \sum_{k=0}^{\infty} \frac{(-1)^k}{(2k+1)!} x^{2k+1}. \quad (1.9.107)$$

Ebenso finden wir für den Kosinus

$$f(x) = \cos x, \quad f'(x) = -\sin x, \quad f''(x) = -\cos x, \quad \dots, \quad (1.9.108)$$

und damit die Taylorreihe

$$\cos x = 1 - \frac{x^2}{2!} + \frac{x^4}{4!} - \frac{x^6}{6!} + \dots = \sum_{k=0}^{\infty} \frac{(-1)^k}{(2k)!} x^{2k}. \quad (1.9.109)$$

Es ist klar, dass wegen $-1 \leq \sin x \leq 1$ und $-1 \leq \cos x \leq 1$ aufgrund der Restgliedformel die Taylorreihen (1.9.107) und (1.9.109) sicher überall in $\mathbb{R}$ und tatsächlich gegen $\sin x$ bzw. $\cos x$ konvergieren.

Die Additionstheoreme der trigonometrischen Funktionen können ebenfalls leicht mit Hilfe von (1.9.102) bewiesen werden. Für den $\cos$ bemerken wir, dass

$$f^{(k)}(a) = \begin{cases} (-1)^{j+1} \sin a & \text{falls } k = 2j + 1, \\ (-1)^j \cos a & \text{falls } k = 2j, \end{cases} \quad j \in \{0, 1, 2, \dots\}. \quad (1.9.110)$$

Daraus folgt mit (1.9.102)

$$\cos(a + y) = \cos a \sum_{j=0}^{\infty} (-1)^j \frac{y^{2j}}{(2j)!} + \sin a \sum_{j=0}^{\infty} (-1)^{j+1} \frac{y^{2j+1}}{(2j+1)!} = \cos a \cos y - \sin a \sin y, \quad (1.9.111)$$

wobei wir (1.9.107) und (1.9.109) benutzt haben. Außerdem haben wir stillschweigend die unendliche Reihe umgeordnet, indem wir gerade und ungerade Potenzen von y zusammengefasst haben. Dies ist erlaubt, weil Potenzreihen im Inneren ihres Definitionsbereichs absolut konvergent sind. Auf einen formalen Beweis verzichten wir hier.

Wir bemerken, dass wir offenbar die Taylor-Reihen von beliebig oft stetig differenzierbaren Funktionen gliedweise differenzieren und integrieren dürfen, d.h. die Summation der unendlichen Reihe vertauscht in diesem Fall mit der Differentiation und Integration.

Betrachten wir als ein weiteres Beispiel die geometrische Reihe

$$\sum_{k=0}^{\infty} x^k = \frac{1}{1-x} \quad \text{für} \quad -1 < x < 1. \quad (1.9.112)$$

Dies lässt sich unmittelbar aus (1.3.39) herleiten, denn demnach gilt

$$\sum_{k=0}^{\infty} x^k = 1 + \sum_{k=1}^{\infty} x^k = 1 + \frac{x}{1-x} = \frac{1}{1-x}. \quad (1.9.113)$$

Natürlich können wir (1.9.112) auch durch Taylorentwicklung der Funktion $f : (-1, 1)$, $f(x) = 1/(1-x)$ herleiten (*Übung!*). Das Argument der Konvergenz aufgrund der Taylorschen Restgliedformel gilt hier offensichtlich nur für abgeschlossene Intervalle $[a, b]$ mit $-1 < a < b < 1$, denn die Funktion und ihre Ableitungen sind offenbar bei $x = 1$ singulär und unbeschränkt für $x \rightarrow 1$.

Nun gilt (*nachrechnen!*)

$$\int_0^x dx' \frac{1}{1-x'} = -\ln(1-x) \quad \text{für} \quad x < 1. \quad (1.9.114)$$

Da die Potenzreihe (1.9.112) nur für $|x| < 1$ konvergiert, dürfen wir diese Integrationsformel nur für solche $|x|$ anwenden. Dann erhalten wir

$$\ln(1-x) = -\sum_{k=0}^{\infty} \frac{x^{k+1}}{k+1} = -\sum_{k=1}^{\infty} \frac{x^k}{k} \quad \text{für} \quad |x| < 1. \quad (1.9.115)$$

Dabei haben wir im letzten Schritt einfach $k+1$ durch k ersetzt und die Summationsgrenzen entsprechend angepasst. Ersetzen wir in (1.9.115) x durch $-x$ erhalten wir

$$\ln(1+x) = \ln[1-(-x)] = \sum_{k=1}^{\infty} \frac{(-1)^{k+1}}{k} x^k \quad \text{für} \quad |x| < 1. \quad (1.9.116)$$

Dies können wir auch wieder durch Taylorentwicklung der Funktion $\ln(1+x)$ beweisen (*Übung*).

Wir bemerken nun, dass wir auch Funktionen durch Potenzreihen definieren können. Sei zunächst durch eine formale Potenzreihe

$$f(x) = \sum_{k=0}^{\infty} c_k x^k \quad (1.9.117)$$

eine Funktion f definiert. Wir studieren nun das Konvergenzverhalten solcher Potenzreihen genauer. Wir interessieren uns für die absolute Konvergenz.

Satz: Sei die Reihe (1.9.121) für ein $x = r > 0$ konvergent. Dann ist sie auf $[-r, r]$ gleichmäßig und absolut konvergent.

Beweis: Wir definieren $f_k : [-r, r] \rightarrow \mathbb{R}$, $f_k(x) = c_k x^k$. Nun gilt

$$\|f_k\| = \sup_{x \in [-r, r]} |c_k x^k| \leq |c_k| r^k. \quad (1.9.118)$$

Nach dem Majorantenkriterium ist demnach die Reihe $\sum_k \|f_k\|$ konvergent und nach dem Weierstraßschen Konvergenzkriterium folglich (1.9.117) auf $[-r, r]$ absolut und gleichmäßig konvergent, und das war zu zeigen.

Definieren wir nun die Menge

$$K = \{x \in \mathbb{R} \mid (1.9.117) \text{ ist absolut konvergent}\}. \quad (1.9.119)$$

Ist dann K nach oben beschränkt, existiert nach dem Satz vom Supremum

$$R = \sup K. \quad (1.9.120)$$

Nach dem gerade bewiesenen Satz ist dann (1.9.117) im *offenen* Intervall $(-R, R)$ absolut konvergent und in jedem ganz in diesem Intervall gelegenen *abgeschlossenen* Intervall absolut und gleichmäßig konvergent. Man nennt R daher den **Konvergenzradius** der Potenzreihe.

Falls K nicht nach oben beschränkt ist, ist die Reihe für alle $x \in \mathbb{R}$ absolut konvergent, und in jedem endlichen Intervall gleichmäßig und absolut konvergent.

Zur praktischen Berechnung von R kann man sehr oft das Quotientenkriterium (1.3.44) heranziehen. Demnach ist die Reihe konvergent falls

$$\lim_{k \rightarrow \infty} \left| \frac{c_{k+1} x^{k+1}}{c_k x^k} \right| = |x| \lim_{k \rightarrow \infty} \frac{c_{k+1}}{c_k} < 1 \quad (1.9.121)$$

ist. Ist der Limes auf der rechten Seite 0, bedeutet dies, dass die Potenzreihe für alle $x \in \mathbb{R}$ konvergiert. Andernfalls ist die Reihe für

$$|x| < R \quad \text{mit} \quad R = \lim_{k \rightarrow \infty} \left| \frac{c_k}{c_{k+1}} \right| \quad (1.9.122)$$

konvergent. Falls der Grenzwert nicht existiert, muss die Konvergenz der Potenzreihe (1.9.116) gesondert untersucht werden. Da dies in der Praxis selten vorkommt, gehen wir nicht genauer darauf ein⁴.

⁴Haben wir es mit Potenzreihen zu tun, bei denen nur gerade oder ungerade Potenzen vorkommen, wie bei den Potenzreihen für $\sin$ und $\cos$ (1.9.107) und (1.9.109), können wir das Quotientenkriterium auf c_j/c_{j+2} anwenden. Dann ist offenbar $\lim_{j \rightarrow \infty} |c_j/c_{j+1}| = R^2$, falls der Limes existiert und demnach die Reihe für $x^2 < R^2$ konvergent ist, d.h. für $|x| < R$ konvergent.

Satz: Sei der Konvergenzradius der Potenzreihe (1.9.117) $R > 0$. Dann ist auch der Konvergenzradius der Potenzreihen

$$\tilde{f}(x) = \sum_{k=1}^{\infty} k c_k x^{k-1}, \quad (1.9.123)$$

$$\tilde{F}(x) = \sum_{k=0}^{\infty} \frac{c_k}{k+1} x^{k+1} \quad (1.9.124)$$

wieder R . Außerdem ist f differenzierbar und integrierbar, und es gilt

$$f'(x) = \tilde{f}(x), \quad (1.9.125)$$

$$F(x) = \int_0^x dx' f(x') = \tilde{F}(x). \quad (1.9.126)$$

Kurz gesagt dürfen Potenzreihen im Inneren ihres Konvergenzbereichs gliedweise differenziert und integriert werden.

Beweis: Sei $r < R$. Nach den obigen Betrachtungen zur gleichmäßigen und absoluten Konvergenz von Potenzreihen ist (1.9.117) auf dem abgeschlossenen Intervall $[-r, r]$ gleichmäßig konvergent. Da die Reihenglieder $f_k(x) = c_k x^k$ allesamt stetig sind, darf nach dem entsprechenden Satz über gleichmäßig konvergente Reihen Integration und Summation vertauscht werden. Damit ist bereits die Behauptung bzgl. der Integration bewiesen.

Hinsichtlich der Ableitung müssen wir nur zeigen, dass der Konvergenzradius der Reihe (1.9.123) der gleiche wie der der ursprünglichen Reihe (1.9.117) ist, denn dann ist (1.9.123) auf jedem kompakten Intervall innerhalb des Konvergenzbereichs gleichmäßig konvergent, und nach dem entsprechenden Satz zur Ableitung gleichmäßig konvergenter Funktionenfolgen gilt (1.9.125). Seien nun r und r' Zahlen, für die $0 < r < r' < R$ gilt. Dann ist $\sum_k c_k r'^k$ absolut konvergent, und es existiert $\max_k |c_k r'^k| = M > 0$. Sei nun $q = r/r' < 1$. Für alle $x \in [-r, r]$ gilt dann

$$|k c_k x^{k-1}| \leq k |c_k| r^{k-1} \leq k |c_k| r'^{k-1} k q^{k-1} \leq M k q^{k-1}. \quad (1.9.127)$$

Die Reihe $\sum_k k q^{k-1}$ ist nach dem Quotientenkriterium konvergent, denn mit $a_k k q^{k-1}$ gilt

$$\lim_{k \rightarrow \infty} \frac{a_{k+1}}{a_k} = \lim_{k \rightarrow \infty} \frac{k+1}{k} q = q < 1. \quad (1.9.128)$$

Mit $g_k(x) = k c_k x^{k-1}$ ist also

$$\|g_k\| \leq M k q^{k-1} \quad (1.9.129)$$

und folglich (1.9.123) tatsächlich auf $[r, r]$ gleichmäßig und absolut konvergent, und damit ist die Behauptung damit bewiesen.

1.10 Die strikte Definition der trigonometrischen Funktionen

Bisher haben wir die trigonometrischen Funktionen Sinus, Kosinus und Tangens nur geometrisch definiert, und bei der Herleitung von Rechenregeln wie Additionstheoremen und Ableitungen geometrische Betrachtungen verwendet, die nicht streng bewiesen wurden. Wir geben daher eine Definition der Funktionen $\cos$ und $\sin$ durch ihre Potenzreihen und leiten alle Eigenschaften dieser Funktionen mit Hilfe der oben bewiesenen Sätze über Potenzreihen her.

Wir definieren also die Funktionen durch die Potenzreihen (1.9.107) und (1.9.109), die sich oben unter Verwendung der heuristisch begründeten Ableitungsregeln ergeben haben:

$$\cos x = \sum_{k=0}^{\infty} \frac{(-1)^k x^{(2k)}}{(2k)!}, \quad (1.10.1)$$

$$\sin x = \sum_{k=0}^{\infty} \frac{(-1)^k x^{2k+1}}{(2k+1)!}. \quad (1.10.2)$$

Den Bereich absoluter Konvergenz können wir mit dem Quotientenkriterium ermitteln. Sei für die Kosinus-Reihe $a_k = (-1)^k x^{2k} / (2k)!$ dann gilt

$$\lim_{k \rightarrow \infty} \left| \frac{a_{k+1}}{a_k} \right| = \lim_{k \rightarrow \infty} \frac{|x|^{2k+2} (2k)!}{|x|^{2k} (2k+2)!} = x^2 \lim_{k \rightarrow \infty} \frac{1}{(2k+1)(2k+2)} = 0. \quad (1.10.3)$$

Das bedeutet, dass aufgrund der im vorigen Abschnitt bewiesenen Sätze die Potenzreihe für alle $x \in \mathbb{R}$ absolut und somit in jedem beschränkten Intervall gleichmäßig konvergiert und folglich gliedweise differenziert und integriert werden darf. Auch die Sinusreihe ist überall konvergent, wie man auf analoge Weise wie eben bei der Kosinusreihe nachrechnet.

Daraus ergeben sich sogleich die Ableitungsregeln, die wir vorher auf naive Weise berechnet haben:

$$\cos' x = -\sin x, \quad \sin' x = \cos x. \quad (1.10.4)$$

Als nächstes wollen wir einige Eigenschaften dieser Funktionen herleiten. Als erstes definieren wir $\pi/2$ als Nullstelle des Kosinus im Intervall $[0, 2]$. Dazu zeigen wir, dass der Kosinus dort genau eine Nullstelle besitzt. Aus der Potenzreihe folgt sofort, dass $\cos 0 = 1$ ist. Wir zeigen nun, dass $\cos 2 < 0$ ist. Es sei

$$a_k = \frac{2^{2k}}{(2k)!} - \frac{2^{2k+2}}{(2k+2)!} = \frac{2^{2k}}{(2k)!} \left(1 - \frac{4}{(2k+1)(2k+2)} \right). \quad (1.10.5)$$

Die Klammer ist monoton wachsend als Funktion von k , und für $k = 1$ wird sie $2/3$, d.h. es gilt $a_k > 0$ für $k \geq 1$. Damit gilt

$$\cos 2 = 1 - a_1 - a_3 - \dots = 1/3 - \sum_{k=1}^{\infty} a_{2k+1} < 0. \quad (1.10.6)$$

Da der Kosinus aufgrund der Definition als Potenzreihe stetig ist, gibt es aufgrund des Zwischenwertsatzes wegen $\cos 0 = 1 > 0$ und $\cos 2 < 0$ wenigstens ein $\xi \in (0, 2)$ mit $\cos \xi = 0$.

Um zu zeigen, dass es nur genau eine Nullstelle in diesem Intervall gibt, zeigen wir, dass $\cos x$ dort streng monoton fallend ist. Dazu müssen wir nur zeigen, dass $\cos' x = -\sin x$ überall in $(0, 2)$ negativ ist. In der Tat ist für alle $x \in (0, 2)$ und alle $k \in \mathbb{N}_0$

$$b_k = \frac{x^{2k+1}}{(2k+1)!} - \frac{x^{2k+3}}{(2k+3)!} = \frac{x^{2k+1}}{(2k+1)!} \left(1 - \frac{x^2}{(2k+2)(2k+3)} \right) > 0. \quad (1.10.7)$$

Damit ist für $x \in (0, 2)$

$$\sin x = \sum_{k=0}^{\infty} b_{2k} > 0 \Rightarrow \cos' x = -\sin x < 0 \quad (1.10.8)$$

und folglich $\cos x$ in $(0, 2)$ strikt monoton fallend. Damit gibt es also in diesem Intervall genau eine Nullstelle, die wir mit $\pi/2$ bezeichnen.

1.10. Die strikte Definition der trigonometrischen Funktionen

Die Additionstheoreme lassen sich nun über den Taylorschen Satz herleiten. Z.B. ist

$$\cos(x_1 + x_2) = \sum_{k=0}^{\infty} \frac{x_2^k}{k!} \frac{d^k}{dx_1^k} \cos x_1. \quad (1.10.9)$$

Nun ist

$$\frac{d^{2j+1}}{dx_1^{2j+1}} \cos x_1 = (-1)^{j+1} \sin x_1 \quad \text{und} \quad \frac{d^{2j}}{dx_1^{2j}} \cos x_1 = (-1)^j \cos x_1. \quad (1.10.10)$$

Die Potenzreihe (1.10.9) ist für alle x_2 absolut konvergent, wie man wie oben bei der Kosinus-Reihe nachrechnet. Folglich kann man diese Reihe nach dem Umordnungssatz beliebig umordnen, ohne dass sich an ihrem Wert etwas ändert. Wir können also in (1.10.9) alle Glieder mit geraden und mit ungeraden k zusammenfassen. Mit (1.10.10) ergibt dies

$$\cos(x_1 + x_2) = \cos x_1 \sum_{j=0}^{\infty} \frac{(-1)^j x_2^{2j}}{(2j)!} - \sin x_1 \sum_{j=0}^{\infty} \frac{(-1)^j x_2^{2j+1}}{(2j+1)!} = \cos x_1 \cos x_2 - \sin x_1 \sin x_2. \quad (1.10.11)$$

Dabei haben wir uns im letzten Schritt der Definitionen (1.10.1) und (1.10.2) des Kosinus und Sinus bedient. Aus den Ableitungsregeln folgt daraus sofort

$$\sin(x_1 + x_2) = -\frac{d}{dx_1} \cos(x_1 + x_2) = \sin x_1 \cos x_2 + \cos x_1 \sin x_2. \quad (1.10.12)$$

Als nächstes beweisen wir die Formel

$$\cos^2 x + \sin^2 x = 1. \quad (1.10.13)$$

Dazu bilden wir die Ableitung mit Hilfe der Produktregel:

$$\frac{d}{dx} (\cos^2 x + \sin^2 x) = -2 \sin x \cos x + 2 \sin x \cos x = 0. \quad (1.10.14)$$

Damit ist klar, dass $\cos^2 x + \sin^2 x = \text{const}$ ist, und für $x = 0$ folgt, dass die Konstante 1 sein muss und damit (1.10.13).

Wegen $\cos(\pi/2) = 0$ ist also $\sin^2(\pi/2) = 1$, und da für $x \in (0, 2)$ stets $\sin x > 0$ ist, muss demnach $\sin \pi/2 = 1$ gelten. Daraus folgt dann mit den Additionstheoremen (1.10.11)

$$\cos(x + \pi/2) = \cos x \cos(\pi/2) - \sin x \sin(\pi/2) = -\sin x, \quad (1.10.15)$$

$$\sin(x + \pi/2) = \sin x \cos(\pi/2) + \cos x \sin(\pi/2) = \cos x. \quad (1.10.16)$$

Da $\sin x \geq 0$ für $x \in [0, \pi/2]$ ist $\cos x \leq 0$ für $x \in [\pi/2, \pi]$, und wegen $\cos x \geq 0$ für $x \in [0, \pi/2]$ ist $\sin x \geq 0$ für $x \in [\pi/2, \pi]$. Wegen der Ableitungsregeln für Sinus und Kosinus folgt daraus weiter, dass in $[\pi/2, \pi]$ der Sinus und Kosinus beide monoton fallend sind. Setzt man in (1.10.15) und (1.10.16) jeweils $x = \pi/2$ folgt

$$\cos \pi = -\sin(\pi/2) = -1, \quad \sin \pi = \cos(\pi/2) = 0. \quad (1.10.17)$$

Damit haben wir wieder mit den Additionstheoremen (1.10.11) und (1.10.12)

$$\cos(x + \pi) = -\cos x, \quad \sin(x + \pi) = -\sin x. \quad (1.10.18)$$

Folglich ist in $[\pi, 3\pi/2]$ stets $\cos x \leq 0$ und $\sin x \leq 0$ und der Kosinus monoton wachsend und der Sinus monoton fallend. Verwendet man diese Information in (1.10.15) und (1.10.16) folgt, dass im Intervall $[3\pi/2, 2\pi]$ stets $\cos x \geq 0$ und $\sin x \leq 0$ und Kosinus und Sinus monoton wachsend sind.

Für $x = \pi/2$ bzw. $x = \pi$ folgt aus (1.10.18) mit (1.10.17)

$$\cos(3\pi/2) = 0, \quad \sin(3\pi/2) = -1, \quad \cos(2\pi) = 1, \quad \sin(2\pi) = 0. \quad (1.10.19)$$

Daraus folgt, dass Kosinus und Sinus **periodische Funktionen** mit der Periode 2π sind, d.h.

$$\cos(x + 2\pi) = \cos x, \quad \sin(x + 2\pi) = \sin x. \quad (1.10.20)$$

Außerdem ist aufgrund der eben hergeleiteten Monotonieeigenschaften dieser Funktionen 2π die kleinste Periode beider Funktionen. Damit haben wir die wesentlichsten Eigenschaften von Kosinus und Sinus streng hergeleitet.

Dass diese Funktionen die bisher zu ihrer Definition verwendete geometrische Bedeutung besitzen, wird in Kapitel 3 bei der Behandlung der Parameterdarstellung von Kurven in der Ebene und im Raum gezeigt.

Kapitel 2

Lineare Algebra

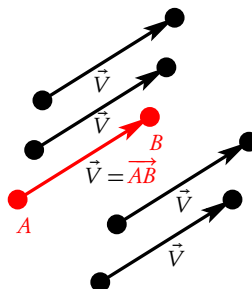
Die **Vektorrechnung** oder **lineare Algebra** und die **Vektoranalysis** stellen die wesentliche mathematische Grundlage zur Beschreibung physikalischer Systeme dar. Wir gehen von der Euklidischen Geometrie des physikalischen dreidimensionalen Raumes aus, die wir als bekannt voraussetzen. Wir werden schrittweise anhand des Begriffs der **Verschiebung** im physikalischen Anschauungsraum die Algebra von Vektoren (Addition und Multiplikation mit reellen Zahlen) motivieren und dann abstrahieren.

2.1 Geometrische Einführung von Euklidischen Vektoren

Wir gehen davon aus, dass die Grundlagen der **Euklidischen Geometrie** von der Schulmathematik her bekannt sind und entwickeln deren Formulierung als **analytische Geometrie** mit Hilfe von Vektoren. Dies ist für die gesamte moderne Naturwissenschaft, wie sie von Galilei und Newton im 17. Jh. begründet worden ist, von entscheidender Bedeutung, denn die analytische Geometrie macht die Analyse von **Bewegungen** von Körpern im Raum der Analysis, also der **Differential- und Integralrechnung**, zugänglich.

2.1.1 Definition von Vektoren als Verschiebungen

Wir führen den Begriff des **Vektors** anhand der Beschreibung von **Verschiebungen** ein. Seien also A und B zwei Punkte. Aufgrund der Axiome, die der Euklidischen Geometrie zugrunde liegen, können wir diese beiden Punkte durch eine **gerade Strecke** verbinden, die eine Länge besitzt. Wir markieren zugleich die Richtung mit einem Pfeil, legen also fest, dass wir die Verschiebung des Punktes A nach B entlang der geraden Strecke meinen.

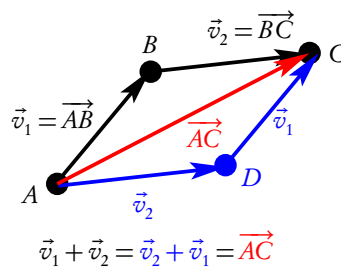


Dieses Objekt, das die Verschiebung von A nach B durch **Länge und Richtung** bzw. eine **gerichtete Strecke** festlegt, bezeichnen wir als den Vektor $\overrightarrow{AB}$.

Nun interessieren wir uns oft nicht für die konkreten beiden Punkte, die wir ineinander verschieben, sondern lediglich für Länge und Richtung der Verschiebung selbst. Daher identifizieren wir alle Pfeile, die aus $\overrightarrow{AB}$ durch **Parallelverschiebung** hervorgehen und bezeichnen diese Klasse von Pfeilen als den Vektor $\vec{v} = \overrightarrow{AB}$. Wir bezeichnen die Menge aller Vektoren in einer Ebene mit E^2 bzw. E^3 , wobei das E für „Euklidischer Vektorraum“ und der hochgestellte Index 2 bzw. 3 die Dimensionszahl angibt.

2.1.2 Vektoraddition

Geben wir zwei Vektoren $\vec{v}_1 = \overrightarrow{AB}$ und $\vec{v}_2 = \overrightarrow{BC}$ vor, so definieren wir die Hintereinanderausführung der Verschiebungen (s. nebenstehende Skizze), die direkt zur Verschiebung $\overrightarrow{AC}$ führt, als die **Summe der Vektoren**: $\vec{v}_1 + \vec{v}_2 = \overrightarrow{AC}$.



Durch Parallelverschiebung von $\vec{v}_2$, so dass sein Anfangspunkt in A zu liegen kommt, ergibt den Punkt D vermöge $\vec{v}_2 = \overrightarrow{AD}$. Nach den Gesetzen der Euklidischen Geometrie ist dann $\overrightarrow{DC} = \vec{v}_1$. Entsprechend folgt $\vec{v}_2 + \vec{v}_1 = \overrightarrow{AD} + \overrightarrow{DC} = \overrightarrow{AC} = \vec{v}_1 + \vec{v}_2$. Das bedeutet, dass die **Vektoraddition kommutativ** ist, d.h. die Summe zweier Vektoren hängt nicht von der Reihenfolge der Summation ab:

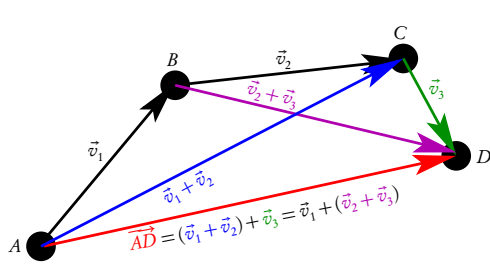
$$\vec{v}_1 + \vec{v}_2 = \vec{v}_2 + \vec{v}_1. \quad (2.1.1)$$

Nun führen wir noch den (nur scheinbar sinnlosen) **Nullvektor** $\overrightarrow{AA} = \vec{0}$ ein. Es ist klar, dass das im Sinne von Verschiebungen bedeutet, dass gar keine Verschiebung ausgeführt wird. Im Sinne unserer Äquivalenzklassenbildung gilt für jeden anderen Punkt B ebenfalls, dass $\overrightarrow{BB} = \vec{0}$ ist.

Entsprechend folgt für die Addition $\overrightarrow{AA} + \overrightarrow{AB} = \overrightarrow{AB}$ bzw. $\vec{0} + \vec{v}_1 = \vec{v}_1 + \vec{0} = \vec{v}_1$. Der Nullvektor ist also das **neutrale Element der Vektoraddition**.

Es ist auch klar, dass wir zu jedem Verschiebungsvektor $\vec{v} = \overrightarrow{AB}$ den die umgekehrte Verschiebung kennzeichnenden Vektor $(-\vec{v}) = \overrightarrow{BA}$ zuordnen können. Der Summe dieser beiden Vektoren entspricht gerade die Verschiebung von A nach B und dann wieder zurück zu A . Insgesamt haben wir also gar keine Verschiebung ausgeführt. Es ist also $\vec{v} + (-\vec{v}) = \overrightarrow{AB} + \overrightarrow{BA} = \overrightarrow{AA} = \vec{0}$. Es ist also $(-\vec{v})$ das bzgl. der Addition inverse Element von $\vec{v}$.

2.1. Geometrische Einführung von Euklidischen Vektoren



Betrachten wir nun drei Vektoren $\vec{v}_1 = \overrightarrow{AB}$, $\vec{v}_2 = \overrightarrow{BC}$ und $\vec{v}_3 = \overrightarrow{CD}$. Dann ist

$$\vec{v}_1 + \vec{v}_2 = \overrightarrow{AB} + \overrightarrow{BC} = \overrightarrow{AC}. \quad (2.1.2)$$

und folglich

$$(\vec{v}_1 + \vec{v}_2) + \vec{v}_3 = \overrightarrow{AC} + \overrightarrow{CD} = \overrightarrow{AD}. \quad (2.1.3)$$

Addieren wir jetzt diese drei Vektoren in einer etwas anderen Reihenfolge, und zwar bilden wir zuerst die Summe

$$\vec{v}_2 + \vec{v}_3 = \overrightarrow{BC} + \overrightarrow{CD} = \overrightarrow{BD}. \quad (2.1.4)$$

Dann folgt

$$\vec{v}_1 + (\vec{v}_2 + \vec{v}_3) = \overrightarrow{AB} + \overrightarrow{BD} = \overrightarrow{AD}. \quad (2.1.5)$$

Vergleichen wir dies mit (2.1.3), ergibt sich das **Assoziativgesetz der Vektoraddition**

$$(\vec{v}_1 + \vec{v}_2) + \vec{v}_3 = \vec{v}_1 + (\vec{v}_2 + \vec{v}_3). \quad (2.1.6)$$

Dies zeigt, dass wir hinsichtlich der Addition mit Vektoren formal genauso wie mit reellen Zahlen rechnen können. Insbesondere können wir auch Gleichungen lösen. Seien z.B. $\vec{a}$ und $\vec{b}$ vorgegebene Vektoren. Wir suchen nun einen Vektor $\vec{x}$, der die Gleichung $\vec{a} + \vec{x} = \vec{b}$ erfüllt. Hätten wir Zahlen vorliegen, könnten wir einfach $\vec{a}$ auf beiden Seiten der Gleichungen abziehen, um $\vec{x}$ zu finden. Aufgrund der eben hergeleiteten Rechenregeln funktioniert das auch für Vektoren, denn es gilt

$$\vec{b} + (-\vec{a}) = (\vec{a} + \vec{x}) + (-\vec{a}) = (-\vec{a}) + (\vec{a} + \vec{x}) = [(-\vec{a}) + \vec{a}] + \vec{x} = \vec{0} + \vec{x} = \vec{x}. \quad (2.1.7)$$

Es ist eine gute *Übung* sich zu vergewissern, welche der oben hergeleiteten Rechenregeln bei den einzelnen Umformungsschritten verwendet wurden! Entsprechend definieren wir die **Subtraktion von Vektoren** in der naheliegenden Weise als $\vec{b} - \vec{a} = \vec{b} + (-\vec{a})$.

2.1.3 Länge (Norm) von Vektoren

Bisher haben wir nicht von der Eigenschaft von Vektoren Gebrauch gemacht, dass sie auch eine **Länge** besitzen. Die Länge des Vektors $\vec{v} = \overrightarrow{AB}$ ist dabei natürlich einfach durch die Länge der Strecke $|\vec{v}| = |\overline{AB}|$ im Sinne der Euklidischen Geometrie definiert. Man nennt $|\vec{v}|$ auch die **Euklidische Norm** des Vektors $\vec{v}$ oder (vor allem in der Physik) auch einfach den **Betrag** oder die **Länge** des Vektors $\vec{v}$. Der Betrag ist eine positive reelle Zahl.

Dass für die Länge von Strecken *reelle* Zahlen benötigt werden und nicht etwa rationale Zahlen ausreichen, ist keinesfalls trivial. Erst David Hilbert (1862-1943) hat Ende des 19. Jh. bemerkt, dass die Manipulationen mit Lineal und Zirkel, wie sie Euklid im Altertum ausgeführt bzw. axiomatisch begründet hat, die reellen Zahlen erfordern, also auch irrationale Zahlen benötigt werden.

Die Euklidische Norm von Vektoren erbt nun naturgemäß einige Eigenschaften vom Längenbegriff der Euklidischen Geometrie. Z.B. ist die Länge des Nullvektors $\vec{0}$: $|\vec{0}| = 0$, denn ein Punkt besitzt definitionsgemäß keine Ausdehnung. Ist umgekehrt $\vec{v}$ ein Vektor mit $|\vec{v}| = 0$ ist offenbar $\vec{v} = \vec{0}$.

Weniger trivial ist die **Dreiecksungleichung**. Sind nämlich A, B und C drei beliebige nicht auf einer Gerade gelegene Punkte, dann gilt für die Seiten des von ihnen definierten Dreiecks stets $|AB| + |BC| > |AC|$. Seien also $\vec{v}_1 = \overrightarrow{AB}$, $\vec{v}_2 = \overrightarrow{BC}$, so gilt

$$|\overrightarrow{AC}| = |\vec{v}_1 + \vec{v}_2| < |\vec{v}_1| + |\vec{v}_2|. \quad (2.1.8)$$

Liegen die drei Punkte auf einer Geraden und B zwischen A und C , so gilt offenbar $|AB| + |BC| = |AC|$. In diesem Fall gilt also $|\vec{v}_1 + \vec{v}_2| = |\vec{v}_1| + |\vec{v}_2|$.

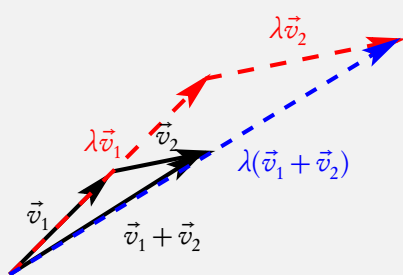
Es gilt also für alle Vektoren $\vec{v}_1$ und $\vec{v}_2$ immer die **Dreiecksungleichung**

$$|\vec{v}_1 + \vec{v}_2| \leq |\vec{v}_1| + |\vec{v}_2|. \quad (2.1.9)$$

Das Gleichheitszeichen gilt offenbar dann und nur dann, wenn die Vektoren $\vec{v}_1$ und $\vec{v}_2$ parallel zueinander sind.

Nun gibt es in der Euklidischen Geometrie zu zwei Punkten A und B und jeder reellen Zahl $\lambda > 0$ einen Punkt auf der durch A und B eindeutig festgelegten Geraden einen Punkt C , so dass $|AC| = \lambda|AB|$, wobei wir festlegen, dass für $\lambda < 1$ der Punkt C zwischen A und B und für $\lambda > 1$ der Punkt B zwischen A und C liegen soll. Entsprechend definieren wir die Multiplikation des Vektors $\vec{v} = \overrightarrow{AB}$ mit der reellen positiven Zahl λ durch $\lambda\vec{v} = \overrightarrow{AC}$. Anders ausgedrückt bedeutet die Verschiebung um den Vektor $\lambda\vec{v}$ eine Verschiebung in die gleiche Richtung wie die durch $\vec{v}$ vorgegebene Verschiebung, aber um eine um den Faktor λ verschiedene Länge.

Wir wollen eine solche Multiplikation von Vektoren mit reellen Zahlen auch für $\lambda < 0$ definieren. Wie wir gleich sehen werden, ist es sinnvoll, in diesem Fall $\lambda\vec{v} = -(|\lambda|\vec{v})$ zu setzen. Dies liegt nahe, denn für $\lambda < 0$ ist $\lambda = -|\lambda|$. Wir verschieben in diesem Fall also um eine um den Faktor λ geänderte Strecke in entgegengesetzter Richtung zu $\vec{v}$. Schließlich definieren wir noch, dass $0\vec{v} = \vec{0}$ sein soll. Man macht sich schnell klar, dass für zwei Zahlen $\lambda_1, \lambda_2 \in \mathbb{R}$ das **Assoziativgesetz** $\lambda_1(\lambda_2\vec{v}) = (\lambda_1\lambda_2)\vec{v}$ gilt.



Es ist unmittelbar einsichtig, dass $(\lambda_1 + \lambda_2)\vec{v} = \lambda_1\vec{v} + \lambda_2\vec{v}$. Es ergeben sich aus diesen Rechenregeln sofort die unmittelbar einleuchtende Formeln wie $\vec{v} + \vec{v} = 2\vec{v}$, d.h. führt man zweimal dieselbe Verschiebung hintereinander aus, erhält man eine Verschiebung in die gleiche Richtung aber um die doppelte Länge. Aus der nebenstehenden Skizze entnehmen wir, dass aufgrund des Strahlensatzes der Euklidischen Geometrie auch das **Distributivgesetz** $\lambda(\vec{v}_1 + \vec{v}_2) = \lambda\vec{v}_1 + \lambda\vec{v}_2$ gilt.

2.1.4 Lineare Unabhängigkeit von Vektoren und Basen

Seien $\vec{b}_1$ und $\vec{b}_2$ zwei nichtparallele Vektoren. Das bedeutet, dass es keine reelle Zahl λ gibt, für die $\lambda\vec{b}_1 = \vec{b}_2$ ist. Man nennt solche Vektoren **linear unabhängig**. Eine etwas allgemeinere Definition ist, dass zwei Vektoren linear unabhängig voneinander sind, genau dann wenn aus $\lambda_1\vec{b}_1 + \lambda_2\vec{b}_2 = \vec{0}$ folgt, dass notwendig $\lambda_1 = \lambda_2 = 0$ sein muss. Beide Definitionen sind offenbar äquivalent. Gilt nämlich $\lambda\vec{b}_1 = \vec{b}_2$, so ist $\lambda\vec{b}_1 - \vec{b}_2 = \vec{0}$. Es ist also zumindest $\lambda_2 = -1 \neq 0$, so dass die Vektoren linear abhängig sind. Ist umgekehrt $\lambda_1\vec{b}_1 + \lambda_2\vec{b}_2 = \vec{0}$ und $\lambda_2 \neq 0$, so gilt $\vec{b}_2 = -(\lambda_1/\lambda_2)\vec{b}_1$, d.h. die Vektoren sind auch gemäß der ersten Definition linear abhängig.

Dies lässt sich nun auf beliebig viele Vektoren verallgemeinern.

Wir nennen eine beliebige endliche Menge von Vektoren $\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n\}$ **voneinander linear unabhängig**, genau dann wenn aus

$$\lambda_1 \vec{v}_1 + \lambda_2 \vec{v}_2 + \dots + \lambda_n \vec{v}_n = \sum_{j=1}^n \lambda_j \vec{v}_j = 0 \quad (2.1.10)$$

notwendig $\lambda_1 = \lambda_2 = \dots = \lambda_n = 0$ folgt. Andernfalls heißen die Vektoren **voneinander linear abhängig**.

Betrachten wir nun als Beispiel Vektoren in einer Ebene. Seien $\vec{b}_1$ und $\vec{b}_2$ zwei beliebige voneinander linear unabhängige Vektoren. Dann können wir jeden beliebigen Vektor $\vec{v}$ durch **Linearkombination** aus diesen **Basisvektoren** zusammensetzen. Wir schreiben die entsprechenden Zahlen, die **Komponenten** von $\vec{v}$ als v_1 und v_2 , d.h. wir können stets Zahlen v_j ($j \in \{1, 2\}$) finden, so dass

$$\vec{v} = v_1 \vec{b}_1 + v_2 \vec{b}_2 = \sum_{j=1}^2 v_j \vec{b}_j. \quad (2.1.11)$$

Es ist nun auch klar, dass diese Zahlen eindeutig sind. Die Vektoren $\vec{b}_1$ und $\vec{b}_2$ sind nämlich linear unabhängig voneinander, denn sie sind nicht parallel zueinander, weisen also in verschiedene Richtungen in der Ebene. Seien nun λ_1 und λ_2 irgendwelche Komponenten von $\vec{v}$, folgt nämlich

$$\vec{0} = \vec{v} - \vec{v} = (v_1 - \lambda_1) \vec{b}_1 + (v_2 - \lambda_2) \vec{b}_2. \quad (2.1.12)$$

Da $\vec{b}_1$ und $\vec{b}_2$ linear unabhängig sind, folgt daraus notwendig, dass $v_1 - \lambda_1 = 0$, also $v_1 = \lambda_1$, und $v_2 - \lambda_2 = 0$, also $v_2 = \lambda_2$ sein muss.

Man nennt nun eine Menge von Vektoren $\{\vec{b}_1, \dots, \vec{b}_n\}$ **vollständig**, wenn man jeden Vektor $\vec{v}$ als Linearkombination dieser Vektoren darstellen kann. Falls diese Menge zusätzlich auch noch linear unabhängig ist, ist diese Linearkombination für jeden Vektor eindeutig, was man genauso beweist wie für unser Beispiel mit zwei Vektoren in der Ebene, und man nennt entsprechend jede vollständige Menge linear unabhängiger Vektoren **eine Basis** des Vektorraumes. Für E^2 bestehen alle Basen offensichtlich aus genau zwei Vektoren.

Genauso bestehen alle Basen im dreidimensionalen Raum offensichtlich aus beliebigen Mengen von genau drei linear unabhängigen Vektoren. Man nennt einen Vektorraum, der eine Basis aus endlich vielen Vektoren besitzt, einen **endlichdimensionalen Vektorraum**. Offenbar bilden die geometrischen Vektoren wie wir sie in diesem Abschnitt definiert haben, in einer Ebene einen zweidimensionalen bzw. im Raum einen dreidimensionalen Vektorraum.

Wir merken hier nur an, dass es Vektorräume beliebiger endlicher Dimension aber auch solche unendlicher Dimension gibt. In diesem Manuskript befassen wir uns nur mit endlichdimensionalen Vektorräumen, und zwar vornehmlich mit den in der Euklidischen Geometrie des physikalischen Raumes der Newtonschen Mechanik auftretenden zwei- und dreidimensionalen Vektorräumen. Beschränkt man sich auf Vektoren entlang einer Geraden, hat man es auch mit einem eindimensionalen Vektorraum zu tun.

2.1.5 Der Vektorraum $\mathbb{R}^3$

Wir haben im vorigen Abschnitt gesehen, dass wir durch Einführung einer Basis $\{\vec{b}_1, \vec{b}_2, \vec{b}_3\}$ jeden räumlichen Vektor $\vec{v}$ durch seine drei Komponenten v_1, v_2 und v_3 eindeutig als Linearkombination dieser Basisvektoren

$$\vec{v} = v_1 \vec{b}_1 + v_2 \vec{b}_2 + v_3 \vec{b}_3 = \sum_{j=1}^3 v_j \vec{b}_j \quad (2.1.13)$$

darstellen können.

Da solche Summenbildungen im folgenden ständig auftreten, lässt man oft auch die Summenzeichen einfach weg. Diese Konvention geht auf Einstein zurück, der sie bei der Formulierung der Allgemeinen Relativitätstheorie eingeführt hat. Man spricht daher auch von der **Einsteinschen Summationskonvention**.

Es ist klar, dass umgekehrt auch durch beliebige drei Zahlen (v_j) ($j \in \{1, 2, 3\}$) durch (2.1.13) ein Vektor $\vec{v}$ definiert ist. Haben wir also einmal eine Basis festgelegt, können wir genauso gut mit diesen **geordneten Zahlentripeln** arbeiten, und zwar ordnen wir diese Zahlentripel gewöhnlich in einer Spalte an

$$\vec{v} \mapsto \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} \equiv (v_j) \equiv \underline{v}. \quad (2.1.14)$$

Wir bezeichnen die so definierten **Spaltenvektoren** mit demselben Symbol wie die geometrischen Vektoren, zur Unterscheidung aber mit einem Unterstrich anstelle eines Pfeilchens. Man muss sich dabei immer vergewissern, bzgl. welcher Basis eine solche Darstellung als Spaltenvektor gemeint ist.

Seien nun $\vec{v}$ und $\vec{w}$ beliebige räumliche Vektoren. Dann muss sich die Summe dieser Vektoren eindeutig durch die Spalten $\underline{v}$ und $\underline{w}$ darstellen lassen. Dies ist in der Tat einfach zu sehen, denn es gilt

$$\vec{v} + \vec{w} = \sum_{j=1}^3 v_j \vec{b}_j + \sum_{j=1}^3 w_j \vec{b}_j = \sum_{j=1}^3 (v_j \vec{b}_j + w_j \vec{b}_j) = \sum_{j=1}^3 (v_j + w_j) \vec{b}_j. \quad (2.1.15)$$

Es ist also der Summe der beiden Vektoren eindeutig der Spaltenvektor

$$\underline{v} + \underline{w} = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} + \begin{pmatrix} w_1 \\ w_2 \\ w_3 \end{pmatrix} = \begin{pmatrix} v_1 + w_1 \\ v_2 + w_2 \\ v_3 + w_3 \end{pmatrix} \quad (2.1.16)$$

zugeordnet. Es werden also einfach die entsprechenden Komponenten addiert.

Genauso zeigt man (*Übung*), dass $\lambda \vec{v}$ der Spaltenvektor

$$\lambda \underline{v} = \lambda \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} = \begin{pmatrix} \lambda v_1 \\ \lambda v_2 \\ \lambda v_3 \end{pmatrix} \quad (2.1.17)$$

zugeordnet ist, d.h. es werden die Komponenten einfach mit der Zahl λ multipliziert.

Ebenso ist es leicht einzusehen (*Übung*), dass die in Spalten angeordneten Zahlentripel mit den Rechenregeln (2.1.16) und (2.1.17) genauso wie die geometrisch definierten Vektoren einen dreidimensionalen reellen Vektorraum bilden, für den bzgl. Addition von Vektoren und Multiplikation von Vektoren mit reellen Zahlen dieselben Rechenregeln gelten. Da die Vektorkomponenten reelle Zahlen sind, nennt man diesen Vektorraum mit den so definierten Rechenoperationen $\mathbb{R}^3$. Wir haben also eine umkehrbar eindeutige Abbildung zwischen dem geometrischen Vektorraum E^3 und dem aus den Zahlentripeln des $\mathbb{R}^3$ gebildeten Vektorraum gefunden. Die Zahlentripel $\mathbb{R}^3$ bilden zudem die gleiche **algebraische Struktur** wie der geometrische Vektorraum. Man spricht bei solchen umkehrbar eindeutigen Abbildungen zwischen zwei solcherart gleichartigen algebraischen Strukturen, für die sich die Rechenoperationen zudem noch umkehrbar eindeutig entsprechen von **Isomorphismen**. Vom Standpunkt einer rein axiomatischen Definition eines Vektorraumes sind die durch einen Isomorphismus verknüpften algebraischen Strukturen nicht unterscheidbar. Sie sind vollständig zueinander **äquivalent** bzw. **isomorph**.

2.1.6 Basistransformationen

Es ist klar, dass wir die Betrachtungen des vorigen Abschnitts mit einer beliebigen Basis ausführen können. Es ist allerdings sehr wichtig, stets in Erinnerung zu behalten, dass diese Abbildung von der willkürlichen Wahl der **Basisvektoren** abhängig ist. In diesem Abschnitt befassen wir uns daher mit **Basistransformationen**, also der Frage, wie wir die Komponenten von Vektoren bzgl. einer Basis $\{\vec{b}_j\}_{j \in \{1,2,3\}}$ in die Komponenten des gleichen Vektors bzgl. einer anderen Basis $\{\vec{b}'_j\}_{j \in \{1,2,3\}}$ umrechnen können. Definitionsgemäß gilt

$$\vec{v} = \sum_{j=1}^3 v_j \vec{b}_j = \sum_{k=1}^3 v'_k \vec{b}'_k. \quad (2.1.18)$$

Da die Abbildung des Vektors $\vec{V} \in E^3$ auf die $\mathbb{R}^3$ -Vektoren $\underline{v}$ bzw. $\underline{v}'$ jeweils umkehrbar eindeutig sind, muss es entsprechend eine umkehrbar eindeutige Abbildung dieser beiden $\mathbb{R}^3$ -Vektoren geben.

Dazu müssen wir nur bedenken, dass es offenbar neun eindeutig bestimmte Zahlen T_{kj} gibt, so dass

$$\vec{b}_j = \sum_{k=1}^3 T_{kj} \vec{b}'_k \quad (2.1.19)$$

gilt. Das ist so, weil definitionsgemäß die Vektoren $\vec{b}'_k$ eine Basis bilden und somit jeder Vektor eine eindeutige Linearkombination dieser drei Basisvektoren ist, insbesondere also auch $\vec{b}_j$. Setzen wir (2.1.19) in (2.1.18) ein, ergibt sich

$$\vec{v} = \sum_{j=1}^3 v_j \vec{b}_j = \sum_{j=1}^3 \sum_{k=1}^3 T_{kj} v_j \vec{b}'_k = \sum_{k=1}^3 \left(\sum_{j=1}^3 T_{kj} v_j \right) \vec{b}'_k = \sum_{k=1}^3 v'_k \vec{b}'_k. \quad (2.1.20)$$

Da die Linearkombination von $\vec{v}$ bzgl. der Basis $\vec{b}'_k$ eindeutig ist, folgt daraus, dass notwendig

$$v'_k = \sum_{j=1}^3 T_{kj} v_j \quad (2.1.21)$$

sein muss. Kennen wir also die neun Zahlen T_{kj} , können wir direkt die Komponenten bzgl. der alten Basis in diejenigen der neuen umrechnen.

Gewöhnlich ordnet man die T_{kj} in ein rechteckiges 3×3 -Zahlenschema, eine sog. **Matrix** an:

$$\hat{T} \equiv (T_{kj}) = \begin{pmatrix} T_{11} & T_{12} & T_{13} \\ T_{21} & T_{22} & T_{23} \\ T_{31} & T_{32} & T_{33} \end{pmatrix}. \quad (2.1.22)$$

Die Menge dieser 3×3 -Matrizen mit reellen Zahlen als **Matrixelementen** bezeichnen wir mit $\mathbb{R}^{3 \times 3}$. Die Transformation (2.1.21) schreibt man dann kurz in der Form

$$\underline{v}' = \hat{T} \cdot \underline{v} \equiv \hat{T} \underline{v}. \quad (2.1.23)$$

Ausführlich geschrieben lautet die Rechenvorschrift

$$\begin{pmatrix} v'_1 \\ v'_2 \\ v'_3 \end{pmatrix} = \begin{pmatrix} T_{11} & T_{12} & T_{13} \\ T_{21} & T_{22} & T_{23} \\ T_{31} & T_{32} & T_{33} \end{pmatrix} \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} = \begin{pmatrix} T_{11}v_1 + T_{12}v_2 + T_{13}v_3 \\ T_{21}v_1 + T_{22}v_2 + T_{23}v_3 \\ T_{31}v_1 + T_{32}v_2 + T_{33}v_3 \end{pmatrix}. \quad (2.1.24)$$

Die einzelnen Einträge haben also die Form „Zeile $\times$ Spalte“, wobei man die Multiplikation jeweils als die angegebene Summe von Produkten anzusehen hat.

Die gleiche Argumentation führt auch umgekehrt zur Berechnung der Komponenten $\underline{v}$ aus den Komponenten $\underline{v}'$. Dazu müssen wir nur die entsprechende Matrix via

$$\vec{b}'_k = \sum_{j=1}^3 U_{jk} \vec{b}_j \quad (2.1.25)$$

eingeführen. Dann folgt

$$\underline{v} = \hat{U} \underline{v}'. \quad (2.1.26)$$

Kombinieren wir (2.1.26) mit (2.1.23) folgt für alle $\underline{v} \in \mathbb{R}^3$

$$\underline{v} = \hat{U} \underline{v}' = \hat{U} \hat{T} \underline{v}. \quad (2.1.27)$$

Zum einen führt dies das **Produkt zweier Matrizen** ein: $\hat{U} \hat{T}$ ist wieder eine $\mathbb{R}^{3 \times 3}$ -Matrix, und man bildet sie wieder nach dem Schema „Zeile $\times$ Spalte“. Zum anderen folgt aber aus (2.1.27), dass $\hat{U} \hat{T}$ eine Matrix sein muss, die alle Vektoren $\underline{v} \in \mathbb{R}^3$ auf sich selbst abbildet. Das kann nur die **Einheitsmatrix**

$$\mathbb{I}_3 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \equiv \text{diag}(1, 1, 1) \quad (2.1.28)$$

sein.

Es gilt also notwendig

$$\hat{U}\hat{T} = \mathbb{1}_3. \quad (2.1.29)$$

Es ist also $\hat{U}$ die **inverse Matrix** zu $\hat{T}$ im Sinne der Matrizenmultiplikation. Man schreibt daher auch

$$\hat{U} = \hat{T}^{-1}. \quad (2.1.30)$$

Wir können natürlich auch umgekehrt von (2.1.23) ausgehen und dann mit (2.1.26) auf

$$\underline{v}' = \hat{T}\underline{v} = \hat{T}\hat{U}\underline{v}' \quad (2.1.31)$$

schließen. Es gilt also mit denselben Argumenten wie eben auch

$$\hat{T}\hat{U} = \mathbb{1}_3 \Rightarrow \hat{T} = \hat{U}^{-1}. \quad (2.1.32)$$

Wir weisen bereits hier darauf hin, dass i.a. die *Matrizenmultiplikation nicht kommutativ ist*, d.h. für irgendwelche Matrizen $\hat{A}, \hat{B} \in \mathbb{R}^{3 \times 3}$ gilt i.a.

$$\hat{A}\hat{B} \neq \hat{B}\hat{A}! \quad (2.1.33)$$

Außerdem ist beileibe nicht jede Matrix invertierbar, d.h. haben wir irgendeine Matrix $\hat{A}$ gegeben, braucht die inverse Matrix $\hat{A}^{-1}$ nicht zu existieren. Das kennen wir schon von den reellen Zahlen: hier ist auch die 0 bzgl. der Multiplikation nicht invertierbar. Wir werden uns mit der Matrizenrechnung in Abschnitt 2.6 noch ausführlich beschäftigen und diese Fragen genauer studieren.

2.1.7 Das Skalarprodukt

Wir haben nun zwar schon einen sehr beachtlichen Teil der Euklidischen Geometrie in die Sprache der Vektoren übersetzt und damit als „Analytische Geometrie“ für die Physik bequemer handhabbar gemacht. Offensichtlich fehlt aber noch die Behandlung von **Winkeln**.

Dazu benötigen wir eine weitere Rechenoperation für zwei Vektoren $\vec{v}$ und $\vec{w}$, das **Skalarprodukt**. In der modernen mathematischen Literatur spricht man auch von einem **inneren Produkt**. Wir geben einfach die Definition des Skalarproduktes an und untersuchen dann seine Eigenschaften:

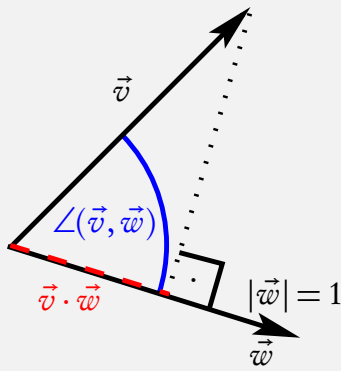
$$\vec{v} \cdot \vec{w} = |\vec{v}||\vec{w}| \cos[\angle(\vec{v}, \vec{w})]. \quad (2.1.34)$$

Dabei ist $\angle(\vec{v}, \vec{w}) = \angle(\vec{w}, \vec{v}) \in [0, \pi]$ der Winkel zwischen den beiden Vektoren, wenn man sie so verschiebt, dass ihre Anfangspunkte aufeinanderfallen (s. die folgende Abbildung).

Da der Cosinus im Intervall $[0, \pi]$ streng monoton fallend ist, wird durch das Skalarprodukt und die Länge der Vektoren der Winkel eindeutig definiert:

$$\angle(\vec{v}, \vec{w}) = \arccos\left(\frac{\vec{v} \cdot \vec{w}}{|\vec{v}||\vec{w}|}\right), \quad \vec{v}, \vec{w} \neq \vec{0}.$$

Falls mindestens einer der beiden Vektoren im Skalarprodukt der Nullvektor ist, ist der Winkel zwischen diesen Vektoren unbestimmt. Wir definieren daher noch zusätzlich, dass für alle Vektoren $\vec{w}$ stets $\vec{0} \cdot \vec{w} = \vec{w} \cdot \vec{0} = 0$ gelten soll. Insbesondere ist natürlich auch $\vec{0} \cdot \vec{0} = 0$.



Die geometrische Bedeutung des Skalarproduktes wird verständlich, wenn wir für $\vec{w}$ einen **Einheitsvektor**, also einen Vektor der Länge 1 wählen. Sei also $|\vec{w}| = 1$. Dann ist

$$\vec{v} \cdot \vec{w} = |\vec{v}| \cos[\angle(\vec{v}, \vec{w})] \quad \text{falls} \quad |\vec{w}| = 1. \quad (2.1.35)$$

Dies ist dem Betrag nach gerade die Länge der senkrechten Projektion des Vektors $\vec{v}$ auf die Richtung von $\vec{w}$ (vgl. Abbildung). Wegen des cos gilt hinsichtlich des Vorzeichens

$$\vec{v} \cdot \vec{w} \begin{cases} > 0 & \text{falls} \quad \angle(\vec{v}, \vec{w}) \in [0, \pi/2), \\ = 0 & \text{falls} \quad \angle(\vec{v}, \vec{w}) = \pi/2, \\ < 0 & \text{falls} \quad \angle(\vec{v}, \vec{w}) \in (\pi/2, \pi]. \end{cases} \quad (2.1.36)$$

Das Skalarprodukt verschwindet also entweder, wenn $\vec{v} = \vec{0}$ oder $\vec{w} = \vec{0}$ ist oder wenn die Vektoren aufeinander senkrecht stehen (denn es ist $\cos(\pi/2) = \cos 90^\circ = 0$). Für $\vec{v} \neq \vec{0}$ und $\vec{w} \neq \vec{0}$ schreibt man dann $\vec{v} \perp \vec{w}$ („ $\vec{v}$ steht senkrecht auf $\vec{w}$ “).

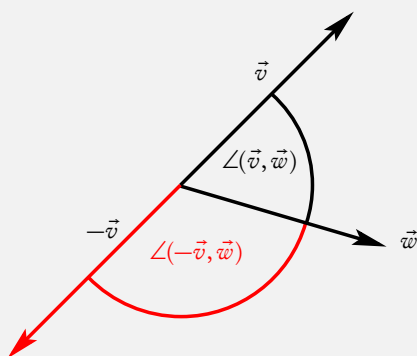
Dem Skalarprodukt eines Vektors mit sich selbst kommt eine besondere Bedeutung zu. Wegen $\cos 0 = 1$ folgt nämlich

$$\vec{v} \cdot \vec{v} \equiv \vec{v}^2 = |\vec{v}|^2. \quad (2.1.37)$$

Daraus folgt sofort

$$\vec{v} \cdot \vec{v} \geq 0, \quad \vec{v} \cdot \vec{v} = 0 \Leftrightarrow \vec{v} = \vec{0}. \quad (2.1.38)$$

Man sagt daher, dass das Skalarprodukt **positiv definit** ist.



Aus der Definition (2.1.34) ist unmittelbar klar, dass das Skalarprodukt **kommutativ** ist, d.h. es kommt auf die Reihenfolge der Multiplikation nicht an

$$\vec{v} \cdot \vec{w} = \vec{w} \cdot \vec{v}. \quad (2.1.39)$$

Weiter ist es auch linear in beiden Argumenten, d.h. es gilt

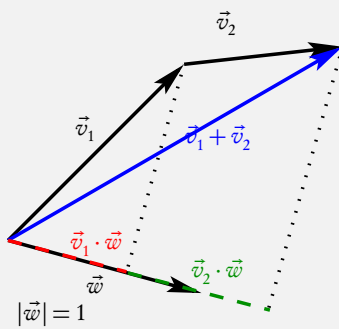
$$(\lambda \vec{v}) \cdot \vec{w} = |\lambda| |\vec{v}| |\vec{w}| \cos[\angle(\lambda \vec{v}, \vec{w})]. \quad (2.1.40)$$

Nun gilt aber gemäß der nebenstehenden Abbildung

$$\angle(\lambda \vec{v}, \vec{w}) = \begin{cases} \angle(\vec{v}, \vec{w}) & \text{falls} \quad \lambda > 0, \\ \pi - \angle(\vec{v}, \vec{w}) & \text{falls} \quad \lambda < 0. \end{cases} \quad (2.1.41)$$

Wegen $\cos(\pi - \alpha) = -\cos \alpha$ folgt also für $\lambda \neq 0$ aus (2.1.40)

$$(\lambda \vec{v}) \cdot \vec{w} = \text{sign } \lambda |\lambda| |\vec{v}| |\vec{w}| \cos[\angle \vec{v}, \vec{w}] = \lambda (\vec{v} \cdot \vec{w}) \equiv \lambda \vec{v} \cdot \vec{w}. \quad (2.1.42)$$



Außerdem entnimmt man der nebenstehenden Abbildung, dass für $|\vec{w}| = 1$ auch das Distributivgesetz, also

$$(\vec{v}_1 + \vec{v}_2) \cdot \vec{w} = \vec{v}_1 \cdot \vec{w} + \vec{v}_2 \cdot \vec{w} \quad (2.1.43)$$

gilt. Falls $\vec{w} = \vec{0}$ ist, gilt diese Gleichung sicher. Für einen Vektor $\vec{w} \neq \vec{0}$, der kein Einheitsvektor ist, können wir stets $\vec{w} = |\vec{w}| \vec{w}/|\vec{w}|$ schreiben. Nun ist $\vec{w}/|\vec{w}|$ ein Einheitsvektor (*warum?*), und wegen (2.1.42) folgt

$$\begin{aligned} (\vec{v}_1 + \vec{v}_2) \cdot \vec{w} &= |\vec{w}| (\vec{v}_1 + \vec{v}_2) \cdot \frac{\vec{w}}{|\vec{w}|} \\ &= |\vec{w}| \left(\vec{v}_1 \cdot \frac{\vec{w}}{|\vec{w}|} + \vec{v}_2 \cdot \frac{\vec{w}}{|\vec{w}|} \right) \\ &= \vec{v}_1 \cdot \vec{w} + \vec{v}_2 \cdot \vec{w}. \end{aligned} \quad (2.1.44)$$

Wegen des Kommutativgesetzes gilt dies freilich auch, wenn die Summe im zweiten Argument steht.

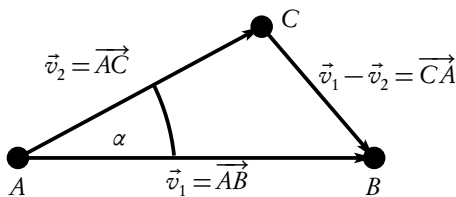
Das Skalarprodukt ist daher auch eine **symmetrische Bilinearform**. Symmetrisch heißt es deshalb, weil das Kommutativgesetz gilt und bilinear, weil es bzgl. beider Argumente eine Lineare Abbildung (in die reellen Zahlen) ist. Wir können nämlich (2.1.42) und (2.1.44) zusammenfassen zu

$$(\lambda_1 \vec{v}_1 + \lambda_2 \vec{v}_2) \cdot \vec{w} = \lambda_1 \vec{v}_1 \cdot \vec{w} + \lambda_2 \vec{v}_2 \cdot \vec{w}. \quad (2.1.45)$$

Mit der Definition des Skalarprodukts ist die Struktur des **Euklidischen Vektorraumes** nunmehr vollständig beschrieben. Ein Vektorraum heißt demnach euklidisch, wenn neben der Vektoralgebra mit den Operationen der Addition von Vektoren und der Multiplikation mit reellen Zahlen auch noch eine positiv definite Bilinearform definiert ist.

2.1.8 Geometrische Anwendungen des Skalarprodukts

Wir können nun das Skalarprodukt verwenden, um mit den Mitteln der analytischen Geometrie bekannte Sätze der Geometrie zu beweisen.



Als erstes beweisen wir den **Kosinus-Satz**. Seien drei Punkte A, B und C gegeben, die nicht auf einer Geraden liegen. Wir setzen dazu $\vec{v}_1 = \overrightarrow{AB}$ und $\vec{v}_2 = \overrightarrow{AC}$. Dann ist offenbar $\vec{v}_1 - \vec{v}_2 = \overrightarrow{AB} + \overrightarrow{CA} = \overrightarrow{CB} = \vec{v}_3$. Für die Länge der Seite $\overline{BC}$ gilt demnach

$$\begin{aligned} |BC|^2 &= \vec{v}_3 \cdot \vec{v}_3 = \vec{v}_3^2 \\ &= (\vec{v}_1 - \vec{v}_2)^2 \\ &= \vec{v}_1^2 + \vec{v}_2^2 - 2\vec{v}_1 \cdot \vec{v}_2 \\ &= |AB|^2 + |AC|^2 - 2|AB||AC|\cos\alpha, \end{aligned} \quad (2.1.46)$$

wobei $\alpha = \angle(\overrightarrow{AB}, \overrightarrow{AC})$ der von den Seiten $\overline{AB}$ und $\overline{AC}$ eingeschlossene Winkel ist, und das ist der Kosinus-Satz. Dabei haben wir ausgenutzt, dass wir mit dem Vektorprodukt formal wie mit Zahlen rechnen können und insbesondere durch Ausmultiplikation die gewohnten **binomischen Formeln** analog wie bei Zahlen gelten.

Falls $\alpha = \pi/2$, liegt offenbar ein **rechtwinkliges Dreieck** vor, und dann wird (2.1.46) wegen $\cos(\pi/2) = 0$

$$|BC|^2 = |AB|^2 + |AC|^2, \quad (2.1.47)$$

und das ist der **Satz des Pythagoras**.

Schließlich gilt wegen $|\cos \alpha| \leq 1$ stets die **Cauchy-Schwarzsche Ungleichung**

$$|\vec{v}_1 \cdot \vec{v}_2| \leq |\vec{v}_1| |\vec{v}_2|. \quad (2.1.48)$$

Das Gleichheitszeichen gilt nur, falls $\cos \alpha = 1$, d.h. $\alpha = 0$ (denn definitionsgemäß soll ja $\alpha \in [0, \pi]$ liegen) oder falls $\cos \alpha = -1$, d.h. $\alpha = \pi$ ist. MaW. das Gleichheitszeichen in (2.1.48) gilt genau dann, wenn $\vec{v}_1 \parallel \vec{v}_2$ ist.

Wir können nun die **Dreiecksungleichung** aus der positiven Definitheit des Skalarproduktes beweisen, denn es gilt

$$\begin{aligned} |\vec{v}_1 + \vec{v}_2|^2 &= (\vec{v}_1 + \vec{v}_2) \cdot (\vec{v}_1 + \vec{v}_2) \\ &= \vec{v}_1^2 + \vec{v}_2^2 + 2\vec{v}_1 \cdot \vec{v}_2 \\ &\leq \vec{v}_1^2 + \vec{v}_2^2 + 2|\vec{v}_1| |\vec{v}_2| \\ &\leq \vec{v}_1^2 + \vec{v}_2^2 + 2|\vec{v}_1| |\vec{v}_2| = (|\vec{v}_1| + |\vec{v}_2|)^2, \end{aligned} \quad (2.1.49)$$

bzw., weil immer nur positive Zahlen quadriert werden,

$$|\vec{v}_1 + \vec{v}_2| \leq |\vec{v}_1| + |\vec{v}_2|. \quad (2.1.50)$$

Umgekehrt folgt aus der positiven Definitheit des Skalarproduktes auch die Cauchy-Schwarzsche Ungleichung (2.1.48).

2.1.9 Das Skalarprodukt im $\mathbb{R}^3$

Wegen der Bilinearität können wir das Skalarprodukt zwischen zwei Vektoren $\vec{v}$ und $\vec{w}$ auch mit Hilfe der Vektorkomponenten $\underline{v}$ und $\underline{w}$ bzgl. einer beliebigen Basis ausdrücken. Dazu müssen wir allerdings offenbar die Skalarprodukte zwischen den Basisvektoren kennen. Dies ergibt wieder eine Matrix.

Wir definieren also die **Metrikkomponenten** bzgl. der gegebenen Basis $\vec{b}_j$ gemäß

$$g_{jk} = \vec{b}_j \cdot \vec{b}_k. \quad (2.1.51)$$

Für zwei beliebige Vektoren $\vec{v}$ und $\vec{w}$ mit den Komponenten $\underline{v} = (v_j)$ und $\underline{w} = (w_k)$ folgt dann wegen der Bilinearität des Skalarproduktes

$$\vec{v} \cdot \vec{w} = \left(\sum_{j=1}^3 v_j \vec{b}_j \right) \cdot \left(\sum_{k=1}^3 w_k \vec{b}_k \right) = \sum_{j=1}^3 \sum_{k=1}^3 v_j w_k \vec{b}_j \cdot \vec{b}_k = \sum_{j=1}^3 \sum_{k=1}^3 g_{jk} v_j w_k = g_{jk} v_j w_k. \quad (2.1.52)$$

Dabei haben wir im letzten Schritt von der Einsteinschen Summenkonvention Gebrauch gemacht, wonach über in einem Ausdruck doppelt auftretende Indizes zu summieren ist.

2.1. Geometrische Einführung von Euklidischen Vektoren

Wir schreiben nun (g_{jk}) als Matrix, machen aber die Tatsache, dass es sich diemal nicht um die Transformationsmatrix zwischen Basisvektoren sondern um die **darstellende Matrix einer Bilinearform** (in diesem Fall des Skalarproduktes) handelt, durch einen übergestellten Doppelpfeil kenntlich¹:

$$\overleftrightarrow{g} = (g_{jk}) = \begin{pmatrix} g_{11} & g_{12} & g_{13} \\ g_{21} & g_{22} & g_{23} \\ g_{31} & g_{32} & g_{33} \end{pmatrix}. \quad (2.1.53)$$

Wir können nun (2.1.52) mit Hilfe von Matrix-Vektorprodukten schreiben, wenn wir den linken Vektor als Zeile schreiben. Dazu definieren wir

$$\underline{v}^T = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix}^T = (v_1 \ v_2 \ v_3) \equiv (v_1, v_2, v_3). \quad (2.1.54)$$

Man nennt diese Operation an Spaltenvektoren des $\mathbb{R}^3$ **Transposition**. Diese Operation kann man auch mit Matrizen vornehmen. Man schreibt einfach alle Spalten als Zeilen:

$$\overleftrightarrow{g}^T = \begin{pmatrix} g_{11} & g_{12} & g_{13} \\ g_{21} & g_{22} & g_{23} \\ g_{31} & g_{32} & g_{33} \end{pmatrix} \Leftrightarrow (g^T)_{jk} = g_{kj}. \quad (2.1.55)$$

Da das Skalarprodukt kommutativ ist, gilt allerdings

$$g_{jk} = \vec{b}_j \cdot \vec{b}_k = \vec{b}_k \cdot \vec{b}_j = g_{kj} \Leftrightarrow \overleftrightarrow{g}^T = \overleftrightarrow{g}. \quad (2.1.56)$$

Mit diesen Definitionen folgt, dass

$$\vec{v} \cdot \vec{w} = \underline{v}^T \overleftrightarrow{g} \underline{w} = \underline{w}^T \overleftrightarrow{g} \underline{v}. \quad (2.1.57)$$

Es ist nun einfach, die Metrikkomponenten bzgl. einer beliebigen anderen Basis $\vec{b}'_k$ durch die der Basis $\vec{b}_j$ auszudrücken. Dazu benötigen wir lediglich (2.1.25).

Wegen der Bilinearität des Skalarprodukts folgt dann

$$g'_{jk} = \vec{b}'_j \cdot \vec{b}'_k = \sum_{l,m=1}^3 U_{lj} U_{mk} \vec{b}_l \cdot \vec{b}_m = \sum_{l,m=1}^3 U_{lj} U_{mk} g_{lm}. \quad (2.1.58)$$

In Matrixschreibweise ergibt dies (*Übung!*)

$$\overleftrightarrow{g}' = \hat{U}^T \overleftrightarrow{g} \hat{U}. \quad (2.1.59)$$

Dass diese Transformationsvorschrift tatsächlich dazu führt, dass das Skalarprodukt von der Wahl der Basis unabhängig ist, ergibt sich auch aus (2.1.26) für die Vektorkomponenten:

$$\underline{v}^T \overleftrightarrow{g} \underline{w} = (\hat{U} \underline{v}')^T \overleftrightarrow{g} \hat{U} \underline{w}' = \underline{v}'^T \hat{U}^T \overleftrightarrow{g} \hat{U} \underline{w}' \stackrel{2.1.59}{=} \underline{v}'^T \overleftrightarrow{g}' \underline{w}'. \quad (2.1.60)$$

¹Es ist sehr wichtig, eine solche Unterscheidung zu treffen, denn wie wir gleich sehen werden, unterscheidet sich das Transformationsverhalten einer solchen Bilinearform (auch **Tensor zweiter Stufe** genannt) unter einem Basiswechsel entscheidend von dem der Vektorkomponenten. Eigentlich wäre es daher auch adäquater gewesen, eine Unterscheidung bereits bei der Bezeichnung der Komponenten zu treffen. Dies geschieht in der weiterführenden Literatur durch die Verwendung von hoch- und tiefgestellten Indizes. Dies habe ich in diesem Skript vermieden, um die Kompatibilität mit der in der Theorie-Vorlesung gewählten Schreibweise zu wahren.

Dabei haben wir verwendet, dass für Matrizenmultiplikationen und Transposition die Reihenfolge beim Transponieren umgekehrt werden muss. In unserem Beispiel ist in der Tat

$$(\hat{U}\underline{v}')^T = (U_{1k}v'_k, U_{2k}v'_k, U_{3k}v'_k) = (v'_1, v'_2, v'_3) \begin{pmatrix} U_{11} & U_{21} & U_{31} \\ U_{12} & U_{22} & U_{32} \\ U_{13} & U_{23} & U_{33} \end{pmatrix} = \underline{v}'^T \hat{U}^T. \quad (2.1.61)$$

2.2 Axiomatische Begründung der linearen Algebra und Geometrie

Bislang haben wir die diversen Rechenoperationen mit Vektoren geometrisch motiviert und uns dabei auf die dreidimensionale (räumliche) Euklidische Geometrie berufen. Es ist nun gerade für die theoretische Physik wichtig, diese Betrachtungen auf abstraktere Füße zu stellen. In der Vorlesung Theoretische Physik 3 werden wir z.B. die **Spezielle Relativitätstheorie** kennenlernen, die sich am elegantesten und leicht verständlichsten in einem vierdimensionalen Punktraum und entsprechenden vierdimensionalen Vektoren beschreiben lässt, der sog. **Minkowski-Raumzeit**. Die Algebra der Vektoren ist dabei nicht wesentlich komplizierter als die aus der soeben geometrischen Anschauung gewonnenen Rechenregeln für den Euklidischen dreidimensionalen Raum. Nicht zuletzt ist die folgende Zusammenstellung der Axiome auch eine gute Merkhilfe der wichtigsten Formeln.

Wir beginnen mit der axiomatischen Begründung des **Vektorraums**. Ein Vektorraum ist zunächst eine abstrakte Menge V von Objekten, für die zunächst eine Abbildung $V \times V \rightarrow V$, die **Addition**, definiert wird. Diese mathematische Pfeilschreibweise für eine Abbildung bedeutet hierbei folgendes: auf der linken Seite steht der **Definitionsbereich** oder **Urbildbereich** der Abbildung. In unserem Fall sind das **geordnete Paare** von Elementen aus der Menge V . Hier steht $V \times V$ nämlich für ein **Mengenprodukt**, welches einfach eine neue Menge ist, die aus **geordneten Paaren** von Elementen aus V besteht. Dabei ist ausdrücklich i.a. $(\vec{v}_1, \vec{v}_2) \neq (\vec{v}_2, \vec{v}_1)$, d.h. die Reihenfolge bei dem geordneten Paar soll ausdrücklich berücksichtigt werden. Rechts vom Pfeil steht die **Wertemenge** oder der **Bildbereich**. Für unser Fall besagt dies, dass jedem geordneten Paar von Vektoren aus V eindeutig ein Vektor, ebenfalls aus V , zugeordnet wird. Die Addition schreiben wir dabei mit einem Pluszeichen. Dies drückt man formal wie folgt aus $(\vec{v}_1, \vec{v}_2) \mapsto \vec{v}_1 + \vec{v}_2$.

Dies mag auf den ersten Blick ein wenig pedantisch und kompliziert erscheinen. Es ist jedoch oft von großem Nutzen, sich klar zu machen, welchen Definitionsbereich und Wertebereich Rechenoperationen eigentlich haben, und dafür ist diese pedantische Schreibweise der Mathematiker äußerst nützlich.

Die Rechenoperationen müssen nun noch weiter spezifiziert werden. Dabei ist es aus mathematischer Sicht unerheblich, um welche konkreten Objekte es sich bei dem, was wir „Vektoren“ nennen, eigentlich handelt. Es werden einfach „Spielregeln“, die **Axiome**, aufgestellt, die formale Operationen festlegen. Für unseren Fall der Addition von Vektoren sind dies die folgenden Axiome:

1. Für beliebige drei Vektoren $\vec{v}_1, \vec{v}_2, \vec{v}_3 \in E^3$ gilt stets das **Assoziativgesetz**

$$(\vec{v}_1 + \vec{v}_2) + \vec{v}_3 = \vec{v}_1 + (\vec{v}_2 + \vec{v}_3). \quad (2.2.1)$$

2. Es existiert ein **neutrales Element der Addition**, das wir mit $\vec{0} \in V$ bezeichnen. Es zeichnet sich dadurch aus, dass für alle Vektoren $\vec{v} \in V$

$$\vec{v} + \vec{0} = \vec{v} \quad (2.2.2)$$

gilt.

3. Zu jedem Vektor $\vec{v} \in V$ existiert ein bzgl. der Addition **inverses Element** $(-\vec{v}) \in V$, für das

$$\vec{v} + (-\vec{v}) = \vec{0}. \quad (2.2.3)$$

4. Für beliebige zwei Vektoren $\vec{v}_1, \vec{v}_2 \in E^3$ gilt das **Kommutativgesetz**

$$\vec{v}_1 + \vec{v}_2 = \vec{v}_2 + \vec{v}_1. \quad (2.2.4)$$

Es führt für diese Physikvorlesung zu weit, den gesamten Apparat der linearen Algebra, die wir oben durch Berufung auf die geometrische Anschauung hergeleitet haben, aus diesen Axiomen zu beweisen. Dies ist Gegenstand der entsprechenden Mathematikvorlesung.

Als ein Beispiel möge der wichtige Satz dienen, dass das neutrale Element der Addition $\vec{0}$ eindeutig ist. Nehmen wir dazu an, es sei auch $\vec{0}'$ ein neutrales Element der Addition. Wegen der Eigenschaft (2.2.2) des neutralen Elements $\vec{0}$ gilt

$$\vec{0}' = \vec{0}' + \vec{0} \stackrel{(2.2.4)}{=} \vec{0} + \vec{0}'. \quad (2.2.5)$$

Da nun aber $\vec{0}'$ voraussetzungsgemäß ebenfalls ein neutrales Element ist und gemäß (2.2.2) gilt, ist demnach

$$\vec{0}' = \vec{0} + \vec{0}' = \vec{0}, \quad (2.2.6)$$

d.h. dass notwendig die beiden neutralen Elemente der Addition gleich sind.

Wir bemerken noch, dass eine algebraische Struktur, die auf einer Menge wie die eben definierte Addition definiert ist und (2.2.1)-(2.2.3) genügt als **Gruppe** bezeichnet wird. So einfach diese Axiome auch anmuten mögen, so liefern sie doch eine sehr reichhaltige mathematische Struktur. Die Gruppentheorie ist ein riesiges Teilgebiet der Mathematik, und Teile der Gruppentheorie spielen übrigens für die Physik eine herausragende Rolle, wie im Verlauf des Physikstudiums noch klar wird. Nimmt man zu diesen drei Axiomen noch das Kommutativgesetz (2.2.4) hinzu, spricht man von einer **Abelschen Gruppe**. Wir können also die Axiome der Vektoraddition auch kurz zusammenfassen, indem wir bemerken, dass die Vektoraddition eine Abelsche Gruppe auf der Menge der Vektoren V bildet.

Um die so definierte additive Gruppe zu einem Vektorraum zu erweitern, benötigen wir die Multiplikation von Vektoren mit einem **Skalar**. Dabei sind die Skalare Elemente eines **Zahlenkörpers**, der in unserem Fall stets der Körper der reellen Zahlen $\mathbb{R}$ ist. Auch diese algebraische Struktur kann wiederum axiomatisch begründet werden. Wir wollen dies hier der Übersichtlichkeit halber aber unterlassen. Die Multiplikation mit Skalaren erfüllt die folgenden Axiome:

1. Für alle $\lambda, \mu \in \mathbb{R}$ und alle $\vec{v} \in V$ gilt $(\lambda + \mu)\vec{v} = \lambda\vec{v} + \mu\vec{v}$ und $\lambda(\mu\vec{v}) = (\lambda\mu)\vec{v}$.
2. Für alle $\vec{v} \in V$ ist $1\vec{v} = \vec{v}$.
3. für alle $\vec{v}, \vec{w} \in V$ und alle $\lambda \in \mathbb{R}$ gilt $\lambda(\vec{v} + \vec{w}) = \lambda\vec{v} + \lambda\vec{w}$.

Man weist leicht einige weitere einfache Rechenregeln nach (*Übung!*), z.B.

$$0\vec{v} = \vec{0}, \quad \lambda\vec{0} = \vec{0}, \quad (-1)\vec{v} = -\vec{v}, \dots \quad (2.2.7)$$

Mit diesen Rechenregeln ist der Begriff des Vektorraumes bereits vollständig bestimmt. Als wichtige Begriffe ergeben sich daraus, wie oben hergeleitet, die **Linearkombinationen** und **Basen** und im Zusammenhang damit auch der **Dimension** eines Vektorraums. Wir werden unten als weitere wichtige Strukturen die **linearen Abbildungen** verschiedener Art kennenlernen, deren Untersuchung weitere Eigenschaften der Vektorräume charakterisieren. Es ist übrigens ohnehin ein Merkmal der modernen axiomatischen Methode der Mathematik, Strukturen einzuführen und dann entsprechende diese Strukturen charakterisierende Abbildungen zwischen diesen Strukturen zu analysieren.

Für die Geometrie und damit auch für deren Anwendung in der Physik als Modell für den physikalischen Raum ist weiter die Einführung des Skalarprodukts wichtig. Es handelt sich um eine Abbildung $V \times V \rightarrow \mathbb{R}$, $\vec{v}, \vec{w} \mapsto \vec{v} \cdot \vec{w}$, das die folgenden Axiome erfüllt

1. Für alle $\vec{v}, \vec{w}, \vec{x} \in E^3$ gilt $\vec{v} \cdot (\vec{w} + \vec{x}) = \vec{v} \cdot \vec{w} + \vec{v} \cdot \vec{x}$.
2. Für alle $\lambda \in \mathbb{R}$ und alle $\vec{v}, \vec{w} \in E^3$ gilt $(\lambda\vec{v}) \cdot \vec{w} = \lambda(\vec{v} \cdot \vec{w})$.
3. Für alle $\vec{v}, \vec{w} \in E^3$ gilt $\vec{v} \cdot \vec{w} = \vec{w} \cdot \vec{v}$.
4. Für alle $\vec{v} \in E^3$ gilt $\vec{v}^2 := \vec{v} \cdot \vec{v} \geq 0$.
5. Es ist $\vec{v}^2 = 0$ genau dann wenn, $\vec{v} = \vec{0}$ ist.

Eine Abbildung $V \times V \rightarrow V$, welche die ersten drei Axiome erfüllt, heißt **symmetrische Bilinearform**. Gilt auch noch das vierte Axiom, heißt die Bilinearform **positiv semidefinit**. Ist schließlich auch noch das fünfte Axiom erfüllt, heißt die Bilinearform **positiv definit**.

Wir wollen als Beispiel mit dem Umgang des so axiomatisch definierten Skalarprodukts die **Cauchy-Schwarzsche Ungleichung** (2.1.48) beweisen, wobei der **Betrag eines Vektors** durch

$$|\vec{v}| = \sqrt{\vec{v}^2} \geq 0 \quad (2.2.8)$$

definiert ist. Wegen der positiven Definitheit des Skalarprodukts gilt für alle Vektoren $\vec{v} \neq \vec{0}$ und $\vec{w} \in V$

$$P(\lambda) = \frac{1}{\vec{v}^2} (\lambda\vec{v} + \vec{w})^2 \geq 0. \quad (2.2.9)$$

Dabei dürfen wir wegen der positiven Definitheit (5. Axiom des Skalarprodukts) durch $\vec{v}^2 \neq 0$ dividieren. Wir können weiter das Skalarprodukt ausmultiplizieren (zur *Übung* mache man sich das anhand der obigen Axiome klar):

$$P(\lambda) = \lambda^2 + \frac{2\vec{v} \cdot \vec{w}}{\vec{v}^2} \lambda + \frac{\vec{w}^2}{\vec{v}^2} \geq 0. \quad (2.2.10)$$

Dieses quadratische Polynom können wir nun noch weiter umformen zu

$$P(\lambda) = \left(\lambda + \frac{\vec{v} \cdot \vec{w}}{\vec{v}^2} \right)^2 + \frac{\vec{w}^2}{\vec{v}^2} - \frac{(\vec{v} \cdot \vec{w})^2}{(\vec{v}^2)^2} \geq 0. \quad (2.2.11)$$

Da diese Ungleichung für alle $\lambda \in \mathbb{R}$ gilt, insbesondere auch für $\lambda = -\vec{v} \cdot \vec{w} / \vec{v}^2$ muss notwendig

$$\frac{\vec{w}^2}{\vec{v}^2} - \frac{(\vec{v} \cdot \vec{w})^2}{(\vec{v}^2)^2} \geq 0 \quad (2.2.12)$$

gelten.

Mit $(\vec{v}^2)^2 > 0$ multipliziert ergibt sich wegen der Definition des Betrages von Vektoren (2.2.8) schließlich

$$(\vec{v} \cdot \vec{w})^2 \leq \vec{v}^2 \vec{w}^2 \Rightarrow |\vec{v} \cdot \vec{w}| \leq |\vec{v}| |\vec{w}|. \quad (2.2.13)$$

Aus dem Beweis folgt auch, dass das Gleichheitszeichen genau dann gilt, wenn es ein $\lambda \in \mathbb{R}$ gibt, so dass $P(\lambda) = 0$ ist, also wenn $\vec{w} = -\lambda \vec{v}$ ist, d.h. wenn die Vektoren **kollinear** sind.

Mit der Rechnung (2.1.49) haben wir gezeigt, dass für den via (2.2.8) definierten Betrag die **Dreiecksungleichung** gilt und damit der Betrag die Axiome einer **Norm** erfüllt, die wir in Abschnitt 2.1.3 besprochen haben. Wegen der Gültigkeit der Cauchy-Schwarzschen Ungleichung (2.2.13) ist es auch sinnvoll, den Winkel zwischen zwei nicht verschwindenden Vektoren gemäß

$$\cos[\angle(\vec{v}, \vec{w})] = \frac{\vec{v} \cdot \vec{w}}{|\vec{v}| |\vec{w}|} \quad (2.2.14)$$

einzuführen und damit die Winkelmessung zu begründen.

Man kann zeigen, dass sich mit diesen Regeln alle Sätze der Euklidischen Geometrie ergeben. Dazu muss man nur noch die Punkte als Menge einführen und postulieren, dass es entsprechend den oben im umgekehrten Sinne aus den Eigenschaften des Euklidischen Punktraumes erschlossenen Regeln für Verschiebungsvektoren $\vec{AB}$ gelten. Wir wollen auch diese Betrachtungen nicht im Einzelnen hier ausführen. Wir bemerken nur, dass ein solches Konstrukt aus einem (reellen endlichdimensionalen) Euklidischen Vektorraum und einer Menge von Punkten als **Euklidischer affiner Punktraum** bezeichnet wird. Eine Verallgemeinerung dieser mathematisch-abstrakten Definition spielt später in der Begründung der **Speziellen Relativitätstheorie** noch eine wichtige Rolle, wo die Beschreibung der Strukturen von Raum und Zeit durch die sog. **Minkowski-Raumzeit** erfolgt. Die Minkowski-Raumzeit entpuppt sich dabei als vierdimensionaler affiner Punktraum, wobei im Vektorraum allerdings kein positiv definites Skalarprodukt sondern eine indefinite nicht ausgeartete Bilinearform definiert ist. Es handelt sich also um einen sog. **pseudo-Euklidischen affinen Punktraum**.

2.3 Kartesische Basen und orthogonale Transformationen

Mit der Einführung eines Skalarprodukts gibt es allerdings auch eine besonders bequeme Klasse von **Basen**, die **Orthonormalbasen** oder **Kartesischen Basen**. Dazu wählt man als Basisvektoren beliebige drei paarweise zueinander senkrechte Einheitsvektoren, ein sog. **Dreibein**. Anschaulich ist unmittelbar klar, dass es beliebig viele solcher Dreibeine gibt und damit auch beliebig viele Orthonormalbasen. Wie wir gleich sehen werden, vereinfachen sie das Rechnen mit Vektorkomponenten und dem Matrixprodukt bzw. den dazugehörigen Metrikkomponenten sowie die Transformation von einem Orthonormalsystem zu einem anderen erheblich.

Es sei also $\{\vec{e}_j\}_{j \in \{1,2,3\}}$ eine beliebige Orthonormalbasis. Voraussetzungsgemäß sind diese drei Vektoren auf 1 normiert und stehen paarweise aufeinander senkrecht. Es gilt also

$$g_{jk} = \vec{e}_j \cdot \vec{e}_k = \delta_{jk} = \begin{cases} 1 & \text{falls } j = k, \\ 0 & \text{falls } j \neq k. \end{cases} \quad (2.3.1)$$

Man nennt δ_{jk} das **Kronecker-Symbol** (Leopold Kronecker 1823-1891). Es ist klar, dass es die Komponenten der Einheitsmatrix repräsentiert.

Bzgl. einer kartesischen Basis wird die darstellende Matrix des Skalarproduktes also extrem bequem, denn es gilt gemäß (2.3.1)

$$\overleftrightarrow{g} = \overleftrightarrow{\delta} = \mathbb{1}_3. \quad (2.3.2)$$

Für kartesische Koordinaten gilt also gemäß (2.1.57)

$$\vec{v} \cdot \vec{w} = \underline{v} \cdot \underline{w} = \underline{v}^T \overleftrightarrow{g} \underline{w} = \underline{v}^T \underline{w} = \sum_{j,k=1}^3 \delta_{jk} v_j w_k = \sum_{j=1}^3 v_j w_j = v_j w_j = v_1 w_1 + v_2 w_2 + v_3 w_3. \quad (2.3.3)$$

Wir untersuchen nun wieder die Frage, wie sich Basiswechsel auswirken. Wir betrachten zunächst die Transformation zwischen den Basisvektoren einer kartesischen Basis $\{\vec{e}_j\}_{j \in \{1,2,3\}}$ und einer anderen beliebigen Basis $\{\vec{b}'_j\}_{j \in \{1,2,3\}}$. Es gilt wieder wie für ganz allgemeine Basen (2.1.20) und (2.1.25):

$$\vec{e}_j = \sum_{k=1}^3 T_{kj} \vec{b}'_k, \quad \vec{b}'_k = \sum_{j=1}^3 U_{jk} \vec{e}_j, \quad (2.3.4)$$

wobei die Matrizen $\hat{T}$ und $\hat{U}$ zueinander invers sind: $\hat{T} \hat{U} = \hat{U} \hat{T} = \mathbb{1}_3$ bzw. $\hat{U} = \hat{T}^{-1}$. Insbesondere sind diese Matrizen notwendig invertierbar, wenn sowohl die $\vec{e}_j$ als auch die $\vec{b}'_j$ jeweils eine Basis bilden, also vollständig und linear unabhängig sind.

Für die Vektorkomponenten gelten entsprechend (2.1.21) bzw. (2.1.22) und (2.1.27):

$$\underline{v}' = \hat{T} \underline{v}, \quad \underline{v} = \hat{U} \underline{v}' = \hat{T}^{-1} \underline{v}'. \quad (2.3.5)$$

Die Tatsache, dass die $\vec{e}_j$ hierbei ein Orthonormalsystem bilden, führt in diesem Fall noch nicht zu irgendwelchen Vereinfachungen. Betrachten wir nun jedoch, wie sich die darstellenden Matrizen für das Skalarprodukt verhalten. Hier ist (2.3.2) entscheidend, denn für kartesische Basen ist demnach die darstellende Matrix einfach die Einheitsmatrix! Die Transformationsvorschrift (2.1.59) vereinfacht sich nämlich zu

$$\overleftrightarrow{g}' = \hat{U}^T \overleftrightarrow{g} \hat{U} = \hat{U}^T \mathbb{1}_3 \hat{U} = \hat{U}^T \hat{U}. \quad (2.3.6)$$

Daraus lässt sich die für spätere Überlegungen sehr wichtige Schlussfolgerung ziehen, dass die Metrikkomponenten bzgl. einer beliebigen Basis eine invertierbare Matrix bilden. Eine Bilinearform, deren Metrikkomponenten diese Eigenschaft besitzen, heißt **nicht ausgeartete Bilinearform**.

Das **Skalarprodukt** ist also eine **nicht ausgeartete positiv definite symmetrische Bilinearform**.

Jede Bilinearform mit diesen Eigenschaften kann als Skalarprodukt verwendet werden und definiert dann auf dem jeweils betrachteten Vektorraum genau die geometrische Struktur, die diesen Vektorraum zu einem Euklidischen Vektorraum macht, d.h. sie induziert eine Norm, die die Länge von Vektoren festlegt, den Begriff des Winkels zwischen zwei Vektoren und insbesondere die Eigenschaft, dass Vektoren orthogonal zueinander sein können (also aufeinander senkrecht stehen).

Um nun also zu zeigen, dass

$$\overleftrightarrow{g}' = (\overleftrightarrow{g})^{-1} \quad (2.3.7)$$

2.3. Kartesische Basen und orthogonale Transformationen

existiert, müssen wir uns bereits ein wenig in die Matrizenrechnung einarbeiten. Wir wissen, dass aufgrund der Tatsache, dass $\hat{U}$ und $\hat{T}$ zueinander inverse Basistransformationen repräsentieren, dass $\hat{U}$ eine invertierbare Matrix ist und dass $\hat{U}^{-1} = \hat{T}$ ist. Wir wissen weiter, dass dies impliziert, dass $\hat{T}\hat{U} = \hat{U}\hat{T} = \mathbb{1}_3$ ist. Dies ist keinesfalls selbstverständlich, weil i.a. die *Matrizenmultiplikation nicht kommutativ* ist. Wenn aber die Matrizen zueinander invers sind, ergibt ihr Produkt unabhängig von der Reihenfolge stets die Einheitsmatrix. Seien nun $\hat{A}$ und $\hat{B}$ beide invertierbar. Dann ist auch das Produkt $\hat{A}\hat{B}$ eine invertierbare Matrix, und es gilt

$$(\hat{A}\hat{B})^{-1} = \hat{B}^{-1}\hat{A}^{-1}. \quad (2.3.8)$$

Man muss hierbei strikt auf die *Reihenfolge* der Matrixmultiplikation achten! Zum Beweis rechnen wir

$$(\hat{A}\hat{B})(\hat{B}^{-1}\hat{A}^{-1}) = \hat{A}(\hat{B}\hat{B}^{-1})\hat{A}^{-1} = \hat{A}\mathbb{1}_3\hat{A}^{-1} = \hat{A}\hat{A}^{-1} = \mathbb{1}_3, \quad (2.3.9)$$

d.h. $\hat{B}^{-1}\hat{A}^{-1}$ ist tatsächlich die inverse Matrix zu $\hat{A}\hat{B}$, wie in (2.3.8) behauptet.

Weiter behaupten wir, dass mit $\hat{U}$ auch $\hat{U}^T$ invertierbar ist, und dass

$$(\hat{U}^T)^{-1} = (\hat{U}^{-1})^T = \hat{T}^T \quad (2.3.10)$$

ist. Zum Beweis müssen wir nur das entsprechende Matrixprodukt in Komponenten ausschreiben:

$$(\hat{U}^T \hat{T}^T)_{jl} = \sum_{k=1}^3 (\hat{U}^T)_{jk} (\hat{T}^T)_{kl} = \sum_{k=1}^3 U_{kj} T_{lk} = (\hat{T} \hat{U})_{lj} = \delta_{lj} = \delta_{jl}. \quad (2.3.11)$$

Das bedeutet aber

$$\hat{U}^T \hat{T}^T = \mathbb{1}_3, \quad (2.3.12)$$

womit (2.3.10) bewiesen ist.

Damit ist wegen (2.3.6) aber $\vec{g} \leftrightarrow$ invertierbar. Wenden wir nämlich (2.3.8) und (2.3.10) auf das entsprechende Produkt an, erhalten wir

$$\vec{g} \leftrightarrow (\vec{g}')^{-1} = (\hat{U}^T \hat{U})^{-1} = \hat{U}^{-1} (\hat{U}^T)^{-1} = \hat{T} \hat{T}^T. \quad (2.3.13)$$

Nun betrachten wir den Fall, dass auch die zweite Basis eine Orthonormalbasis ist, d.h. $\vec{b}'_j = \vec{e}'_j$ sei auch eine Orthonormalbasis. Dann gilt auch $\vec{g}' \leftrightarrow$. Dann gilt aber gemäß (2.3.6)

$$\vec{g}' \leftrightarrow \hat{U}^T \hat{U} = \mathbb{1}_3. \quad (2.3.14)$$

Damit folgt aber für diesen speziellen Fall, für die Matrix des Basiswechsels zwischen zwei orthonormalen Basen

$$\hat{T} = \hat{U}^{-1} = \hat{U}^T. \quad (2.3.15)$$

Man nennt Matrizen mit dieser Eigenschaft **Orthogonale Matrizen**. Natürlich gilt auch für die Umkehrtransformation

$$\hat{T}^{-1} = \hat{U} = (\hat{U}^T)^T \stackrel{(2.3.15)}{=} \hat{T}^T. \quad (2.3.16)$$

Also ist auch $\hat{T}$ eine orthogonale Matrix.

Es ist auch sehr einfach, diese Transformationsmatrizen zu bestimmen, wenn man die beiden Orthonormalbasen $\vec{e}_j$ und $\vec{e}'_k$ gegeben hat. Man muss sich natürlich stets vergewissern, dass beide Basen auch tatsächlich Orthonormalbasen sind! Dann ist nämlich

$$\vec{e}_j = \sum_{k=1}^3 T_{kj} \vec{e}'_k. \quad (2.3.17)$$

Dann folgt wegen $\vec{e}'_l \cdot \vec{e}'_k = \delta_{lk}$

$$\vec{e}'_l \cdot \vec{e}_j = \vec{e}'_l \cdot \sum_{k=1}^3 T_{kj} \vec{e}'_k = \sum_{k=1}^3 T_{kj} \vec{e}'_l \cdot \vec{e}'_k = \sum_{k=1}^3 T_{kj} \delta_{lk} = T_{lj}. \quad (2.3.18)$$

Wenden wir nun (2.3.16) an. Daraus folgt

$$U_{jk} = (\hat{T}^T)_{jk} = T_{kj} = \vec{e}'_k \cdot \vec{e}_j = \vec{e}_j \cdot \vec{e}'_k. \quad (2.3.19)$$

Das folgt natürlich sofort auch aus

$$\vec{e}'_k = \sum_{j=1}^3 U_{jk} \vec{e}_j \quad (2.3.20)$$

und der Orthonormalität der $\vec{e}_j$ (*nachrechnen!*).

Die Matrizen für die Basiswechsel zwischen **Orthonormalbasen** sind **orthogonale Matrizen**.

2.4 Das Vektorprodukt

Schließlich ist noch das sog. **Vektorprodukt** zwischen zwei Vektoren des E^3 sehr nützlich. Wir führen es zunächst rein algebraisch ein und beschäftigen uns mit der geometrischen Bedeutung später.

Zunächst soll das Kreuzprodukt zweier Vektoren $\vec{u}$ und $\vec{v}$ wieder einen *Vektor* ergeben: $\vec{w} = \vec{u} \times \vec{v}$, und zwar soll $\vec{w}$ sowohl auf $\vec{u}$ als auch auf $\vec{v}$ senkrecht stehen, d.h. es gilt

$$\vec{w} \cdot \vec{u} = 0, \quad \vec{w} \cdot \vec{v} = 0. \quad (2.4.1)$$

Außerdem soll das Kreuzprodukt linear in beiden Argumenten sein, d.h.

$$(\lambda_1 \vec{u}_1 + \lambda_2 \vec{u}_2) \times \vec{v} = \lambda_1 \vec{u}_1 \times \vec{v} + \lambda_2 \vec{u}_2 \times \vec{v} \quad (2.4.2)$$

und analog für das zweite Argument. Schließlich soll das Vektorprodukt **antisymmetrisch** sein, d.h.

$$\vec{u} \times \vec{v} = -\vec{v} \times \vec{u}. \quad (2.4.3)$$

Es ist also die *Reihenfolge* der Argumente im Vektorprodukt wichtig, und wenn man beim Rechnen diese Reihenfolge umkehrt, muss man den Vorzeichenwechsel sorgfältig beachten.

Insbesondere folgt aus (2.4.3), dass das Kreuzprodukt eines Vektors mit sich selbst verschwindet, denn vertauscht man die beiden gleichen Vektoren, ändert sich einerseits gar nichts, aber andererseits gilt eben (2.4.3) und folglich

$$\vec{u} \times \vec{u} = -\vec{u} \times \vec{u} \Rightarrow \vec{u} \times \vec{u} = 0. \quad (2.4.4)$$

2.4. Das Vektorprodukt

Kommen wir nun zur geometrischen Bedeutung des Vektorprodukts. Dazu betrachten wir eine kartesische Basis $(\vec{e}_1, \vec{e}_2, \vec{e}_3)$. Es liegen also drei Vektoren der Länge 1 vor, die paarweise aufeinander senkrecht stehen. Aus der Definition des Vektorprodukts folgt, dass $\vec{e}_1 \times \vec{e}_2 = \pm \vec{e}_3$ ist, denn $\pm \vec{e}_3$ sind offenbar die einzigen Einheitsvektoren, die auf *beiden* Vektoren $\vec{e}_1$ und $\vec{e}_2$ senkrecht stehen. Es ist klar, dass der Fall $\vec{e}_1 \times \vec{e}_2 = +\vec{e}_3$ besonders bequem ist.

Geometrisch anschaulich wird diese Uneindeutigkeit des Vorzeichens wie folgt erklärt:

Wir vereinbaren willkürlich, dass das positive Vorzeichen gilt, wenn die drei kartesischen Basisvektoren gemäß der **Rechte-Hand-Regel** (RHR) orientiert sind, d.h. richtet man den Daumen der rechten Hand in Richtung von $\vec{e}_1$, den Zeigefinger in Richtung von $\vec{e}_2$, muss der Mittelfinger in Richtung von $\vec{e}_3$ weisen. Dann legen wir fest, dass $\vec{e}_1 \times \vec{e}_2 = +\vec{e}_3$ ist. Man nennt solche kartesischen Basen entsprechend **rechtshändige kartesische Basen**.

Man macht sich leicht anschaulich klar, dass im Fall, dass $\vec{e}_3$ in die andere Richtung weist, die entsprechende **Linke-Hand-Regel** gilt. In der theoretischen Physik benutzen wir aber vereinbarungsgemäß ausschließlich rechtshändige Basen, weil alles andere zu viel Verwirrung Anlass geben kann.

Es ist anschaulich auch klar, dass für rechtshändige Basen die Formel $\vec{e}_1 \times \vec{e}_2 = \vec{e}_3$ auch unter **zyklischer Vertauschung** der Indizes gilt, d.h. wir haben die drei Gleichungen

$$\vec{e}_1 \times \vec{e}_2 = \vec{e}_3, \quad \vec{e}_2 \times \vec{e}_3 = \vec{e}_1, \quad \vec{e}_3 \times \vec{e}_1 = \vec{e}_2. \quad (2.4.5)$$

Damit können wir nun die Vektorprodukte beliebiger Vektoren durch ihre Komponenten ausdrücken:

$$\vec{u} \times \vec{v} = (u_1 \vec{e}_1 + u_2 \vec{e}_2 + u_3 \vec{e}_3) \times (v_1 \vec{e}_1 + v_2 \vec{e}_2 + v_3 \vec{e}_3). \quad (2.4.6)$$

Ausmultiplizieren und Anwendung von (2.4.5) ergibt unter Berücksichtigung der Antisymmetrie des Kreuzprodukts

$$\vec{u} \times \vec{v} = \vec{e}_1(u_2 v_3 - u_3 v_2) + \vec{e}_2(u_3 v_1 - u_1 v_3) + \vec{e}_3(u_1 v_2 - u_2 v_1). \quad (2.4.7)$$

Schreibt man dies mithilfe der entsprechenden $\mathbb{R}^3$ -Spaltenvektoren, folgt

$$\underline{u} \times \underline{v} = \begin{pmatrix} u_1 \\ u_2 \\ u_3 \end{pmatrix} \times \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} = \begin{pmatrix} u_2 v_3 - u_3 v_2 \\ u_3 v_1 - u_1 v_3 \\ u_1 v_2 - u_2 v_1 \end{pmatrix}. \quad (2.4.8)$$

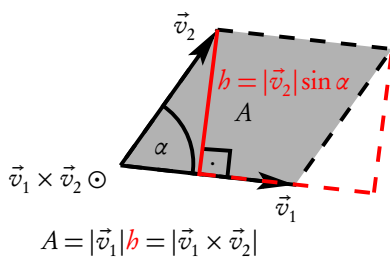
Nun werden oft auch Formeln benötigt, in denen mehrere Skalar- und Vektorprodukte vorkommen. Mit Hilfe der Formel (2.4.8) für kartesische Komponenten kann man durch einfaches Nachrechnen, was lediglich etwas Fleißarbeit erfordert, zeigen, dass für drei beliebige Vektoren stets

$$\vec{u} \cdot (\vec{v} \times \vec{w}) = (\vec{u} \times \vec{v}) \cdot \vec{w} \quad (2.4.9)$$

und

$$\vec{u} \times (\vec{v} \times \vec{w}) = \vec{v}(\vec{u} \cdot \vec{w}) - \vec{w}(\vec{u} \cdot \vec{v}) \quad (2.4.10)$$

gelten (*Übung!*).



Nun können wir eine weitere wichtige geometrische Eigenschaft des Vektorprodukts zeigen. Wir rechnen dazu die Länge des Vektorprodukts zweier Vektoren bzw. dessen Quadrat aus. Dabei wenden wir nacheinander (2.4.9) und (2.4.10) an:

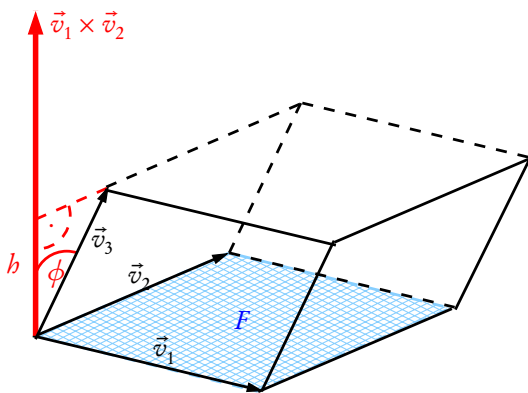
$$\begin{aligned}
 |\vec{v}_1 \times \vec{v}_2|^2 &= (\vec{v}_1 \times \vec{v}_2) \cdot (\vec{v}_1 \times \vec{v}_2) \\
 &= [(\vec{v}_1 \times \vec{v}_2) \times \vec{v}_1] \cdot \vec{v}_2 = [\vec{v}_2(\vec{v}_1 \cdot \vec{v}_1) - \vec{v}_1(\vec{v}_1 \cdot \vec{v}_2)] \cdot \vec{v}_2 \\
 &= |\vec{v}_2|^2 |\vec{v}_1|^2 - (\vec{v}_1 \cdot \vec{v}_2)^2 \\
 &= |\vec{v}_1|^2 |\vec{v}_2|^2 [1 - \cos^2 \angle(\vec{v}_1, \vec{v}_2)] \\
 &= |\vec{v}_1|^2 |\vec{v}_2|^2 \sin^2 \angle(\vec{v}_1, \vec{v}_2).
 \end{aligned} \tag{2.4.11}$$

Da $\angle(\vec{v}_1, \vec{v}_2) \in [0, \pi]$ ist, ist $\sin \angle(\vec{v}_1, \vec{v}_2) \geq 0$ und damit

$$|\vec{v}_1 \times \vec{v}_2| = |\vec{v}_1| |\vec{v}_2| \sin \angle(\vec{v}_1, \vec{v}_2). \tag{2.4.12}$$

Anhand der nebenstehenden Skizze ergibt sich, dass das Kreuzprodukt $\vec{v}_1 \times \vec{v}_2$ vom Betrag her der **Fläche** des von den beiden Vektoren aufgespannten **Parallelogramms** entspricht.

2.5 Das Spatprodukt



$$\text{vol}(\vec{v}_1, \vec{v}_2, \vec{v}_3) = (\vec{v}_1 \times \vec{v}_2) \cdot \vec{v}_3 = \pm h F$$

Es ist klar, dass wir nun auch kombinierte Produkte bilden können. Von besonderer Bedeutung ist das **Spatprodukt** aus drei Vektoren $(\vec{v}_1 \times \vec{v}_2) \cdot \vec{v}_3$. Seine geometrische Bedeutung wird aus der nebenstehenden Zeichnung deutlich. Dazu bemerken wir, dass

$$\begin{aligned}
 \text{vol}(\vec{v}_1, \vec{v}_2, \vec{v}_3) &= (\vec{v}_1 \times \vec{v}_2) \cdot \vec{v}_3 \\
 &= |\vec{v}_1 \times \vec{v}_2| |\vec{v}_3| \cos \phi
 \end{aligned} \tag{2.5.1}$$

mit $\cos \phi$, $\phi = \angle(\vec{v}_1 \times \vec{v}_2, \vec{v}_3)$. Aus der Zeichnung wird deutlich, dass der Betrag das Volumen des durch die drei Vektoren

aufgespannten **Parallelepipeds** oder **Spats** ist. Das Spatprodukt ist offenbar positiv, wenn die drei Vektoren eine rechtshändige und negativ wenn sie eine linkshändige Basis bilden.

Falls die Vektoren linear abhängig sind, d.h. liegen die Vektoren alle in einer Ebene oder sind sogar alle parallel zueinander, verschwindet das Spatprodukt. Sind nämlich die drei Vektoren linear abhängig, gibt es Zahlen λ_1 und λ_2 , so dass $\vec{v}_3 = \lambda_1 \vec{v}_1 + \lambda_2 \vec{v}_2$. Nun ist das Vektorprodukt $\vec{v}_1 \times \vec{v}_2$ ein sowohl zu $\vec{v}_1$ als auch zu $\vec{v}_2$ senkrechter Vektor, und das Skalarprodukt mit $\vec{v}_3$ verschwindet demnach. Nehmen wir umgekehrt an, dass das Spatprodukt der drei Vektoren verschwindet, bedeutet dies, dass $\vec{v}_3$ senkrecht auf $\vec{v}_1 \times \vec{v}_2$ steht, und das besagt, dass $\vec{v}_3$ in der von $\vec{v}_1$ und $\vec{v}_2$ aufgespannten Ebene liegt und also wieder $\vec{v}_3 = \lambda_1 \vec{v}_1 + \lambda_2 \vec{v}_2$ gilt.

Daraus folgt, dass drei Vektoren im E^3 dann und nur dann linear abhängig sind, wenn das Spatprodukt verschwindet.

2.5. Das Spatprodukt

Allgemein wichtig für das Rechnen mit Vektorkomponenten bzgl. rechtshändiger kartesischer Basen ist noch das **Levi-Civita-Symbol**. Es wird durch das Spatprodukt beliebiger dreier rechtshändiger kartesischer Basisvektoren definiert:

$$\epsilon_{ijk} = (\vec{e}_i \times \vec{e}_j) \cdot \vec{e}_k. \quad (2.5.2)$$

Wir zeigen nun, dass dieses Symbol sein Vorzeichen ändert, wenn wir zwei beliebige Indizes vertauschen. Da das Vektorprodukt antikommutativ ist, trifft dies sicher auf die beiden ersten Indizes zu:

$$\epsilon_{jik} = (\vec{e}_j \times \vec{e}_i) \cdot \vec{e}_k = -(\vec{e}_i \times \vec{e}_j) \cdot \vec{e}_k = -\epsilon_{ijk}. \quad (2.5.3)$$

Mit (2.4.9) folgt

$$\epsilon_{ikj} = (\vec{e}_i \times \vec{e}_k) \cdot \vec{e}_j = \vec{e}_i \cdot (\vec{e}_k \times \vec{e}_j) = -\vec{e}_i \cdot (\vec{e}_j \times \vec{e}_k) = -(\vec{e}_i \times \vec{e}_j) \cdot \vec{e}_k = -\epsilon_{ijk}. \quad (2.5.4)$$

Die letzte Möglichkeit ist noch, den ersten mit dem dritten Index zu vertauschen. Um zu zeigen, dass auch dabei einfach das Vorzeichen wechselt, benötigen wir aber nur (2.5.3) und (2.5.4):

$$\epsilon_{kji} = -\epsilon_{kij} = +\epsilon_{ikj} = -\epsilon_{ijk}. \quad (2.5.5)$$

Da nun

$$\epsilon_{123} = (\vec{e}_1 \times \vec{e}_2) \cdot \vec{e}_3 = \vec{e}_3 \cdot \vec{e}_3 = 1 \quad (2.5.6)$$

gilt, ist schließlich

$$\epsilon_{123} = \epsilon_{312} = \epsilon_{231} = 1, \quad \epsilon_{213} = \epsilon_{321} = \epsilon_{132} = -1 \quad (2.5.7)$$

und $\epsilon_{ijk} = 0$ falls wenigstens zwei Indizes gleich sind.

Wir können nun mit Hilfe des Levi-Civita-Symbols die im vorigen Abschnitt in der Spaltennotation geschriebenen Formeln mit Hilfe des Indexkalküls formulieren. Es sei wieder

$$\vec{w} = \vec{u} \times \vec{v}. \quad (2.5.8)$$

Dann gilt

$$\begin{aligned} w_i &= \vec{e}_i \cdot \vec{w} = \vec{e}_i \cdot (\vec{u} \times \vec{v}) \\ &= \vec{e}_i \cdot \sum_{j,k=1}^3 u_j v_k (\vec{e}_j \times \vec{e}_k) \\ &= \sum_{j,k=1}^3 u_j v_k \vec{e}_i \cdot (\vec{e}_j \times \vec{e}_k) \\ &= \sum_{j,k=1}^3 u_j v_k (\vec{e}_i \times \vec{e}_j) \cdot \vec{e}_k \\ &= \sum_{j,k=1}^3 \epsilon_{ijk} u_j v_k. \end{aligned} \quad (2.5.9)$$

Daraus folgt für das Spatprodukt

$$\vec{t} \cdot (\vec{u} \times \vec{v}) = \sum_{i=1}^3 t_i \vec{e}_i \cdot (\vec{u} \times \vec{v}) = \sum_{i=1}^3 \epsilon_{ijk} t_i u_j v_k. \quad (2.5.10)$$

Wichtig ist auch die folgenden Summenformel für zwei Levi-Civita-Symbole

$$\begin{aligned}\sum_{k=1}^3 \epsilon_{ijk} \epsilon_{klm} &= \sum_{k=1}^3 [(\vec{e}_i \times \vec{e}_j) \cdot \vec{e}_k][(\vec{e}_k \times \vec{e}_l) \cdot \vec{e}_m] \\ &= \sum_{k=1}^3 [(\vec{e}_i \times \vec{e}_j) \cdot \vec{e}_k][\vec{e}_k \cdot (\vec{e}_l \times \vec{e}_m)] \\ &= (\vec{e}_i \times \vec{e}_j) \cdot (\vec{e}_l \times \vec{e}_m).\end{aligned}\tag{2.5.11}$$

Dabei haben wir im letzten Schritt verwendet, dass für beliebigen Vektoren

$$\vec{a} \cdot \vec{b} = \sum_{k=1}^3 a_k b_k = \sum_{k=1}^3 (\vec{a} \cdot \vec{e}_k)(\vec{e}_k \cdot \vec{b})\tag{2.5.12}$$

gilt. Nun können wir (2.5.11) weiter vereinfachen

$$\begin{aligned}\sum_{k=1}^3 \epsilon_{ijk} \epsilon_{klm} &\stackrel{(2.4.9)}{=} \vec{e}_i \cdot [\vec{e}_j \times (\vec{e}_l \times \vec{e}_m)] \\ &\stackrel{(2.4.10)}{=} \vec{e}_i \cdot [\vec{e}_l(\vec{e}_j \cdot \vec{e}_m) - \vec{e}_m(\vec{e}_j \cdot \vec{e}_l)] \\ &= \delta_{il} \delta_{jm} - \delta_{im} \delta_{jl}.\end{aligned}\tag{2.5.13}$$

Dies können wir wegen $\epsilon_{klm} = +\epsilon_{lmk}$ auch in der Form

$$\sum_{k=1}^3 \epsilon_{ijk} \epsilon_{lmk} = \delta_{il} \delta_{jm} - \delta_{im} \delta_{jl}\tag{2.5.14}$$

schreiben. Setzen wir in dieser Formel $m = j$ und summieren dann über j , ergibt sich weiter

$$\sum_{j,k=1}^3 \epsilon_{ijk} \epsilon_{ljk} = \sum_{j=1}^3 (\delta_{il} \delta_{jj} - \delta_{ij} \delta_{jl}) = 3\delta_{il} - \delta_{il} = 2\delta_{il}.\tag{2.5.15}$$

Setzen wir auch noch $l = i$ und summieren dann über i , finden wir schließlich

$$\sum_{i,j,k=1}^3 \epsilon_{ijk} \epsilon_{ijk} = 2 \sum_{i=1}^3 \delta_{ii} = 2 \cdot 3 = 6.\tag{2.5.16}$$

2.6 Lineare Gleichungssysteme und Determinanten

Wir verlassen für die nächsten Abschnitte den dreidimensionalen Euklidischen Vektorraum und betrachten das Problem, **lineare Gleichungssysteme** zu lösen, was uns auf die wichtigsten Begriffe der **Matrizenrechnung** führt.

2.6.1 Lineare Gleichungssysteme

Ein **lineares Gleichungssystem** stellt uns vor die Aufgabe, n Gleichungen der Form

$$\begin{aligned}A_{11}x_1 + A_{12}x_2 + \cdots + A_{1n}x_n &= y_1, \\ A_{21}x_1 + A_{22}x_2 + \cdots + A_{2n}x_n &= y_2, \\ &\vdots \\ A_{n1}x_1 + A_{n2}x_2 + \cdots + A_{nn}x_n &= y_n\end{aligned}\tag{2.6.1}$$

zu lösen. Dabei sind die $A_{jk} \in \mathbb{R}$ und $y_j \in \mathbb{R}$ gegeben und die x_j gesucht.

Wir können diese Gleichung kurz mit Hilfe der Matrix-Vektor-Schreibweise formulieren:

$$\hat{A}\underline{x} = \underline{y}, \quad (2.6.2)$$

wobei $\underline{x}, \underline{y} \in \mathbb{R}^n$ Spaltenvektoren mit n reellen Komponenten sind und $\hat{A} \in \mathbb{R}^{n \times n}$ die **quadratische Matrix**

$$\hat{A} = \begin{pmatrix} A_{11} & A_{12} & \dots & A_{1n} \\ A_{21} & A_{22} & \dots & A_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ A_{n1} & A_{n2} & \dots & A_{nn} \end{pmatrix} \quad (2.6.3)$$

bezeichnen.

Wir betrachten nun zunächst einige einfache nichttriviale Beispiele für den Fall $n = 2$. Beginnen wir mit dem Gleichungssystem

$$\begin{aligned} 3x_1 + x_2 &= 5, \\ x_1 - 3x_2 &= 7. \end{aligned} \quad (2.6.4)$$

Man kann in solch einem einfachen Fall auf vielerlei Arten vorgehen, um das Gleichungssystem zu lösen. Wir wollen hier als systematisches Verfahren den sog. **Gauß-Algorithmus** (Carl-Friedrich Gauß, 1777-1855) betrachten. Die grundlegende Strategie ist dabei, durch lineare Operationen an den einzelnen Gleichungen eine Art **Dreiecksform** des Gleichungssystems zu erzielen, so dass man leicht die unbestimmten Variablen nacheinander bestimmen kann. In dem hier vorliegenden zweidimensionalen Fall benötigen wir dazu nur zwei einfache Schritte. Zuerst multiplizieren wir die zweite Gleichung mit 3. Das ergibt

$$\begin{aligned} 3x_1 + x_2 &= 5, \\ 3x_1 - 9x_2 &= 21. \end{aligned} \quad (2.6.5)$$

Diese Strategie hat dazu geführt, dass der Koeffizient vor x_1 in beiden Gleichungen gleich geworden ist. Wir behalten nun die erste Gleichung aus (2.6.4) bei und ersetzen die zweite Gleichung durch die Gleichung, die entsteht, wenn man von ihr die erste Gleichung in (2.6.5) abzieht. Dadurch eliminieren wir x_1 aus dieser Gleichung. Es entsteht

$$\begin{aligned} 3x_1 + x_2 &= 5, \\ -10x_2 &= 16. \end{aligned} \quad (2.6.6)$$

Dividieren wir die letzte Gleichung durch (-10) , erhalten wir bereits die Lösung $x_2 = -16/10 = -8/5 = -1,6$.

Nun müssen wir diese Lösung nur in die erste Gleichung einzusetzen und können dann leicht nach x_1 auflösen:

$$3x_1 - 8/5 = 5 \Rightarrow 3x_1 = 5 + 8/5 = 33/5 \Rightarrow x_1 = 11/5 = 2,2. \quad (2.6.7)$$

In dem Fall erhalten wir also eine **eindeutige Lösung**, nämlich $x_1 = -8/5$, $x_2 = 11/5$.

Es ist nun offenbar unnötig, bei diesen Rechnungen beständig die unbekannten Variablen mitzuschreiben. Stattdessen können wir einfach die Matrix aus den Koeffizienten

$$\hat{A} = \begin{pmatrix} 3 & 1 \\ 1 & -3 \end{pmatrix} \quad (2.6.8)$$

und den Spaltenvektor

$$\underline{y} = \begin{pmatrix} 5 \\ 7 \end{pmatrix} \quad (2.6.9)$$

zu der $\mathbb{R}^{n \times n+1}$ -Matrix

$$\begin{pmatrix} 3 & 1 & 5 \\ 1 & -3 & 7 \end{pmatrix} \quad (2.6.10)$$

ergänzen und mittels linearer Manipulationen mit Zeilen auf die einfache Form bringen, so dass für die Teilmatrix $\hat{A}$ zunächst nur eine **obere Dreiecksmatrix** übrig bleibt. Dies geschieht genau wie eben mit den Gleichungen demonstriert, indem man die obere Zeile von der mit 3 multiplizierten unteren Zeile abzieht. Daraus entsteht die Matrix

$$\begin{pmatrix} 3 & 1 & 5 \\ 0 & -10 & 16 \end{pmatrix}. \quad (2.6.11)$$

Wir können nun auch systematisch noch die verbliebenen Außerdiagonalelemente der Untermatrix $\hat{A}$ eliminieren, indem wir geeignete Vielfache bilden und die weiter unten stehende Zeile geeignet abziehen. Hier müssen wir zum 10-fachen der ersten Zeile nur die zweite Zeile addieren und danach schließlich die erste Zeile noch durch 30 und die zweite durch -10 teilen. Daraus entsteht nacheinander

$$\begin{pmatrix} 30 & 0 & 66 \\ 0 & -10 & 16 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & 2,2 \\ 0 & 1 & -1,6 \end{pmatrix}. \quad (2.6.12)$$

Diese abkürzende Schreibweise bedeutet nun einfach, dass das Gleichungssystem zu dem Gleichungssystem identisch ist, dass wir die beiden ersten Spalten als neue Koeffizientenmatrix $\hat{A}' = \mathbb{I}_2$ auffassen und die verbliebene dritte Spalte als neuen Vektor $\underline{y}'$. Dann ist aber $\hat{A}'\underline{x} = \mathbb{I}_2\underline{x} = \underline{x} = \underline{y}'$. Damit können wir das Ergebnis $x_1 = -2,2$ und $x_2 = -1,6$ sofort an dieser Matrix ablesen!

Wir können auf diese Art auch die **Umkehrmatrix** der vorgegebenen Koeffizientenmatrix A finden. Wir müssen nur die beiden Gleichungen $\hat{A}\bar{x}_1 = \underline{e}_1 := (1,0)^T$ und $\hat{A}\bar{x}_2 = \underline{e}_2 := (0,1)^T$ lösen. Dann bilden die Spalten $\bar{x}_1$ und $\bar{x}_2$ offensichtlich die Umkehrmatrix, denn dann ist

$$\hat{A}(\bar{x}_1, \bar{x}_2) = (\hat{A}\bar{x}_1, \hat{A}\bar{x}_2) = (\underline{e}_1, \underline{e}_2) = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \mathbb{I}_2. \quad (2.6.13)$$

Statt die beiden Gleichungssysteme einzeln zu lösen, können wir dies wieder simultan erledigen, indem wir zunächst neben die Koeffizientenmatrix $\hat{A}$ die Einheitsmatrix $\mathbb{I}_2$ schreiben. Daraus ergibt sich in unserem Beispiel die $\mathbb{R}^{n \times 2n}$ -Matrix und ausgehend davon die folgende Rechensequenz (die Manipulationen sind wieder dieselben wie eben bei der Lösung des linearen Gleichungssystems)

$$\begin{pmatrix} 3 & 1 & 1 & 0 \\ 1 & -3 & 0 & 1 \end{pmatrix} \rightarrow \begin{pmatrix} 3 & 1 & 1 & 0 \\ 0 & -10 & -1 & 3 \end{pmatrix} \rightarrow \begin{pmatrix} 30 & 0 & 9 & 3 \\ 0 & -10 & -1 & 3 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & 3/10 & 1/10 \\ 0 & 1 & 1/10 & -3/10 \end{pmatrix} \quad (2.6.14)$$

Die zweite Hälfte dieser Matrix ist demnach die gesuchte Umkehrmatrix

$$\hat{A}^{-1} = \begin{pmatrix} 3/10 & 1/10 \\ 1/10 & -3/10 \end{pmatrix}. \quad (2.6.15)$$

In der Tat rechnet man nach, dass

$$\begin{pmatrix} 3 & 1 \\ 1 & -3 \end{pmatrix} \begin{pmatrix} 3/10 & 1/10 \\ 1/10 & -3/10 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \mathbb{I}_2 \quad (2.6.16)$$

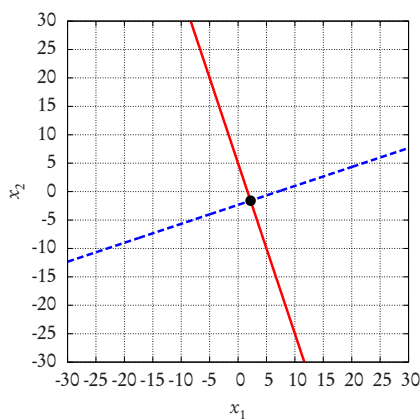
ist, wie erwartet.

Wir bemerken noch, dass für zwei beliebige Matrizen $\hat{A}, \hat{B} \in \mathbb{R}^{n \times n}$ aus $\hat{A}\hat{B} = \mathbb{1}_n$ auch $\hat{B}\hat{A} = \mathbb{1}_n$ folgt. Demnach ist in diesem Fall $\hat{A}$ invertierbar, und es ist $\hat{B} = \hat{A}^{-1}$. Dabei definieren wir $\hat{A}^{-1}$ zunächst als **linksinverses Element** zu $\hat{A}$, d.h. es gilt $\hat{A}^{-1}\hat{A} = \mathbb{1}_n$. In der Tat folgt aus $\hat{A}\hat{B} = \mathbb{1}_n$ durch Multiplizieren mit $\hat{A}^{-1}$ von links $\hat{A}^{-1}(\hat{A}\hat{B}) = (\hat{A}^{-1}\hat{A})\hat{B} = \hat{B} = \hat{A}^{-1}$. Damit ist aber in der Tat auch $\hat{B}\hat{A} = \hat{A}^{-1}\hat{A} = \mathbb{1}_n$. Obwohl also die Matrizenmultiplikation i.a. *nicht kommutativ* ist, gilt für die Umkehrmatrix von $\hat{A}$, falls sie existiert, dass sie sowohl das **links- als auch das rechtsinverse Element** ist, d.h. es ist stets $\hat{A}^{-1}\hat{A} = \hat{A}\hat{A}^{-1} = \mathbb{1}_n$.

Schließlich bemerken wir, dass wenn $\hat{A}$ invertierbar ist, also eine inverse Matrix $\hat{A}^{-1}$ existiert, so ist auch diese Matrix ihrerseits invertierbar, und es gilt $(\hat{A}^{-1})^{-1} = \hat{A}$. In der Tat finden wir durch Multiplikation der Gleichung $(\hat{A}^{-1})^{-1}\hat{A}^{-1} = \mathbb{1}$ von rechts mit $\hat{A}$ sofort $[(\hat{A}^{-1})^{-1}\hat{A}^{-1}]\hat{A} = (\hat{A}^{-1})^{-1}(\hat{A}^{-1}\hat{A}) = (\hat{A}^{-1})^{-1}\mathbb{1}_n = (\hat{A}^{-1})^{-1} = \mathbb{1}_n\hat{A} = \hat{A}$.

Betrachten wir nun nochmals das Gleichungssystem (2.6.5). Zunächst stellen wir fest, dass in dem gegebenen Fall das Gleichungssystem für jeden Vektor $\underline{y}$ genau eine Lösung hat, weil nach unserer obigen Rechnung die Koeffizientenmatrix $\hat{A}$ eine Inverse besitzt, denn es gilt ja

$$\hat{A}\underline{x} = \underline{y} \Leftrightarrow \underline{x} = \hat{A}^{-1}\underline{y}. \quad (2.6.17)$$



Jetzt wollen wir die Gleichungen geometrisch deuten. Dazu stellen wir uns vor, dass $\underline{x}$ die kartesischen Komponenten eines geometrischen Vektors sind. Dann bedeutet die erste Gleichung

$$3x_1 + x_2 = 5 \Rightarrow x_2 = -3x_1 + 5, \quad (2.6.18)$$

dass die diese Gleichungen erfüllenden Vektoren auf einer **Geraden** liegen, die durch diese Gleichung bestimmt ist. Dabei bedeutet der Koeffizient (-3) vor x_1 auf der rechten Seite, dass die Gerade die Steigung (-3) besitzt (also monoton von links oben nach rechts unten) in der Zeichenebene fällt und die Konstante 5, dass die Gerade durch den Punkt $(0; 5)$ verläuft, d.h. die x_2 -Achse bei $x_2 = 5$ schneidet (vgl. die rote Gerade in der nebenstehenden Skizze). Ebenso lässt sich die zweite Gleichung als Gerade in der Ebene interpretieren (blau gestrichelte Gerade):

$$x_1 - 3x_2 = 7 \Rightarrow x_2 = \frac{1}{3}x_1 - \frac{7}{3}. \quad (2.6.19)$$

Offenbar bestimmt die Lösung des Gleichungssystems den **Schnittpunkt** dieser beiden Geraden. Nun können sich in der Tat zwei Geraden in der Ebene (und auch im Raum) maximal einmal schneiden, d.h. wenn überhaupt ein Schnittpunkt existiert, dann genau einer. Im vorliegenden Fall schneiden sich die beiden Geraden tatsächlich in dem einen Punkt mit den oben berechneten Koordinaten $x_1 = 1,6$, $x_2 = -2,2$.

Nun kann es aber sein, dass die Geraden parallel verlaufen und sich nie schneiden. Dann ist die eine Zeile der Koeffizientenmatrix ein Vielfaches der ersten Zeile, denn genau dann ergibt sich für die durch sie im obigen Sinne definierten Geraden die gleiche Steigung. Es darf dann aber nicht y_2 dasselbe Vielfache von y_1 sein. Dann besitzt das Gleichungssystem nämlich offenbar keine Lösung, und die Geraden sind Parallelen. In der Tat besitzen die entsprechenden Geraden dieselbe Steigung aber unterschiedliche Schnittpunkte mit der x_2 -Achse. Ist hingegen die eine Gleichung einfach ein Vielfaches der anderen Gleichung, so definieren offensichtlich beide Gleichungen dieselbe Gerade, d.h. sie liegen aufeinander, und das lineare Gleichungssystem besitzt dann beliebig viele Lösungen, die graphisch gerade durch diese Gleichung bestimmt sind. Der Fall, dass die beiden Zeilenvektoren in der Koeffizientenmatrix linear abhängig sind, führt dazu, dass im Gauß-Algorithmus die zweite Zeile bereits im ersten Schritt identisch verschwindet, und wir daher keine Umkehrmatrix finden können, denn wir können ja nicht mehr mit Hilfe der unteren Zeile das obere Außerdiagonalelement

eliminieren. Für den Fall, dass wir ein konkretes Gleichungssystem lösen, ergibt sich entweder nach dem ersten Eliminationsschritt auch das Element in der 3. Spalte zu 0 (dann ergeben sich unendlich viele Lösungen, die im Graphen auf einer Geraden liegen) oder die letzte Spalte verschwindet oder das Element in der 3. Spalte ergibt sich zu $\neq 0$ (dann gibt es keine Lösung für das Gleichungssystem, und die Geraden im Graphen liegen parallel zueinander und schneiden sich demnach nicht).

Dieses Lösungsverhalten verallgemeinert sich auch auf höhere Dimensionen. Wir können dies mit Hilfe des Gauß-Algorithmus' stets im Einzelfall untersuchen, und wir wollen daher darauf nicht genauer eingehen. Betrachten wir stattdessen noch ein dreidimensionales Beispiel zur Lösung eines linearen Gleichungssystems mit drei Unbekannten:

$$\begin{aligned}x_1 + 2x_2 + 5x_3 &= 12, \\2x_1 + 4x_2 + 2x_3 &= 16, \\3x_1 + 10x_2 + x_3 &= 30.\end{aligned}\tag{2.6.20}$$

Als erstes schreiben wir die erweiterte Matrix für dieses Gleichungssystem auf:

$$\begin{pmatrix} 1 & 2 & 5 & 12 \\ 2 & 4 & 2 & 16 \\ 3 & 10 & 1 & 30 \end{pmatrix}.\tag{2.6.21}$$

Wir können uns die Arbeit etwas erleichtern, wenn wir die zweite Zeile durch 2 dividieren:

$$\begin{pmatrix} 1 & 2 & 5 & 12 \\ 1 & 2 & 1 & 8 \\ 3 & 10 & 1 & 30 \end{pmatrix}.\tag{2.6.22}$$

Jetzt beginnen wir mit den Gauß-Eliminationen, um die Koeffizientenmatrix in obere Dreiecksgestalt zu bringen. Wir starten mit der Elimination der ersten Spalte der beiden letzten Zeilen, indem wir die erste Zeile bzw. die dreifache erste Zeile subtrahieren:

$$\begin{pmatrix} 1 & 2 & 5 & 12 \\ 0 & 0 & -4 & -4 \\ 0 & 4 & -14 & -6 \end{pmatrix}.\tag{2.6.23}$$

Beim Lösen linearer Gleichungen (und auch beim Invertieren von Matrizen) dürfen wir auch Zeilen beliebig vertauschen, weil dies nur der Vertauschung der entsprechenden Gleichungen innerhalb des Gleichungssystems entspricht. Nachdem wir die zweite Zeile durch -4 dividiert haben, vertauschen wir das Resultat mit der letzten Zeile. Dann ergibt sich (nachdem wir auch die dann mittlere Zeile noch durch 2 dividiert haben

$$\begin{pmatrix} 1 & 2 & 5 & 12 \\ 0 & 2 & -7 & -3 \\ 0 & 0 & 1 & 1 \end{pmatrix}.\tag{2.6.24}$$

Damit ist die Koeffizientenmatrix bereits in oberer Dreiecksgestalt. Wir wollen nun auch die oberen Außerdiagonalelemente eliminieren. Dazu ziehen wir erst die zweite von der ersten Zeile ab:

$$\begin{pmatrix} 1 & 0 & 12 & 15 \\ 0 & 2 & -7 & -3 \\ 0 & 0 & 1 & 1 \end{pmatrix}.\tag{2.6.25}$$

Schließlich addieren wir das 7-fache der letzten Zeile zur zweiten und dividieren das Resultat durch 2. Zugleich subtrahieren das 12-fache der letzten Zeile von der ersten Zeile:

$$\begin{pmatrix} 1 & 0 & 0 & 3 \\ 0 & 1 & 0 & 2 \\ 0 & 0 & 1 & 1 \end{pmatrix}.\tag{2.6.26}$$

Damit ist der Lösungsvektor durch die letzte Spalte gegeben: $\underline{x} = (3, 2, 1)^T$. Einsetzen in das ursprüngliche Gleichungssystem bestätigt diese Lösung. Diese Probe sollte man zur Kontrolle am Ende der Rechnung stets ausführen!

2.6.2 Determinanten als Volumenform

In diesem Abschnitt führen wir **Determinanten** ein. Wir haben bereits Beispiele für Determinanten kennengelernt, ohne dass wir dies explizit erwähnt haben. So ist bereits das in Abschnitt 2.5 eingeführte Spatprodukt ein Beispiel für eine Determinante. Gewöhnlich definiert man aber Determinanten nicht als **multilineare Abbildungen** von Vektoren, sondern als Funktionen von **quadratischen Matrizen**. Allerdings stellt das Beispiel des Spatprodukts eine schöne Motivation für Determinanten dar, und zwar gestatten sie die Verallgemeinerung des Begriffs von **orientierten Volumeninhalten** auf beliebige Dimensionen. Beim Spatprodukt haben wir gesehen, dass

$$\text{vol}(\vec{v}_1, \vec{v}_2, \vec{v}_3) = (\vec{v}_1 \times \vec{v}_2) \cdot \vec{v}_3 \quad (2.6.27)$$

vom Betrag her das Volumen des von den drei Vektoren bestimmten Parallelepipeds ergibt, und das Vorzeichen bestimmt, ob die Vektoren in der angegebenen Reihenfolge ein rechts- oder linkshändiges System von Basisvektoren darstellen, falls sie linear unabhängig sind. Falls sie linear abhängig sind, verschwindet das Spatprodukt. Betrachten wir nun die Berechnung des Spatprodukts aus den Komponenten der Vektoren bzgl. einer *kartesischen Basis* (2.5.3), also

$$\underline{v}_1 = \begin{pmatrix} v_{11} \\ v_{21} \\ v_{31} \end{pmatrix}, \quad \underline{v}_2 = \begin{pmatrix} v_{12} \\ v_{22} \\ v_{32} \end{pmatrix}, \quad \underline{v}_3 = \begin{pmatrix} v_{13} \\ v_{23} \\ v_{33} \end{pmatrix}, \quad (2.6.28)$$

erhalten wir

$$(\vec{v}_1 \times \vec{v}_2) \cdot \vec{v}_3 = \sum_{j,k,l=1}^3 v_{j1} v_{k2} v_{l3} \epsilon_{jkl}, \quad (2.6.29)$$

wird schon nahegelegt, wie dieser Begriff des orientierten Volumens auf beliebige Dimensionen zu verallgemeinern sein wird:

Wir müssen für den n -dimensionalen Euklidischen Raum $\mathbb{R}^n$ lediglich ein Levi-Civita-Symbol mit n Indizes einführen, das wie im dreidimensionalen Fall total antisymmetrisch unter Vertauschen dieser Indizes ist und für das

$$\epsilon_{12\dots n} = 1 \quad (2.6.30)$$

gilt. Dann können wir für n Vektoren die **Volumenform**

$$\text{vol}(\underline{v}_1, \dots, \underline{v}_n) = \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 j_2 \dots j_n} v_{j_1 1} v_{j_2 2} \cdots v_{j_n n} \quad (2.6.31)$$

eingeführen.

Dabei heißt dieses Konstrukt **Form**, weil es offensichtlich eine **lineare Abbildung** $V^n \rightarrow \mathbb{R}$ ist, die vollständig **antisymmetrisch** beim Vertauschen zweier beliebiger Argumente ist (*warum?*). Diese speziellen linearen Abbildungen spielen eine wichtige Rolle. Wir werden uns in dieser Vorlesung aber nur mit der Volumenform beschäftigen.

Wir können weiter

$$(\underline{v}_1, \dots, \underline{v}_n) = \hat{V} = \begin{pmatrix} v_{11} & v_{12} & \dots & v_{1n} \\ v_{21} & v_{22} & \dots & v_{2n} \\ \vdots & \vdots & \dots & \vdots \\ v_{n1} & v_{n2} & \dots & v_{nn} \end{pmatrix} \quad (2.6.32)$$

auch als $\mathbb{R}^{n \times n}$ -Matrix auffassen.

Dies führt dann auf die Definition der **Determinante einer Matrix**

$$\det \hat{V} = \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 j_2 \dots j_n} v_{j_1 1} v_{j_2 2} \dots v_{j_n n}. \quad (2.6.33)$$

Bleiben wir aber noch bei der geometrischen Bedeutung von (2.6.31). Betrachten wir zunächst den Spezialfall $n = 2$ der Ebene. Dann ist (*nachvollziehen!*)

$$\text{vol}(\underline{v}_1, \underline{v}_2) = v_{11}v_{22} - v_{12}v_{21}. \quad (2.6.34)$$

Dass dies vom Betrag her die **Fläche** des von $\vec{v}_1$ und $\vec{v}_2$ aufgespannten Parallelogramms in der Ebene ist, machen wir uns schnell klar, indem wir diese Betrachtung im dreidimensionalen Raum anstellen. Interpretieren wir also die zweidimensionalen Vektoren als kartesische Komponenten dreidimensionaler Vektoren in der 12-Ebene gemäß

$$\underline{w}_1 = \begin{pmatrix} v_{11} \\ v_{12} \\ 0 \end{pmatrix}, \quad \underline{w}_2 = \begin{pmatrix} v_{21} \\ v_{22} \\ 0 \end{pmatrix}, \quad (2.6.35)$$

erhalten wir für das Kreuzprodukt

$$\underline{w}_1 \times \underline{w}_2 = \begin{pmatrix} 0 \\ 0 \\ v_{11}v_{22} - v_{12}v_{21} \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ \text{vol}(\underline{v}_1, \underline{v}_2) \end{pmatrix}. \quad (2.6.36)$$

Nun ist das Kreuzprodukt zweier dreidimensionaler Vektoren der Länge nach der Flächeninhalt des von diesen beiden Vektoren definierten Parallelogramms, und das Vorzeichen bestimmt sich nach der Rechte-Hand-Regel. Es ist offenbar positiv, wenn man den Vektor $\vec{v}_1$ um den kleineren Drehwinkel entgegen dem Uhrzeigersinn in die Richtung von $\vec{v}_2$ drehen kann und entsprechend negativ, wenn diese Drehung im Uhrzeigersinn erfolgt. Dies definiert eine Orientierung in der Ebene.

Im n -dimensionalen Raum mit $n \geq 4$ ist es etwas schwierig, die Volumenform (2.6.31) geometrisch zu veranschaulichen. Wir können sie aber einfach als Definition des orientierten Volumens eines n -dimensionalen Parallelepipeds im n -dimensionalen Raum auffassen.

Wir befassen uns mit der Frage, wie man das Volumen eines n -dimensionalen Parallelepipeds berechnet, wenn man die entsprechenden Vektorkomponenten bzgl. einer beliebigen nicht-Kartesischen Basis gegeben hat, im nächsten Abschnitt, weil wir dazu die Eigenschaften der Determinanten von Matrizen benötigen.

2.6.3 Determinanten von Matrizen

Wir wollen nun einige Rechenregeln für die Determinante einer Matrix herleiten. Wir gehen dabei von der Definition (2.6.33) aus. Aus der Definition des Levi-Civita-Symbols ist sofort klar, dass die Determinante verschwindet, wenn die Spaltenvektoren der Matrix linear abhängig sind.

Man sieht dies für den Fall, dass zwei Spalten sogar gleich sind, sofort ein. Seien beispielsweise die ersten beiden Spalten gleich, d.h. gilt $A_{j1} = A_{j2}$ für alle $j \in \{1, 2, \dots, n\}$, so wechselt beim Vertauschen dieser Spalten

einerseits wegen $\epsilon_{j_1 j_2 j_3 \dots j_n} = -\epsilon_{j_2 j_1 j_3 \dots j_n}$ die Determinante ihr Vorzeichen. Da andererseits aber die ersten beiden Spalten gleich sind, ändert sich nichts, wenn man im ursprünglichen Ausdruck einfach die j_1 und j_2 bei den Matricelementen vertauscht, d.h. die Determinante ändert andererseits ihr Vorzeichen bei diesem Tausch nicht. Es ist also dann $\det \hat{A} = -\det \hat{A}$, und daraus folgt notwendig $\det \hat{A} = 0$.

Seien nun die Spalten der Matrix linear abhängig. Dann läßt sich eine Spalte (sagen wir die erste) als Linearkombination aller übrigen Spalten darstellen, d.h. es gibt Zahlen λ_k mit $k \in \{2, \dots, n\}$, so dass

$$A_{j_1} = \sum_{k=2}^n \lambda_k A_{j_k}. \quad (2.6.37)$$

Dann folgt

$$\det \hat{A} = \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 \dots j_n} \left(\sum_{k=2}^n \lambda_k A_{j_1 k} \right) A_{j_2} \dots A_{j_n} = \sum_{k=2}^n \lambda_k \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 \dots j_n} A_{j_1 k} A_{j_2} \dots A_{j_n}. \quad (2.6.38)$$

Im letzten Ausdruck ist aber für jedes $k \in \{2, \dots, n\}$ die innere Summe die Determinante einer Matrix, mit zwei gleichen Spalten, d.h. jeder dieser Summanden verschwindet, und damit ist gezeigt, dass tatsächlich $\det \hat{A} = 0$ ist, falls die Spalten von $\hat{A}$ linear abhängig sind.

Es ergibt sich daraus übrigens sofort, dass sich die Determinante nicht ändert, wenn man beliebige Vielfache einer Spalte zu irgendeiner anderen Spalte oder beliebige Vielfache einer Zeile zu irgendeiner anderen Zeile addiert. Durch Manipulationen wie wir sie beim Gaußalgorithmus verwendet haben, ändert sich also die Determinante der Matrix nicht. Wenn wir allerdings einfach eine Zeile oder Spalte mit einer reellen Zahl multiplizieren, multipliziert sich auch die Determinante mit dieser Zahl. Wir werden weiter unten noch sehen, wie man dies ausnutzen kann, um Determinanten effizient zu berechnen. Offensichtlich ändert sich aber das Vorzeichen der Determinante, wenn man zwei Spalten (oder Zeilen) miteinander vertauscht. Schreiben wir $\hat{A} = (\underline{v}_1, \dots, \underline{v}_n)$, gilt nämlich (z.B. beim Vertauschen der ersten und der zweiten Spalte)

$$\begin{aligned} \text{vol}(\underline{v}_2, \underline{v}_1, \underline{v}_3, \dots, \underline{v}_n) &= \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 j_2 \dots j_n} v_{j_1 2} v_{j_2 1} v_{j_3 3} \dots v_{j_n n} \\ &= - \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_2 j_1 \dots j_n} v_{j_1 2} v_{j_2 1} v_{j_3 3} \dots v_{j_n n} \\ &= -\text{vol}(\underline{v}_1, \dots, \underline{v}_n) = -\det \hat{A}. \end{aligned} \quad (2.6.39)$$

Nun betrachten wir die Determinante der **transponierten Matrix**. Es gilt ja $(\hat{A}^T)_{jk} = A_{kj}$, und damit ist

$$\det \hat{A}^T = \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 \dots j_n} (\hat{A}^T)_{j_1 1} \dots (\hat{A}^T)_{j_n 1} = \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 \dots j_n} A_{1 j_1} \dots A_{n j_n}. \quad (2.6.40)$$

Nun liefern gemäß der Definition des Levi-Civita-Symbols nur solche Terme einen von 0 verschiedenen Beitrag zur Determinante, für die $(j_1, \dots, j_n)$ eine **Permutation** (d.h. Anordnung in geänderter Reihenfolge) von $(1, 2, \dots, n)$ ist, und es wird über alle möglichen dieser Anordnungen summiert. Man kann demnach im letzten Ausdruck in (2.6.40) genauso gut die Produkte auch bzgl. der hinteren Indizes „sortieren“ und über alle Permutationen von $(1, 2, \dots, n)$ der vorderen Indizes summieren. Das ist aber gerade wieder die Determinante der ursprünglichen Matrix. Wir erhalten also

$$\det \hat{A}^T = \det \hat{A}. \quad (2.6.41)$$

2. Lineare Algebra

Betrachten wir weiter die Determinante eines Produktes zweier $\mathbb{R}^{n \times n}$ -Matrizen. Der Definition der Determinante gemäß (2.6.33) gilt

$$\det(\hat{A}\hat{B}) = \sum_{j_1, \dots, j_n=1}^n \sum_{k_1, \dots, k_n=1}^n \epsilon_{j_1 \dots j_n} A_{j_1 k_1} B_{k_1 1} A_{j_2 k_2} B_{k_2 2} \cdots A_{j_n k_n} B_{k_n n}. \quad (2.6.42)$$

Wir können nun offenbar die Summenzeichen vertauschen, weil man Summanden in Summen beliebig vertauschen darf. Die Summation über die j_i betrifft nur Matrixkomponenten der Matrix $\hat{A}$. Wir können also zugleich die Matrixelemente der Matrix $\hat{B}$ aus der Summe über die j_i ausklammern. Das führt zu

$$\det(\hat{A}\hat{B}) = \sum_{k_1, \dots, k_n=1}^n B_{k_1 1} B_{k_2 2} \cdots B_{k_n n} \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 \dots j_n} A_{j_1 k_1} A_{j_2 k_2} \cdots A_{j_n k_n}. \quad (2.6.43)$$

Betrachten wir nun die Summen über die j_i , sehen wir, dass es sich um Determinanten von Matrizen handelt, die aus den Spalten der Matrix $\hat{A}$ mit den Nummern $k_1, k_2, \dots, k_n$ bestehen. Diese Determinanten verschwinden allesamt wenn zwei oder mehr k_i 's gleich sind, d.h. es ergibt sich nur ein von 0 verschiedener Wert, wenn unter den der j_i -Summe entsprechenden Determinanten lediglich Vertauschungen der Spalten der ursprünglichen Matrix vorgenommen wurden. Vertauscht man aber zwei Spalten in einer Matrix, ändert deren Determinante nur das Vorzeichen. Insgesamt folgt aufgrund dieser Überlegung

$$\sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 \dots j_n} A_{j_1 k_1} A_{j_2 k_2} \cdots A_{j_n k_n} = \epsilon_{k_1 k_2 \dots k_n} \det \hat{A}. \quad (2.6.44)$$

Setzen wir das in (2.6.43) ein, folgt

$$\det(\hat{A}\hat{B}) = \det \hat{A} \sum_{k_1, \dots, k_n=1}^n \epsilon_{k_1 k_2 \dots k_n} B_{k_1 1} B_{k_2 2} \cdots B_{k_n n} = (\det \hat{A})(\det \hat{B}). \quad (2.6.45)$$

Wir haben also den Satz bewiesen, dass die Determinante des Produktes zweier Matrizen das Produkt der Determinanten der einzelnen Matrizen ist, d.h. $\det \hat{A}\hat{B} = (\det \hat{A})(\det \hat{B})$.

Für **invertierbare** Matrizen folgt daraus sofort, dass

$$\det(\hat{A}\hat{A}^{-1}) = (\det \hat{A})(\det \hat{A}^{-1}) = \det \mathbb{1}_n = 1 \quad (2.6.46)$$

und damit, dass die Determinante einer invertierbaren Matrix stets von 0 verschieden ist und dass

$$\det(\hat{A}^{-1}) = \frac{1}{\det \hat{A}} \quad (2.6.47)$$

gilt.

Jetzt wollen wir eine rekursive Methode zur Berechnung der Determinante herleiten, den sog. **Entwicklungssatz**, und zwar betrachten wir die Entwicklung nach der **ersten Spalte**. Dazu müssen wir in der Definition der Determinante nur die Summation über j_1 „abseparieren“:

$$\det \hat{A} = \sum_{j_1, \dots, j_n=1}^n \epsilon_{j_1 \dots j_n} A_{j_1 1} A_{j_2 2} \cdots A_{j_n n} = \sum_{j_1=1}^n A_{j_1 1} \sum_{j_2, \dots, j_n=1}^n \epsilon_{j_1 j_2 \dots j_n} A_{j_2 2} A_{j_3 3} \cdots A_{j_n n}. \quad (2.6.48)$$

Betrachten wir nun die innere Summe über $j_2, \dots, j_n$. Dies ist bis auf das Vorzeichen die Determinante der Matrix, die aus der ursprünglichen Matrix entsteht, wenn man die j_1 -te Zeile und die erste Spalte streicht. Um das Vorzeichen zu erhalten, betrachten wir die Summanden mit $j_1 = 1$ und $j_1 = 2$. Für $j_1 = 1$ kommt in der inneren Summe $\epsilon_{1j_2 \dots j_n}$ zu stehen. Wenn eines der j_k ($k \in \{2, \dots, n\}$) den Wert 1 annimmt, verschwindet das Levi-Civita-Symbol, und der Summand trägt nichts zur Summe bei, was dem Streichen der $j_1 = 1$ -sten Zeile entspricht. Andernfalls erhält man offenbar einfach die Determinante der Matrix $\hat{A}'_{11}$, also derjenigen $\mathbb{R}^{(n-1) \times (n-1)}$ -Matrix, die aus der Matrix $\hat{A}$ durch Streichen der ersten Zeile und der ersten Spalte entsteht. Für $j_1 = 2$ steht in der inneren Summe das Levi-Civita-Symbol $\epsilon_{2j_2 \dots j_n}$. Für $j_2 = 1, j_3 = 3, \dots, j_n = n$, erhält man den Faktor $\epsilon_{213 \dots n} = -1$. Man macht sich klar, dass wegen $\epsilon_{2j_2 \dots j_n} = -\epsilon_{j_2 2 j_3 \dots j_n}$ die innere Summe gerade $-\det \hat{A}'_{21}$ entspricht. Dabei ist $\hat{A}'_{21}$ die $(n-1) \times (n-1)$ -Matrix, die aus $\hat{A}$ durch Streichen der zweiten Zeile und der ersten Spalte entsteht. Für $j_1 = 3$ ergibt sich wegen $\epsilon_{3j_2 \dots j_n} = +\epsilon_{j_1 j_2 3 j_3 \dots j_n}$ hingegen wieder $+\det \hat{A}'_{31}$ usw. Insgesamt folgt also

$$\det \hat{A} = \sum_{j_1=1}^n (-1)^{j_1+1} A_{j_1 1} \det \hat{A}'_{j_1 1}. \quad (2.6.49)$$

Die Berechnung der Determinante einer $\mathbb{R}^{n \times n}$ -Matrix ist damit auf die Berechnung von Determinanten von $\mathbb{R}^{(n-1) \times (n-1)}$ -Matrizen zurückgeführt.

Man überlegt sich auf dieselbe Weise wie eben, dass man auch nach einer beliebigen anderen Spalte entwickeln kann. Die sorgfältige Analyse der Vorzeichen liefert für die **Entwicklung nach der i -ten Spalte**

$$\det \hat{A} = \sum_{j=1}^n (-1)^{j+i} A_{ji} \det \hat{A}'_{ji}. \quad (2.6.50)$$

Dabei ist $i \in \{1, 2, \dots, n\}$ beliebig aber während der Rechnung fest gewählt.

Wegen (2.6.41) folgt, dass man genauso gut nach einer beliebigen Zeile entwickeln kann, denn der Entwicklung der Determinante $\det(\hat{A}^T)$ nach der i -ten Spalte entspricht einer Entwicklung von $\det A = \det(\hat{A}^T)$ nach der i -ten Zeile:

$$\det \hat{A} = \sum_{j=1}^n (-1)^{j+i} A_{ij} \det \hat{A}'_{ij}. \quad (2.6.51)$$

Eine wichtige Folgerung daraus ist, dass in dem Fall, dass $\det \hat{A} \neq 0$ die Inverse $\hat{B} = \hat{A}^{-1}$ der Matrix $\hat{A}$ existiert, und deren Komponenten durch die Unterdeterminanten wie folgt bestimmt sind:

$$B_{jk} = (\hat{A}^{-1})_{jk} = \frac{(-1)^{j+k}}{\det \hat{A}} \det \hat{A}'_{kj}, \quad (2.6.52)$$

wobei auf die *Reihenfolge* der Indizes der Unterdeterminante zu achten ist: Auf der rechten Seite der Gleichung stehen die Indizes in der umgekehrten Reihenfolge wie auf der linken Seite!

Um diese Behauptung zu beweisen, rechnen wir dies einfach nach. Zunächst ist

$$(\hat{A}\hat{B})_{ik} = \sum_{j=1}^n A_{ij} B_{jk} = \frac{1}{\det \hat{A}} \sum_{j=1}^n (-1)^{j+k} A_{ij} \det \hat{A}'_{kj}. \quad (2.6.53)$$

Ist nun $i = k$, so ergibt sich wegen (2.6.51) auf der rechten Seite 1. Falls $i \neq k$ ist, berechnet man die Determinante der Matrix, die aus $\hat{A}$ entsteht, wenn man die k -te Zeile durch die i -te Zeile ersetzt. In dieser Matrix

sind aber zwei Zeilen gleich, und die Determinante verschwindet daher. Damit folgt

$$(\hat{A}\hat{B})_{ik} = \delta_{ik} \Rightarrow \hat{A}\hat{B} = \mathbb{1}_n, \quad (2.6.54)$$

und das war zu zeigen.

Auch die Lösung des linearen Gleichungssystems $\hat{A}\underline{x} = \underline{y}$ im Fall, dass $\hat{A}$ invertierbar ist, d.h. eine und nur eine Lösung existiert, kann man nun mit Determinanten geschlossen angeben, denn es ist dann offenbar

$$\underline{x} = \hat{A}^{-1}\underline{y}. \quad (2.6.55)$$

Schreiben wir dies in Komponenten aus und verwenden (2.6.52), ergibt sich

$$x_j = \sum_{k=1}^n (\hat{A}^{-1})_{jk} y_k = \frac{1}{\det \hat{A}} \sum_{k=1}^n (-1)^{j+k} \det \hat{A}'_{kj} y_k. \quad (2.6.56)$$

Die Summe ist aber gerade die Determinante derjenigen Matrix, die aus der Koeffizientenmatrix $\hat{A}$ entsteht, wenn man die j -te Spalte durch den Spaltenvektor $\underline{y}$ ersetzt. Schreiben wir die Spaltenvektoren dieser Matrix als

$$\underline{a}_j = \begin{pmatrix} A_{1j} \\ A_{2j} \\ \vdots \\ A_{nj} \end{pmatrix}, \quad (2.6.57)$$

so gilt also die **Cramersche Regel**

$$x_j = \frac{\text{vol}(\underline{a}_1, \dots, \underline{a}_{j-1}, \underline{y}, \underline{a}_{j+1}, \dots, \underline{a}_n)}{\det \hat{A}}, \quad (2.6.58)$$

wobei wir uns der Schreibweise (2.6.31) bedient haben. Es ist natürlich klar, dass es in der Praxis wesentlich ökonomischer ist, nicht die $(n+1)$ Determinanten zu berechnen, die man benötigt, um gemäß (2.6.58) ein lineares Gleichungssystem zu lösen sondern stattdessen das oben beschriebene Gaußsche Eliminationsverfahren anzuwenden.

Schließlich wollen wir noch das Volumen eines n -dimensionalen Parallelepipeds bzgl. beliebiger nicht notwendig kartesischer Koordinaten herleiten. Seien dazu $\underline{x}_1, \dots, \underline{x}_n$ die n Spaltenvektoren aus *kartesischen* Koordinaten der Vektoren $\vec{x}_1, \dots, \vec{x}_n$, die das Parallelepiped aufspannen. Dann ist definitionsgemäß dessen Volumen

$$V = \text{vol}(\underline{x}_1, \dots, \underline{x}_n). \quad (2.6.59)$$

Nun seien wieder die Transformationsmatrizen $\hat{T}$ und $\hat{U}$ zwischen der kartesischen Basis $\vec{e}_j$ und der beliebigen anderen Basis $\vec{b}'_j$ wie in (2.3.4). Dann gilt gemäß (2.3.5) $\underline{x}_j = \hat{U}\underline{x}'_j$. Seien nun die $\hat{X} = (\underline{x}_1, \dots, \underline{x}_n)$ und $\hat{X}' = (\underline{x}'_1, \dots, \underline{x}'_n)$ die aus den Spaltenvektoren gebildeten Matrizen, so gilt $\hat{X} = \hat{U}\hat{X}'$ und folglich wegen (2.6.45)

$$V = \det \hat{X} = \det(\hat{U}\hat{X}') = (\det \hat{U})(\det \hat{X}'). \quad (2.6.60)$$

Wegen (2.3.6) gilt nun

$$\det \hat{g}' = \det(\hat{U}^T \hat{U}) = (\det \hat{U}^T)(\det \hat{U}) = (\det \hat{U})^2. \quad (2.6.61)$$

Dabei haben wir (2.6.41) benutzt. Es gilt also

$$V = \text{sign}(\det \hat{U}) \sqrt{\det \vec{g}'} \det \hat{X}' = \text{sign}(\det \hat{U}) \sqrt{\det \vec{g}'} \text{vol}(\underline{x}'_1, \dots, \underline{x}'_n). \quad (2.6.62)$$

Das Vorzeichen der Determinante der Transformationsmatrix bestimmt also, ob die Volumenform bzgl. der neuen Vektorkomponenten das gleiche Vorzeichen besitzt wie das bzgl. der ursprünglichen kartesischen Basis. Man nennt daher Basiswechsel, für die $\det \hat{U} = 1/\det \hat{T} > 0$ ist **orientierungserhaltende Basistransformationen**. Falls $\det \hat{U} < 0$ ist, brauchen wir nur irgendwelche zwei Basisvektoren der Basis $\vec{b}'_j$ miteinander zu vertauschen, um eine zur kartesischen Basis gleichorientierte Basis zu erhalten. Im Folgenden nehmen wir stets an, dass alle Basen stets gleichartig zueinander **rechtshändig orientiert** sind.

Betrachten wir nun den für uns wichtigsten Spezialfall, dass die neue Basis $\vec{b}'_j = \vec{e}'_j$ gleichfalls eine kartesische Basis ist. Dann ist gemäß Abschnitt 2.3 die Matrix des Basiswechsels eine orthogonale Matrix, d.h. es gilt

$$\hat{U}^T = \hat{U}^{-1} \quad (2.6.63)$$

und damit $\vec{g}' = \mathbb{1}_n$. Dann folgt aus (2.6.62), dass die Volumformen bzgl. der beiden Basen wie zu erwarten übereinstimmen. Außerdem ist für orthogonale Matrizen

$$(\det \hat{U})^2 = \det(\hat{U} \hat{U}^T) = \det \mathbb{1}_n = 1 \Rightarrow \det \hat{U} = \det \hat{T} = \pm 1. \quad (2.6.64)$$

Für **orientierungserhaltende orthogonale Matrizen** ist also $\det \hat{T} = +1$.

2.6.4 Transformationsverhalten des Kreuzprodukts

Wir untersuchen nun, wie sich die Komponenten des Kreuzprodukts zweier Vektoren im dreidimensionalen Raum unter **orthogonalen Transformationen** verhalten. Von der Anschauung her erwarten wir, dass sich die Komponenten des Kreuzprodukts wie ein Vektor transformieren. Wir werden gleich sehen, dass dies mit einer kleinen Einschränkung tatsächlich zutrifft.

Gemäß (2.5.5) gilt für die Komponenten des Vektorprodukts $\vec{x} = \vec{v} \times \vec{w}$ bzgl. einer beliebigen *rechtshändigen, kartesischen* Basis

$$x_j = \sum_{k,l=1}^3 \epsilon_{jkl} v_k w_l. \quad (2.6.65)$$

Für die Komponenten bzgl. einer anderen kartesischen Basis

$$\vec{e}'_a = \sum_{b=1}^3 U_{ba} \vec{e}_b \quad (2.6.66)$$

gilt

$$\vec{v} = \sum_a v'_a \vec{e}'_a = \sum_{a,b} \vec{e}_b U_{ba} v'_a = \sum_b v_b \vec{e}_b \Rightarrow v_b = \sum_a U_{ba} v'_a \quad (2.6.67)$$

bzw. in Matrix-Vektorschreibweise für die Spaltenvektoren aus den jeweiligen Komponenten bzgl. der Basen

$$\underline{v} = \hat{U} \underline{v}' \quad \text{und} \quad \underline{w} = \hat{U} \underline{w}'. \quad (2.6.68)$$

Daraus folgt

$$\begin{aligned}
 x_j &= \sum_{k,l=1}^3 \epsilon_{jkl} (\hat{U} \underline{v}')_k (\hat{U} \underline{w}')_l \\
 &= \sum_{a,b,k,l=1}^3 \epsilon_{jkl} U_{ka} U_{lb} v'_a w'_b \\
 &= \sum_{a,b,c,k,l=1}^3 \epsilon_{ckl} \delta_{jc} U_{ka} U_{lb} v'_a w'_b \\
 &\stackrel{(*)}{=} \sum_{a,b,c,d,k,l=1}^3 \epsilon_{ckl} U_{jd} U_{cd} U_{ka} U_{lb} v'_a w'_b \\
 &= \sum_{a,b,c,d,k,l=1}^3 \epsilon_{dab} \det \hat{U} U_{jd} v'_a w'_b \\
 &= \det \hat{U} \sum_{d=1}^3 U_{jd} (\underline{v}' \times \underline{w}')_d.
 \end{aligned} \tag{2.6.69}$$

Dabei haben wir im Schritt (*) die Orthogonalität der Matrix $\hat{U}$, d.h. $\hat{U} \hat{U}^T = \mathbb{1}$, also $\sum_d U_{jd} U_{cd} = \delta_{cd}$, verwendet.

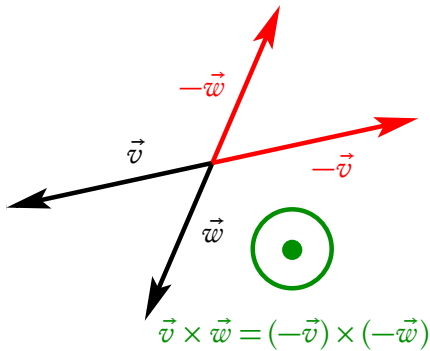
In Matrix-Vektor-Produktschreibweise bedeutet nun aber (2.6.69)

$$\underline{x} = \det \hat{U} \hat{U} \underline{x}' \Rightarrow \underline{v} \times \underline{w} = (\hat{U} \underline{v}') \times (\hat{U} \underline{w}') = \det \hat{U} \hat{U} (\underline{v}' \times \underline{w}'). \tag{2.6.70}$$

Da $\hat{U}$ eine orthogonale Matrix ist, ist $\det U = \pm 1$, denn aus $\hat{U} \hat{U}^T = \mathbb{1}$ folgt

$$\det(\hat{U} \hat{U}^T) = \det \hat{U} \det \hat{U}^T = (\det \hat{U})^2 = \det \mathbb{1} = 1 \Rightarrow \det \hat{U} = \pm 1. \tag{2.6.71}$$

Dann besagt (2.6.70), dass sich für **orthogonale Basistransformationen** die Komponenten des Vektorprodukts wie ein Vektor transformieren, wenn die Basistransformation **orientierungserhaltend** ist und andernfalls ein zusätzliches Vorzeichen auftritt. Man nennt solche Größen daher genauer auch **Pseudovektoren** oder **axiale Vektoren**. Im letzteren Zusammenhang bezeichnet man dann die gewöhnlichen Vektoren auch als **polare Vektoren**.



Diese Benennung ergibt Sinn, wenn man als Spezialfall die orthogonale Transformation betrachtet, die durch $\hat{U} = -\mathbb{1}$ gegeben ist. Dies bedeutet, dass die neue Orthonormalbasis ($\vec{e}'_k$) aus der alten ($\vec{e}_j$) hervorgeht, indem man sich die Vektoren in einem gemeinsamen Anfangspunkt befestigt denkt (dem Ursprung des kartesischen Koordinatensystems) und an diesem Punkt spiegelt. Es ist klar, dass dann aus einem rechtshändigen ein linkshändiges Orthonormalbasissystem wird. Die Komponenten irgendwelcher polaren Vektoren $\vec{v}$ und $\vec{w}$ mit Anfangspunkt im Ursprung werden durch diese Spiegelung einfach mit (-1) multipliziert, d.h. $\underline{v}' = -\underline{v}$ und $\underline{w}' = -\underline{w}$, aber ihr Vektorprodukt wird zu $\underline{v}' \times \underline{w}' = (-\underline{v}) \times (-\underline{w}) = +\underline{v} \times \underline{w}$. Dies folgt auch anschaulich aus der

Rechte-Hand-Regel (s. die nebenstehende Skizze). Während also polare Vektoren unter **Raumspiegelungen** ihr Vorzeichen wechseln, ist dies für ihr Vektorprodukt nicht der Fall. Das Vektorprodukt zweier polarer Vektoren hat also geometrisch zumindest hinsichtlich ihrer Richtung eher mit einem Drehsinn als mit einer Parallelverschiebung zu tun.

Es ist klar, dass man (2.6.70) auch in der Form

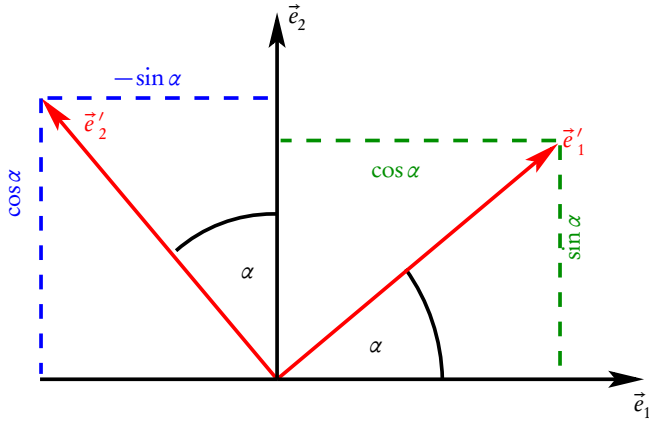
$$\underline{v}' \times \underline{w}' = (\hat{T}\underline{v}) \times (\hat{T}\underline{w}) = \det \hat{T} \hat{T} (\underline{v} \times \underline{w}) \quad (2.6.72)$$

schreiben kann, denn natürlich ist auch $\hat{T} = \hat{U}^{-1} = \hat{U}^T$ eine orthogonale Matrix.

2.7 Drehungen

Wir haben in Abschnitt 2.3 die orthogonalen Transformationen als diejenigen Basistransformationen eingeführt, die eine beliebige Orthonormalbasis in beliebige andere Orthonormalbasen abbilden. Im vorigen Abschnitt haben wir die Transformationen, für deren Transformationsmatrix $\det \hat{T} = \det \hat{U} = +1$ gilt, als orientierungserhaltend erkannt. Anschaulich ist klar, dass dies geometrisch einer **Drehung** des Koordinatensystems in der Ebene oder im Raum entspricht. In diesem Abschnitt betrachten wir die Drehungen etwas genauer.

2.7.1 Drehungen in der Ebene



In der nebenstehenden Zeichnung geht die Orthonormalbasis (kartesische Basis) $(\vec{e}_1, \vec{e}_2)$ durch eine Drehung um den **Drehwinkel** α aus der kartesischen Basis $(\vec{e}_1, \vec{e}_2)$ hervor. Beide Basen sind positiv orientiert, d.h. dreht man $\vec{e}_1$ entgegen dem Uhrzeigersinn (willkürlich als **positive Drehrichtung** in der Ebene definiert) um $\pi/2$ erhält man den Vektor $\vec{e}_2$ und entsprechend für die Basisvektoren $\vec{e}_1'$ und $\vec{e}_2'$. Nun gilt

$$\vec{e}_k' = \sum_{j=1}^2 U_{jk} \vec{e}_j. \quad (2.7.1)$$

Aus der Skizze lesen wir sofort ab, dass

$$\begin{aligned} \vec{e}_1' &= \cos \alpha \vec{e}_1 + \sin \alpha \vec{e}_2, \\ \vec{e}_2' &= -\sin \alpha \vec{e}_1 + \cos \alpha \vec{e}_2 \end{aligned} \quad (2.7.2)$$

ist. Die Transformationsmatrix ist also

$$\hat{U} = \hat{D}(\alpha) = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix}. \quad (2.7.3)$$

Es ist nun leicht nachzuprüfen, dass dies in der Tat eine orthogonale Transformation in der Ebene ist, denn es gilt

$$\begin{aligned} \hat{U} \hat{U}^T &= \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \begin{pmatrix} \cos \alpha & \sin \alpha \\ -\sin \alpha & \cos \alpha \end{pmatrix} \\ &= \begin{pmatrix} \cos^2 \alpha + \sin^2 \alpha & \cos \alpha \sin \alpha - \sin \alpha \cos \alpha \\ -\cos \alpha \sin \alpha + \sin \alpha \cos \alpha & \sin^2 \alpha + \cos^2 \alpha \end{pmatrix} = \mathbb{1}_2, \end{aligned} \quad (2.7.4)$$

denn für alle α gilt bekanntlich $\cos^2 \alpha + \sin^2 \alpha = 1$. Damit ist $\hat{U}^T = \hat{U}^{-1}$ und folglich die Transformationsmatrix eine Orthogonalmatrix. Sie ist offensichtlich auch orientierungserhaltend, denn die Entwicklung der

Determinante nach der ersten Spalte ergibt

$$\det \hat{U} = +\cos^2 \alpha - (-\sin^2 \alpha) = \cos^2 \alpha + \sin^2 \alpha = +1. \quad (2.7.5)$$

Als nächstes betrachten wir die Hintereinanderausführung zweier Drehungen. Es möge also zuerst die Basis $(\vec{e}'_1, \vec{e}'_2)$ durch Drehung um den Winkel α aus der Basis $(\vec{e}_1, \vec{e}_2)$ und dann $(\vec{e}''_1, \vec{e}''_2)$ durch Drehung um den Winkel β aus der Basis $(\vec{e}'_1, \vec{e}'_2)$ hervorgehen. Mit Hilfe der entsprechenden Drehmatrizen geschrieben ergibt sich

$$\vec{e}''_l = \sum_{k=1}^2 D_{kl}(\beta) \vec{e}'_k = \sum_{j,k=1}^2 D_{kl}(\beta) D_{jk}(\alpha) \vec{e}_j. \quad (2.7.6)$$

Die Transformationsmatrix, die direkt die Transformation von den $(\vec{e}_j)$ nach den $(\vec{e}''_l)$ angibt, ist also durch

$$U_{jl} = \sum_{k=1}^2 D_{jk}(\alpha) D_{kl}(\beta) \Rightarrow \hat{U} = \hat{D}(\alpha) \hat{D}(\beta) \quad (2.7.7)$$

gegeben. Anschaulich ist klar, dass $\hat{U} = \hat{D}(\alpha + \beta)$ sein sollte (vgl. die nebenstehende Skizze). Das können wir nun auch einfach nachrechnen:

$$\begin{aligned} \hat{D}(\alpha) \hat{D}(\beta) &= \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \begin{pmatrix} \cos \beta & -\sin \beta \\ \sin \beta & \cos \beta \end{pmatrix} \\ &= \begin{pmatrix} \cos \alpha \cos \beta - \sin \alpha \sin \beta & -\cos \alpha \sin \beta - \sin \alpha \cos \beta \\ \sin \alpha \cos \beta + \cos \alpha \sin \beta & -\sin \alpha \sin \beta + \cos \alpha \cos \beta \end{pmatrix} \\ &= \begin{pmatrix} \cos(\alpha + \beta) & -\sin(\alpha + \beta) \\ \sin(\alpha + \beta) & \cos(\alpha + \beta) \end{pmatrix} = \hat{D}(\alpha + \beta). \end{aligned} \quad (2.7.8)$$

Dabei haben wir die **Additionstheoreme** der trigonometrischen Funktionen

$$\cos(\alpha + \beta) = \cos \alpha \cos \beta - \sin \alpha \sin \beta, \quad \sin(\alpha + \beta) = \sin \alpha \cos \beta + \cos \alpha \sin \beta \quad (2.7.9)$$

verwendet.

Aus (2.7.8) wird auch sofort klar, dass **Drehungen in der Ebene kommutieren**, d.h. es gilt

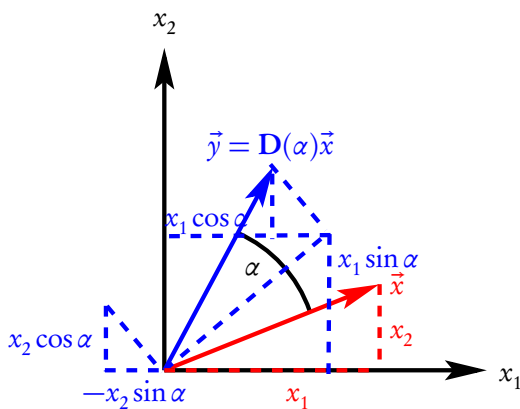
$$\hat{D}(\beta) \hat{D}(\alpha) = \hat{D}(\alpha) \hat{D}(\beta). \quad (2.7.10)$$

Aus (2.7.8) folgt auch, dass

$$\hat{D}(-\alpha) \hat{D}(\alpha) = \hat{D}(0) = \mathbb{1}_2 \quad (2.7.11)$$

ist. Es ist daher sinnvoll, negative Drehwinkel als Drehungen im Uhrzeigersinn (mathematisch negativer Drehsinn) zu betrachten. Andererseits ist klar, dass solche Drehungen auch als Drehungen im positiven Sinne um den Komplementärwinkel $2\pi - |\alpha|$ angesehen werden können. Alle Drehungen in der Ebene werden also durch die Matrizen $\hat{D}(\alpha)$ mit $\alpha \in [0, 2\pi)$ erfasst.

2.7. Drehungen



Wir betrachten nun noch Drehungen in einem etwas anderen Sinn, und zwar als **lineare Abbildung** von Vektoren. Lineare Abbildungen wollen wir mit einem aufrechten fett gedruckten Symbol bezeichnen. Dann soll die Drehung $D(\alpha)$ um den (positiven) Drehwinkel α einer Drehung entgegen dem Uhrzeigersinn um diesen Winkel bedeuten. Setzen wir also

$$\vec{y} = D(\alpha)\vec{x}, \quad (2.7.12)$$

so folgt aus der nebenstehenden Skizze, dass für die Komponenten in einem positiv orientierten kartesischen Koordinatensystem

$$\underline{y} = \hat{D}(\alpha)\underline{x} \quad (2.7.13)$$

stem

gilt.

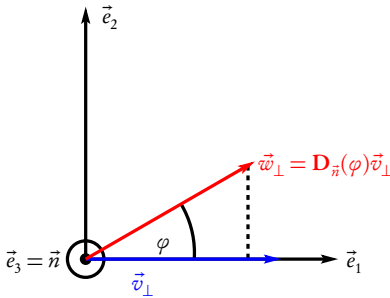
Manchmal bezeichnet man die Drehungen, die auftreten, wenn man die Transformation von **Vektorkomponenten** bei Drehungen einer kartesischen Basis in eine andere betrachtet als **passive Drehungen**, denn in dem Fall bleiben die Vektoren $\vec{v}$ fest, und nur die Komponenten ändern sich, weil man diesen Vektor durch Komponenten bzgl. verschiedener Basen betrachtet. Es geschieht also bei dieser Auffassung von Drehungen nichts mit dem Vektor als geometrischem Objekt sondern nur etwas hinsichtlich der Darstellung als $\mathbb{R}^2$ -Spaltenvektor aufgrund der Änderung der Basis.

Dazu bemerken wir, dass diese beiden Arten von Transformationen aus Sicht der Physik durchaus unterschiedliche Bedeutung besitzen: Wie nämlich im Lauf der Theorie-Vorlesungen noch klar werden wird, spielen die **Symmetrien der Naturgesetze** eine große Rolle in der Physik. Dies wird bereits an dieser Stelle deutlich: Allein dadurch, dass wir den physikalischen Raum als euklidischen affinen Raum modellieren, implizieren wir, dass dieser Raum bestimmte Symmetrien besitzt. So lässt sich, rein mathematisch betrachtet, kein Ort vor irgendeinem anderen Ort auszeichnen. Wir können geometrische Figuren beliebig parallel zueinander verschieben, ohne dass sich deren wesentliche Eigenschaften, wie z.B. die Längen und Winkel eines Dreiecks, irgendwie verändern. Diese Symmetrie unter **Verschiebungen oder Translationen** bezeichnet man als **Homogenität des Raumes**. Daran ist bemerkenswert, dass zunächst einmal eine solche Transformation aufgrund der mathematischen Struktur überhaupt definiert ist und zum Anderen auch bestimmte Eigenschaften von mathematischen Objekten bei einer solchen Transformation ungeändert bleiben. Genauso verändern sich Eigenschaften wie Längen und Winkel auch nicht bei **Drehungen**. Diese Symmetrie bezeichnen wir als **Isotropie des Raumes**. Betrachten wir nun über diese rein mathematischen Gegebenheiten hinausgehend auch **physikalische Gesetzmäßigkeiten**, wie z.B. in der Newtonschen Mechanik: Soll nun die Beschreibung des physikalischen Raumes als Euklidischer affiner Raum vollständig und streng gültig sein, sollten auch die **Naturgesetze**, z.B. die Gleichungen, die die Bewegung von Teilchen beschreiben, unabhängig davon sein, wo man eine solche Bewegung beobachtet und wie man die Versuchsanordnung relativ zu irgendeinem willkürlich gewählten Bezugssystem orientiert. Anders betrachtet bedeutet dies, dass es unmöglich ist, aufgrund des Verhaltens irgendwelcher physikalischer Objekte irgendeinen Ort oder irgendeine Orientierung auszuzeichnen. Dies ist die **aktive Auffassung** von Symmetrien. Andererseits ist es von vornherein klar, dass die Wahl einer Orthonormalbasis und eines Koordinatenursprungs im Raum und die Zuordnung von Koordinaten zu physikalischen Vektoren nichts an den Aussagen der Naturgesetze ändern darf, damit diese Beschreibung überhaupt sinnvoll als Beschreibung eines Naturvorgangs gelten kann. Die Naturgesetze bleiben demnach auch ungeändert, wenn man ein neues Koordinatensystem einführt, das sich lediglich durch eine Verschiebung und/oder eine Drehung der Orthonormalbasis unterscheidet. Dies ist die **passive Auffassung** von Symmetrieprinzipien.

2.7.2 Drehungen im Raum

Wir betrachten nun Drehungen im Raum. Anschaulich ist klar, dass wir jede Drehung im Raum durch eine **Drehachse**, die wir durch einen Einheitsvektor $\vec{n}$ festlegen, und einen **Drehwinkel** $\varphi \in [0, \pi]$ im Raum eindeutig beschreiben können. Die Drehung ist dabei wieder durch die **Rechte-Hand-Regel** festgelegt: Streckt man den Daumen in Richtung von $\vec{n}$, zeigen die gekrümmten Finger in die Richtung der Drehung. Hat man die Richtung so gewählt, dass man einen Drehwinkel $> \pi$ benötigen würde, kann man offenbar einfach statt $\vec{n}$ die Drehachse umorientieren, also stattdessen $-\vec{n}$ als Drehachse wählen, und beschreibt dann die gleiche Drehung mit einem Drehwinkel $< \pi$. Die entsprechende lineare Abbildung bezeichnen wir mit $\mathbf{D}_{\vec{n}}(\varphi)$.

Um die entsprechende (offenbar lineare) Abbildung zu finden, konstruieren wir uns ein rechtshändiges kartesisches Koordinatensystem, das an diese Situation angepasst ist. Dazu nehmen wir an, dass $\vec{v} \times \vec{n} \neq \vec{0}$ ist, d.h. der Vektor soll nicht parallel oder antiparallel zur Drehachse gerichtet sein.



Dann können wir die folgenden rechtshändigen kartesischen Basisvektoren definieren: Es sei $\vec{e}_3 = \vec{n}$, d.h. in diesem Basissystem beschreibt sich die Drehung als Drehung um die 3-Achse. Die beiden anderen Basisvektoren wählen wir als

$$\begin{aligned} \vec{e}_3 &= \vec{n}, & \vec{e}_2 &= \frac{\vec{n} \times \vec{v}}{|\vec{n} \times \vec{v}|}, \\ \vec{e}_1 &= \vec{e}_2 \times \vec{e}_3 = \frac{(\vec{n} \times \vec{v}) \times \vec{n}}{|\vec{n} \times \vec{v}|} = \frac{\vec{v} - \vec{n}(\vec{n} \cdot \vec{v})}{|\vec{n} \times \vec{v}|}. \end{aligned} \quad (2.7.14)$$

Offenbar ist (nachrechnen!)

$$\vec{v}_{\perp} = (\vec{n} \times \vec{v}) \times \vec{n} = \vec{v} - \vec{n}(\vec{n} \cdot \vec{v}), \quad |\vec{v}_{\perp}| = |\vec{v} \times \vec{n}| \quad (2.7.15)$$

die Projektion des Vektors $\vec{v}$ auf die (12)-Ebene, die senkrecht zur Drehachse $\vec{e}_3 = \vec{n}$ ist. Mit der Projektion auf die Richtung der Drehachse $\vec{v}_{\parallel} = \vec{n}(\vec{n} \cdot \vec{v})$ folgt

Aus der nebenstehenden Skizze lesen wir ab, dass

$$\vec{w} = \mathbf{D}_{\vec{n}}(\varphi)\vec{v} = \mathbf{D}(\varphi)[\vec{v}_{\parallel} + \vec{v}_{\perp}] = \underbrace{\vec{n}(\vec{n} \cdot \vec{v})}_{\vec{v}_{\parallel} = \vec{w}_{\parallel}} + \underbrace{|\vec{v}_{\perp}|(\vec{e}_1 \cos \varphi + \vec{e}_2 \sin \varphi)}_{\vec{w}_{\perp}}, \quad (2.7.16)$$

und mit den Definitionen der drei Basisvektoren (2.7.14) folgt schließlich

$$\begin{aligned} \mathbf{D}_{\vec{n}}(\varphi)\vec{v} &= \vec{n}(\vec{n} \cdot \vec{v}) + (\vec{n} \times \vec{v}) \sin \varphi + (\vec{n} \times \vec{v}) \times \vec{n} \cos \varphi \\ &= \vec{n}(\vec{n} \cdot \vec{v})(1 - \cos \varphi) + \vec{v} \cos \varphi + (\vec{n} \times \vec{v}) \sin \varphi. \end{aligned} \quad (2.7.17)$$

Diese Formel gilt freilich auch für den Fall, dass $\vec{v} = \vec{n}(\vec{n} \cdot \vec{v})$ ist, also falls $\vec{v} = \vec{v}_{\parallel}$ und damit $\vec{v}_{\perp} = 0$ und $\vec{v} \times \vec{n} = \vec{0}$ ist, denn dann ergibt die Formel, dass sich der Vektor bei der „Drehung um sich selbst“ nicht ändert, wie es sein muss.

2.7.3 Euler-Winkel

Wir betrachten nun noch eine andere Parametrisierung von Drehungen, die allerdings erst für die Behandlung des Kreisels relevant sein wird. Dort benötigt man sie für die Umrechnung zwischen Vektorkomponenten

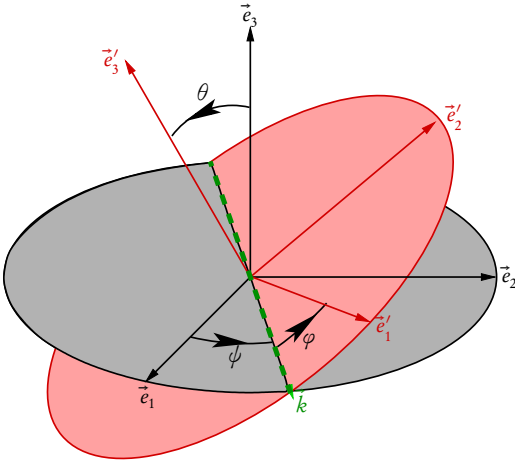
2.7. Drehungen

bzgl. einer rechtshändigen Orthonormalbasis $(\vec{e}_j)$, die im Bezugssystem eines raumfesten Beobachters ruht und einer rechtshändigen Orthonormalbasis $(\vec{e}'_k)$, die fest in dem betrachteten starren Körper verankert ist, wobei nur Drehungen um einen festgehaltenen Punkt, den wir als den Koordinatenursprung wählen, betrachtet werden. Zu irgendeinem festen Zeitpunkt ist offenbar die rechtshändige Orthonormalbasis $(\vec{e}'_k)$ gegenüber der rechtshändigen Orthonormalbasis $(\vec{e}_j)$ verdreht, und wir wollen die entsprechende Drehmatrix $\hat{D}$ finden, so dass

$$\vec{e}'_k = \sum_{l=1}^3 D_{lk} \vec{e}_l \quad (2.7.18)$$

ist. Dazu führen wir zunächst die Drehmatrizen für die Drehung um die drei Koordinatenachsen eines rechtshändigen kartesischen Koordinatensystems ein:

$$\begin{aligned} \hat{D}_1(\varphi) &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \varphi & -\sin \varphi \\ 0 & \sin \varphi & \cos \varphi \end{pmatrix}, \\ \hat{D}_2(\varphi) &= \begin{pmatrix} \cos \varphi & 0 & \sin \varphi \\ 0 & 1 & 0 \\ -\sin \varphi & 0 & \cos \varphi \end{pmatrix}, \\ \hat{D}_3(\varphi) &= \begin{pmatrix} \cos \varphi & -\sin \varphi & 0 \\ \sin \varphi & \cos \varphi & 0 \\ 0 & 0 & 1 \end{pmatrix}. \end{aligned} \quad (2.7.19)$$



Betrachten wir nun die geometrischen Verhältnisse der beiden rechtshändigen Orthonormalsysteme $(\vec{e}_j)$ und $(\vec{e}'_k)$ anhand der nebenstehenden Skizze: Die Schnittlinie der 12- mit der 1'2'-Ebene definiert die sogenannte Knotenlinie, die wir gestrichelt grün eingezeichnet haben. Der in ihre Richtung weisende Einheitsvektor $\vec{k}$ ist nach der Rechte-Hand-Regel so orientiert, dass bei der Drehung des Orthonormalsystems $(\vec{e}_j)$ um diese Achse die 3-Achse in die 3'-Achse gedreht wird, und zwar so, dass der Drehwinkel ins Intervall $\vartheta \in [0, \pi]$ fällt.

Das Orthonormalsystem $(\vec{e}_j)$ wird nun nacheinander durch drei Drehungen in das Dreibein $\vec{e}'_j$ verdreht, und zwar wie folgt: Zunächst erfolgt eine Drehung um die 3-Achse um den Drehwinkel $\psi \in [0, 2\pi)$, so dass die neue Einsachse $\vec{e}''_1$ mit dem Knotenlinienvektor $\vec{k}$ zusammenfällt:

$$\vec{e}''_j = \sum_{k=1}^3 \vec{e}_k [\hat{D}_3(\psi)]_{kj}. \quad (2.7.20)$$

Sodann erfolgt eine Drehung um die Achse $\vec{k} = \vec{e}''_1$ um den Winkel ϑ , die dafür sorgt, dass die neue Dreiachse $\vec{e}'''_3$ nunmehr mit der 3'-Achse zusammenfällt:

$$\vec{e}'''_l = \sum_{j=1}^3 \vec{e}''_j [\hat{D}_1(\vartheta)]_{jl} = \sum_{j,k=1}^3 \vec{e}_k [\hat{D}_3(\psi)]_{kj} [\hat{D}_1(\vartheta)]_{jl}. \quad (2.7.21)$$

Schließlich wird noch um die Achse $\vec{e}'''_3 = \vec{e}'_3$ um den Winkel φ gedreht, so dass schließlich auch die neuen

Eins- und Zweiachsen jeweils auf $\vec{e}'_1$ und $\vec{e}'_2$ zu liegen kommen:

$$\begin{aligned}\vec{e}'_m &= \sum_{l=1}^3 \vec{e}'''_l [\hat{D}_3(\varphi)]_{lm} \\ &= \sum_{j,k,l=1}^3 \vec{e}_k [\hat{D}_3(\psi)]_{kj} [\hat{D}_1(\vartheta)]_{jl} [\hat{D}_3(\varphi)]_{lm} \\ &= \sum_{k=1}^3 \vec{e}_k [\hat{D}_3(\psi) \hat{D}_1(\vartheta) \hat{D}_3(\varphi)]_{km} = \sum_{k=1}^3 \vec{e}_k D_{km}.\end{aligned}\tag{2.7.22}$$

Für einen beliebigen Vektor $\vec{x}$ folgt dann

$$\vec{x} = \sum_{m=1}^3 \vec{e}'_m x'_m = \sum_{k,m=1}^3 \vec{e}_k D_{km} x'_m \Rightarrow \underline{x} = \hat{D}(\psi, \vartheta, \varphi) \underline{x}',\tag{2.7.23}$$

wobei gemäß (2.7.22)

$$\hat{D}(\psi, \vartheta, \varphi) = \hat{D}_3(\psi) \hat{D}_1(\vartheta) \hat{D}_3(\varphi)\tag{2.7.24}$$

ist.

Es ist sehr wichtig zu bemerken, dass die Hintereinanderausführung von *Drehungen um verschiedene Achsen nicht kommutativ* ist, d.h. es ist

$$\hat{D}_3(\alpha) \hat{D}_1(\beta) \neq \hat{D}_1(\beta) \hat{D}_3(\alpha).\tag{2.7.25}$$

Wir schreiben schließlich die in **Euler-Winkeln** $(\psi, \vartheta, \varphi)$ parametrisierte Drehmatrix (2.7.24) noch explizit auf. Man erhält diese Form einfach durch direkte Matrizenmultiplikation:

$$\hat{D} = \begin{pmatrix} \cos \psi \cos \varphi - \sin \psi \cos \vartheta \sin \varphi & -\cos \psi \sin \varphi - \sin \psi \cos \vartheta \cos \varphi & \sin \psi \sin \vartheta \\ \sin \psi \cos \varphi + \cos \psi \cos \vartheta \sin \varphi & -\sin \psi \sin \varphi + \cos \psi \cos \vartheta \cos \varphi & -\cos \psi \sin \vartheta \\ \sin \vartheta \sin \varphi & \sin \vartheta \cos \varphi & \cos \vartheta \end{pmatrix}.\tag{2.7.26}$$

2.8 Lineare Abbildungen

Wir betrachten eine allgemeine lineare Abbildung $\mathbf{M}: V \rightarrow V$, $\vec{v} \mapsto \mathbf{M}\vec{v}$, wobei V ein beliebiger (reeller) d -dimensionaler Vektorraum ist. Die Linearität bedeutet, dass für $\lambda_1, \lambda_2 \in \mathbb{R}$ und $\vec{v}_1, \vec{v}_2 \in V$

$$\mathbf{M}(\lambda_1 \vec{v}_1 + \lambda_2 \vec{v}_2) = \lambda_1 \mathbf{M}\vec{v}_1 + \lambda_2 \mathbf{M}\vec{v}_2\tag{2.8.1}$$

gilt. Wegen der Linearität ist die lineare Abbildung bereits vollständig durch die Abbildung der Vektoren $\vec{b}_j$ einer beliebigen Basis gegeben. Wir können die lineare Abbildung also eindeutig durch eine $(d \times d)$ -Matrix beschreiben:

$$\mathbf{M}\vec{b}_j = M_{kj} \vec{b}_k,\tag{2.8.2}$$

wobei wieder gemäß der Einsteinschen Summenkonvention über den doppelt auftretenden Index k zu summieren ist. Wir bezeichnen die Matrix mit $\hat{M} = (M_{kj})$. Die lineare Abbildung ist dann durch die entsprechende Matrixmultiplikation der Vektorkomponenten bzgl. dieser Basis eindeutig bestimmt:

$$\mathbf{M}\vec{v} = \mathbf{M}(v_j \vec{b}_j) = M_{kj} v_j \vec{b}_j,\tag{2.8.3}$$

d.h.

$$\underline{\mathbf{M}}\underline{v} = \hat{\underline{\mathbf{M}}}\underline{v}. \quad (2.8.4)$$

Wir untersuchen nun, wie sich die Matrix ändert, wenn wir zu einer neuen Basis

$$\vec{b}'_k = U_{jk} \vec{b}_j \quad (2.8.5)$$

übergehen. Um die entsprechende Matrix $\hat{\underline{\mathbf{M}}}'$ zu bestimmen, beachten wir, dass die Transformation zwischen den Vektorkomponenten durch

$$\vec{v} = v_j \vec{b}_j = v'_k \vec{b}'_k = v'_k U_{jk} \vec{b}_j \Rightarrow \underline{v} = \hat{U} \underline{v}' \quad \underline{v}' = \hat{U}^{-1} \underline{v} \quad (2.8.6)$$

gegeben ist. Demnach ist

$$\underline{\mathbf{M}}\underline{v}' = \hat{\underline{\mathbf{M}}}'\underline{v}' = \hat{U}^{-1} \hat{\underline{\mathbf{M}}} \underline{v} = \hat{U}^{-1} \hat{\underline{\mathbf{M}}} \hat{U} \underline{v}', \quad (2.8.7)$$

und das bedeutet

$$\hat{\underline{\mathbf{M}}}' = \hat{U}^{-1} \hat{\underline{\mathbf{M}}} \hat{U}. \quad (2.8.8)$$

2.9 Eigenwertprobleme

In diesem Abschnitt betrachten wir die sogenannten **Eigenwertprobleme** von Matrizen und betrachten dazu ausnahmsweise einen Vektorraum mit **komplexen Zahlen** als Skalare (vgl. Kapitel 4), da dies die Diskussion von Eigenwerten etwas vereinfacht. Dabei ist eine Matrix $\hat{\underline{\mathbf{M}}} \in \mathbb{C}^{d \times d}$ gegeben, und es werden Vektoren $\vec{u} \in \mathbb{C}^d$ und Zahlen $\lambda \in \mathbb{C}$ gesucht, für die

$$\hat{\underline{\mathbf{M}}} \vec{u} = \lambda \vec{u} \quad (2.9.1)$$

gilt. Wir betrachten dabei zunächst komplexe Vektoren und Matrizen, weil dies die Lösung des Problems etwas vereinfacht. Die Vektoren $\vec{u}$ heißen **Eigenvektoren** und die Zahlen λ **Eigenwerte**. Wir werden am Ende noch den wichtigen Spezialfall reeller symmetrischer Matrizen behandeln, wo man mit reellen Vektoren auskommt.

Zunächst betrachten wir ein typisches physikalisches Problem, wo Eigenwerte und Eigenvektoren eine Rolle spielen.

2.9.1 Beispiel: Durch Feder verbundene Massen

Wir betrachten zwei Massen m_1 und m_2 , die reibungsfrei auf einer Stange gleiten können und durch eine masselose Feder verbunden sind. Es seien x_1 und x_2 die Koordinaten von Massenpunkt 1 und Massenpunkt 2 relativ zu einer Lage beider Massen, für die die Feder entspannt ist. Die Bewegungsgleichungen lauten dann wegen des Hookschen Gesetzes

$$m_1 \ddot{x}_1 = -D(x_1 - x_2), \quad m_2 \ddot{x}_2 = -D(x_2 - x_1) \quad (2.9.2)$$

oder mit $\omega_1 = \sqrt{D/m_1}$ und $\omega_2 = \sqrt{D/m_2}$ und in Matrix-Vektorschreibweise ausgedrückt

$$\begin{pmatrix} \ddot{x}_1 \\ \ddot{x}_2 \end{pmatrix} = \begin{pmatrix} -\omega_1^2 & \omega_1^2 \\ \omega_2^2 & -\omega_2^2 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}. \quad (2.9.3)$$

Setzen wir

$$\vec{u} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}, \quad \hat{M} = \begin{pmatrix} -\omega_1^2 & \omega_1^2 \\ \omega_2^2 & -\omega_2^2 \end{pmatrix} \quad (2.9.4)$$

lautet das Differentialgleichungssystem (2.9.3)

$$\ddot{\vec{u}} = \hat{M}\vec{u}. \quad (2.9.5)$$

Die Form der Differentialgleichung für den Vektor $\vec{u}$ legt den **Exponentialansatz**

$$\vec{u}(t) = \vec{a} \exp(\lambda t) \quad (2.9.6)$$

mit $\vec{a} = \text{const}$ nahe. Setzen wir dies in (2.9.5) ein, erhalten wir

$$\ddot{\vec{u}} = \lambda^2 \vec{a} \exp(\lambda t) = \hat{M} \vec{a} \exp(\lambda t) \Rightarrow \hat{M} \vec{a} = \Lambda \vec{a} \quad (2.9.7)$$

mit $\Lambda = \lambda^2$. Wir werden also auf ein **Eigenwertproblem** geführt. Wir suchen also Zahlen Λ und dazugehörige **Eigenvektoren** $\vec{a}$, die (2.9.7) erfüllen. Wir können diese Gleichung offenbar zu

$$(\hat{M} - \Lambda \mathbb{1}) \vec{a} = \vec{0} \quad (2.9.8)$$

umformen. Es ist klar, dass dieses lineare Gleichungssystem mit dem Nullvektor auf der rechten Seite immer die triviale Lösung $\vec{a} = \vec{0}$ besitzt, wobei Λ eine beliebige Zahl sein kann. Diese triviale Lösung hilft uns aber bei der Suche nach der allgemeinen Lösung unseres Differentialgleichungssystems (2.9.5) nicht weiter, d.h. wir suchen **nichttriviale Lösungen** von (2.9.8) mit $\vec{a} \neq \vec{0}$.

Wir wissen nun, dass die Lösung des linearen Gleichungssystems (2.9.8) genau dann eindeutig ist, wenn $\det(\hat{M} - \Lambda \mathbb{1}) \neq 0$ ist. Dann besitzt das Gleichungssystem aber eben nur die triviale Lösung $\vec{a} = \vec{0}$.

Um eine nichttriviale Lösung zu erhalten, müssen wir also fordern, dass

$$\det(\hat{M} - \Lambda \mathbb{1}) = 0 \quad (2.9.9)$$

ist. Dies bestimmt die möglichen **Eigenwerte** der Matrix $\hat{M}$.

Ist nämlich (2.9.9) erfüllt, besitzt (2.9.8) immer mindestens eine nichttriviale Lösung mit $\vec{a} \neq \vec{0}$, d.h. zu jedem Eigenwert existiert wenigstens ein Eigenvektor.

Setzen wir die in (2.9.4) gegebene Matrix in (2.9.9) ein, folgt

$$\det(\hat{M} - \Lambda \mathbb{1}) = \det \begin{pmatrix} -\omega_1^2 - \Lambda & \omega_1^2 \\ \omega_2^2 & -\omega_2^2 - \Lambda \end{pmatrix} = (\omega_1^2 + \omega_2^2)\Lambda + \Lambda^2 \stackrel{!}{=} 0. \quad (2.9.10)$$

Wir finden hier also zwei Eigenwerte

$$\Lambda_1 = 0, \quad \Lambda_2 = -(\omega_1^2 + \omega_2^2) = -\Omega^2. \quad (2.9.11)$$

Weiter benötigen wir die dazugehörigen Eigenvektoren. Dazu lösen wir (2.9.8) für die jeweiligen Eigenwerte. Für $\Lambda_1 = 0$ folgt (*Nachrechnen!*)

$$\hat{M} \vec{a}_1 = 0 \Rightarrow \vec{a}_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix} \quad (2.9.12)$$

Es ist klar, dass jeder Vektor $\vec{a}'_1 = c \vec{a}_1$ ebenfalls eine Lösung ist. Wir benötigen im Folgenden für jeden Eigenwert aber nur *einen* beliebigen von Null verschiedenen Eigenvektor. Für $\Lambda_2 = -\Omega^2$ ergibt sich

$$(\hat{M} - \Lambda_2 \mathbb{1}) \vec{a}_2 = \begin{pmatrix} \omega_2^2 & \omega_1^2 \\ \omega_2^2 & \omega_1^2 \end{pmatrix} \vec{a}_2 = 0 \Rightarrow \vec{a}_2 = \begin{pmatrix} 1 \\ -\omega_2^2/\omega_1^2 \end{pmatrix}. \quad (2.9.13)$$

2.9. Eigenwertprobleme

Wir haben also zwei Eigenwerte und die dazugehörigen Eigenvektoren gefunden. Wir bemerken auch, dass die Eigenvektoren linear unabhängig sind. Wir können also jeden Vektor $\vec{v} \in \mathbb{C}^2$ als eindeutige Linearkombination aus $\vec{a}_1$ und $\vec{a}_2$ darstellen. Das nutzen wir nun aus, um unser Differentialgleichungssystem (2.9.5) zu lösen, indem wir die gesuchte Funktion $\vec{u}(t)$ als Linearkombination der beiden Eigenvektoren schreiben:

$$\vec{u}(t) = u_1(t)\vec{a}_1 + u_2(t)\vec{a}_2. \quad (2.9.14)$$

Setzen wir diesen Ansatz in die Differentialgleichung ein, folgt

$$\ddot{\vec{u}}(t) = \ddot{u}_1(t)\vec{a}_1 + \ddot{u}_2(t)\vec{a}_2 = u_1(t)\hat{M}\vec{a}_1 + u_2(t)\hat{M}\vec{a}_2 = \Lambda_1 u_1(t)\vec{a}_1 + u_2(t)\Lambda_2\vec{a}_2. \quad (2.9.15)$$

Da $\vec{a}_1$ und $\vec{a}_2$ linear unabhängig voneinander sind, muss offenbar gelten

$$\ddot{u}_1(t) = \Lambda_1 u_1(t) = 0, \quad \ddot{u}_2(t) = \Lambda_2 u_2(t) = -\Omega^2 u_2(t). \quad (2.9.16)$$

Wir haben also zwei einfache lineare Differentialgleichungen für f_1 und f_2 gefunden, und diese können wir sehr leicht lösen. Die erste Gleichung müssen wir nur zweimal bzgl. t integrieren, und die zweite Gleichung ist die eines ungedämpften harmonischen Oszillators mit Kreisfrequenz Ω . Die allgemeinen Lösungen sind also

$$u_1(t) = C_{11}t + C_{12}, \quad u_2(t) = C_{21}\cos(\Omega t) + C_{22}\sin(\Omega t). \quad (2.9.17)$$

Dabei sind C_{11} , C_{12} , C_{21} und C_{22} Integrationskonstanten, die ggf. aus den Anfangsbedingungen bestimmt werden können. Durch (2.9.14) mit den Lösungen (2.9.17) für $u_1(t)$ und $u_2(t)$ haben wir damit die allgemeine Lösung unseres mechanischen Problems gefunden. Eine Lösung mit $u_1(t) \neq 0$ und $u_2(t) = 0$ entspricht dabei offenbar der gleichförmigen Bewegung beider Massen, d.h. es wirken keine Kräfte auf die Massenpunkte, und die Feder ist entspannt. Die Massenpunkte laufen also einfach mit der gleichen konstanten Geschwindigkeit in einem Abstand, der der Länge der entspannten Feder entspricht, auf der Stange entlang. Der Lösung mit $u_1(t) = 0$ entspricht eine Schwingung der beiden Massenpunkte um den gemeinsamen Schwerpunkt mit der Kreisfrequenz $\Omega = \sqrt{\omega_1^2 + \omega_2^2}$. Dies nennt man eine **Eigenschwingung** des Systems.

2.9.2 Allgemeine Eigenschaften von Eigenvektoren und Eigenwerten

Wir sehen also anhand dieses Beispiels, dass es sehr nützlich sein kann, Eigenvektoren und Eigenwerte für eine Matrix zu suchen. Insbesondere ist es bequem, wenn eine Matrix $\hat{M} \in \mathbb{C}^{d \times d}$ genau d voneinander linear unabhängige Eigenvektoren $\vec{u}_k$ besitzt, denn dann kann man diese Vektoren als Basis für den Vektorraum verwenden, und die durch die Matrix definierte **lineare Abbildung** nimmt eine besonders einfache Form an. Dann ist nämlich für einen beliebigen Vektor

$$\vec{v} = \sum_{k=1}^d v'_k \vec{u}_k \Rightarrow \hat{M}\vec{v} = \sum_{k=1}^d v'_k \hat{M}\vec{u}_k = \sum_{k=1}^d \lambda_k v'_k \vec{u}_k. \quad (2.9.18)$$

Bzgl. der Basis $\vec{u}_k$ von Eigenvektoren ist also

$$(\hat{M}\vec{v})'_k = \lambda_k v'_k. \quad (2.9.19)$$

Die lineare Abbildung $\vec{v} \mapsto \hat{M}\vec{v}$ wird also bzgl. dieser Basis von Eigenvektoren durch eine **Diagonalmatrix** dargestellt, d.h.

$$\underline{v}' = \begin{pmatrix} v'_1 \\ \vdots \\ v'_d \end{pmatrix} \mapsto \hat{M}' \underline{v}' = \begin{pmatrix} \lambda_1 & 0 & 0 & \cdots & 0 \\ 0 & \lambda_2 & 0 & \cdots & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \cdots & \lambda_d \end{pmatrix} \underline{v}' = \text{diag}(\lambda_1, \dots, \lambda_d) \underline{v}'. \quad (2.9.20)$$

2. Lineare Algebra

Gibt es also eine Basis von Eigenvektoren der Matrix $\hat{M}$, können wir durch eine Basistransformation erreichen, dass die entsprechende lineare Abbildung bzgl. der neuen Basis durch eine **Diagonalmatrix** darstellbar ist. Man sagt dann auch, die Matrix $\hat{M}$ sei **diagonalisierbar**. Es ist wichtig zu bemerken, dass Matrizen i.a., *nicht* diagonalisierbar sein müssen!

Bezeichnen wir die ursprüngliche Basis mit $\vec{b}_j$ und definieren die entsprechende Transformationsmatrix $\hat{T}$ durch

$$\vec{b}_j = \sum_{k=1}^j \vec{u}_k T_{kj} \quad (2.9.21)$$

folgt

$$\vec{v} = \sum_{j=1}^d v_j \vec{b}_j = \sum_{j,k=1}^d \vec{u}_k T_{kj} v_j \Rightarrow \underline{v}' = \hat{T} \underline{v}. \quad (2.9.22)$$

Für die lineare Abbildung folgt

$$\hat{M} \vec{v} = \sum_{j,l=1}^d \vec{b}_l M_{lj} v_j = \sum_{j,k,l=1}^d \vec{u}_k T_{kl} M_{lj} v_j = \sum_{k,m} \vec{u}_k M'_{km} v'_m \quad (2.9.23)$$

Das liefert in Matrix-Vektorschreibweise ausgedrückt

$$\hat{T} \hat{M} \underline{v} = \hat{M}' \underline{v}'. \quad (2.9.24)$$

Da $\hat{T}$ eine Matrix ist, die einen Basiswechsel definiert, ist sie invertierbar, und aus (2.9.22) folgt $\underline{v} = \hat{T}^{-1} \underline{v}'$. Setzen wir dies auf der linken Seite von (2.9.24) ein, folgt

$$\hat{T} \hat{M} \hat{T}^{-1} \underline{v}' = \hat{M}' \underline{v}'. \quad (2.9.25)$$

Da dies offenbar für beliebige $\underline{v}' \in \mathbb{C}^d$ gilt, ist

$$\hat{M}' = \hat{T} \hat{M} \hat{T}^{-1}. \quad (2.9.26)$$

D.h. man kann die Matrix $\hat{M}$ mit einer invertierbaren Matrix $\hat{T}$ gemäß (2.9.26) diagonalisieren, wenn es eine Basis $\vec{u}_k$ von Eigenvektoren von $\hat{M}$ gibt.

Aus (2.9.21) folgt nun, dass

$$\vec{u}_k = \sum_{j=1}^d (T^{-1})_{jk} \vec{b}_j. \quad (2.9.27)$$

Da die $\vec{b}_j$ die üblichen kanonischen Einheitsvektoren sind, d.h. der Spaltenvektor mit der j -Komponente 1 und alle anderen Komponenten 0 ist $\vec{b}_j \cdot \vec{b}_l = \delta_{jl}$ und damit

$$\vec{b}_j \cdot \vec{u}_k = (T^{-1})_{jk}, \quad (2.9.28)$$

d.h. die inverse Transformationsmatrix $\hat{T}^{-1} = (\vec{u}_1, \vec{u}_2, \dots, \vec{u}_d)$: Die Spalten der Matrix $\hat{T}^{-1}$ sind durch die Spalteneigenvektoren $\vec{u}_k$ der Matrix $\hat{M}$ gegeben.

In unserem obigen mechanischen Beispiel ist also (*Nachrechnen!*)

$$\hat{T}^{-1} = \begin{pmatrix} 1 & 1 \\ 1 & -\omega_2^2/\omega_1^2 \end{pmatrix} \Rightarrow \hat{T} = \frac{1}{\Omega^2} \begin{pmatrix} \omega_2 & \omega_1^2 \\ \omega_1^2 & -\omega_2^2 \end{pmatrix} \quad (2.9.29)$$

und in der Tat folgt daraus mit (2.9.4)

$$\hat{T}\hat{M}\hat{T}^{-1} = \begin{pmatrix} 0 & 0 \\ 0 & -\Omega^2 \end{pmatrix} = \text{diag}(0, -\Omega^2). \quad (2.9.30)$$

Die Transformationsmatrix diagonalisiert also $\hat{M}$, wie in (2.9.26) gezeigt.

Ob eine Matrix diagonalisierbar ist, entscheidet sich also daran, ob es eine Basis von Eigenvektoren dieser Matrix gibt. Wie bereits im obigen Beispiel argumentiert wurde, ergeben sich die Eigenwerte der Matrix daraus, dass die Eigenwertgleichung

$$\hat{M}\vec{u} = \lambda\vec{u} \Leftrightarrow (\hat{M} - \lambda\mathbb{1})\vec{u} = \vec{0} \quad (2.9.31)$$

nur dann nichttriviale Lösungen $\vec{u} \neq 0$ besitzen kann, wenn

$$P_{\hat{M}}(\lambda) = \det(\hat{M} - \lambda\mathbb{1}) = 0 \quad (2.9.32)$$

ist. Die Determinante ist offensichtlich ein Polynom vom Grad d , und die Nullstellen dieses Polynoms sind Eigenwerte der Matrix, und zu jedem Eigenwert gibt es dann wenigstens einen Eigenvektor. Das Polynom heißt das **charakteristische Polynom** der Matrix $\hat{M}$. Ohne Beweis bemerken wir, dass es dem **Hauptsatz der Algebra** zufolge für jedes Polynom wenigstens eine komplexe Nullstelle gibt. Durch Induktion kann man daraus sofort folgern, dass jedes Polynom in Linearfaktoren zerfällt, also

$$P_{\hat{M}}(\lambda) = A(\lambda - \lambda_1)(\lambda - \lambda_2) \cdots (\lambda - \lambda_d). \quad (2.9.33)$$

Dabei kann es vorkommen, dass von den Nullstellen mehrere gleich sind. Z.B. hat das Polynom $P(\lambda) = (\lambda - 1)^r$ nur eine Nullstelle $\lambda_1 = 1$. Man nennt dann r die Vielfachheit der Nullstelle bzw. λ_1 ist eine Nullstelle r -ten Grades.

Als nächstes zeigen wir, dass wenn $\vec{u}_1, \vec{u}_2, \dots, \vec{u}_j$ Eigenvektoren von $\hat{M}$ mit lauter *verschiedenen* Eigenwerten $\lambda_1, \dots, \lambda_j$ sind, diese Eigenvektoren untereinander linear unabhängig sind. Das beweisen wir durch vollständige Induktion nach der Anzahl der Eigenvektoren. Haben wir nur einen Eigenvektor $\vec{u}_1$ zum Eigenwert λ_1 ist definitionsgemäß $\vec{u}_1 \neq 0$ und damit linear unabhängig. Nehmen wir nun an, dass die Vektoren $\vec{u}_1, \dots, \vec{u}_{r-1}$ linear unabhängig sind. Dann wollen wir zeigen, dass dann auch die Vektoren $\vec{u}_1, \dots, \vec{u}_r$ linear unabhängig sind. Es sei also

$$\sum_{k=1}^r \mu_k \vec{u}_k = \vec{0}. \quad (2.9.34)$$

Wenden wir darauf die Matrix $\hat{M}$ an, folgt

$$\sum_{k=1}^r \mu_k \hat{M} \vec{u}_k = \sum_{k=1}^r \mu_k \lambda_k \vec{u}_k = \hat{M} \vec{0} = \vec{0}. \quad (2.9.35)$$

Multiplizieren wir andererseits (2.9.34) mit λ_r und subtrahieren die entstehende Gleichung von (2.9.35), folgt

$$\sum_{k=1}^{r-1} \mu_k (\lambda_k - \lambda_r) \vec{u}_k = \vec{0}. \quad (2.9.36)$$

Da nun voraussetzungsgemäß die Vektoren $\vec{u}_1, \dots, \vec{u}_{r-1}$ linear unabhängig sind, folgt $\mu_k (\lambda_k - \lambda_r) = 0$ für $k \in \{1, 2, \dots, r-1\}$. Da für diese k voraussetzungsgemäß $\lambda_k - \lambda_r \neq 0$ ist, folgt notwendig $\mu_k = 0$ für $k \in \{1, \dots, r-1\}$. Damit demnach also (2.9.34) gelten kann, muss $\mu_r \vec{u}_r = \vec{0}$ gelten. Da $\vec{u}_j$ also Eigenvektor von $\hat{M}$ definitionsgemäß nicht der Nullvektor ist, muss auch $\mu_r = 0$ sein. Damit sind also die $\vec{u}_1, \dots, \vec{u}_r$ linear unabhängig.

Daraus ergibt sich sofort, dass eine Matrix $\hat{M}$, deren charakteristisches Polynom $P_{\hat{M}}(\lambda)$ genau d verschiedene Nullstellen besitzt, diagonalisierbar ist, denn dann sind die dazugehörigen Eigenvektoren $\vec{u}_1, \dots, \vec{u}_d$ linear unabhängig und damit eine Basis von $\mathbb{C}^d$ und folglich die Matrix diagonalisierbar.

Ist eine Nullstelle λ_j des charakteristischen Polynoms eine r -fache Nullstelle, kann es sein, dass die Matrix nicht diagonalisierbar ist. Es seien also $\lambda_1, \dots, \lambda_k$ die *verschiedenen* Nullstellen. Sei $g(\lambda_j)$ der Grad der Nullstelle λ_j .

Da jedes Polynom im komplexen Zahlenkörper in Linearfaktoren zerfällt, ist offenbar $g(\lambda_1) + g(\lambda_2) + \dots + g(\lambda_k) = d$, und die Matrix ist demnach genau dann diagonalisierbar, wenn es zu jedem Eigenwert λ_j genau $g(\lambda_j)$ *linear unabhängige* Eigenvektoren gibt.

Das muss offenbar nicht notwendig der Fall sein. Betrachten wir als ein einfaches Beispiel die Matrix

$$\hat{M} = \begin{pmatrix} 2 & 1 \\ 0 & 2 \end{pmatrix} \in \mathbb{C}^{2 \times 2}. \quad (2.9.37)$$

Das charakteristische Polynom ist

$$P_{\hat{M}}(\lambda) = \det(\hat{M} - \lambda \mathbb{1}) = (2 - \lambda)^2. \quad (2.9.38)$$

Das Polynom besitzt also die doppelte Nullstelle $\lambda_1 = \lambda_2 = 2$. Der einzige Eigenwert der Matrix ist also 2. Suchen wir nun die Eigenvektoren:

$$(\hat{M} - 2\mathbb{1})\vec{u} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \vec{u} \stackrel{!}{=} \vec{0} \Rightarrow u_2 = 0. \quad (2.9.39)$$

Es gibt also (bis auf einen Faktor) nur einen einzigen Eigenvektor $\vec{u} = (1, 0)$, d.h. es gibt *keine* Basis aus Eigenvektoren, und die Matrix (2.9.37) ist daher *nicht* diagonalisierbar.

2.9.3 Symmetrische reelle Matrizen und die Hauptachsentransformation

In diesem Abschnitt wollen wir zeigen, dass eine symmetrische reelle Matrix $\hat{M} \in \mathbb{R}^{d \times d}$ mit $\hat{M} = \hat{M}^T$, stets diagonalisierbar ist. Dabei sind alle **Eigenwerte reell**, und es gibt stets eine **kartesische Basis (also eine Orthonormalbasis)** von Eigenvektoren.

Dies impliziert, dass es für symmetrische reelle Matrizen eine **orthogonale Matrix $\hat{T}$** gibt, so dass $\hat{M}' = \hat{T} \hat{M} \hat{T}^{-1} = \hat{T} \hat{M} \hat{T}^T = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_d)$ ist. Man nennt diese Basistransformation dann **Hauptachsentransformation**.

Um dieses wichtige Resultat der linearen Algebra zu beweisen, führen wir die Matrix als darstellende Matrix einer **symmetrischen Bilinearform** ein. Dabei ist $B : V \times V \rightarrow \mathbb{R}$, wobei V ein d -dimensionaler reeller Vektorraum mit Skalarprodukt sein soll. Die Bilinearform ist linear in beiden Argumenten, d.h. es gilt für beliebige $\lambda_1, \lambda_2 \in \mathbb{R}$ und beliebige $\vec{v}, \vec{v}_1, \vec{v}_2, \vec{w}, \vec{w}_1, \vec{w}_2 \in V$

$$B(\lambda_1 \vec{v}_1 + \lambda_2 \vec{v}_2, \vec{w}) = \lambda_1 B(\vec{v}_1, \vec{w}) + \lambda_2 B(\vec{v}_2, \vec{w}) \quad (2.9.40)$$

$$B(\vec{v}, \lambda_1 \vec{w}_1 + \lambda_2 \vec{w}_2) = \lambda_1 B(\vec{v}, \vec{w}_1) + \lambda_2 B(\vec{v}, \vec{w}_2). \quad (2.9.41)$$

Die Bilinearform heißt weiter symmetrisch, wenn für beliebige $\vec{v}, \vec{w} \in V$ stets

$$B(\vec{v}, \vec{w}) = B(\vec{w}, \vec{v}) \quad (2.9.42)$$

2.9. Eigenwertprobleme

ist.

Seien nun $\vec{e}_j$ mit $j \in \{1, 2, \dots, d\}$ ein Orthonormalsystem, d.h. es gelte $\vec{e}_j \cdot \vec{e}_k = \delta_{jk}$. Dann kann man die Bilinearform als Abbildung $\mathbb{R}^d \times \mathbb{R}^d$ der entsprechenden Komponenten der Vektoren auffassen, denn es gilt wegen der Bilinearität

$$B(\vec{v}, \vec{w}) = \sum_{j,k=1}^d B(v_j \vec{e}_j, w_k \vec{e}_k) = \sum_{j,k=1}^d B(\vec{e}_j, \vec{e}_k) v_j w_k. \quad (2.9.43)$$

Führen wir also die Matrix $\hat{M} = (M_{jk})$ mittels

$$M_{jk} = B(\vec{e}_j, \vec{e}_k) \quad (2.9.44)$$

ein, gilt

$$B(\vec{v}, \vec{w}) = \sum_{j,k=1}^d M_{jk} v_j w_k. \quad (2.9.45)$$

Mit den Spaltenvektoren $\underline{v} = (v_j)$ und $\underline{w} = (w_j)$ können wir dies als

$$B(\vec{v}, \vec{w}) = \underline{v}^T \hat{M} \underline{w} = \underline{v} \cdot (\hat{M} \underline{w}) \quad (2.9.46)$$

schreiben. Wobei der Punkt $\cdot$ das übliche Skalarprodukt im $\mathbb{R}^d$ bezeichnet, also $\underline{v} \cdot \underline{w} = \underline{v}^T \underline{w} = v_1 w_1 + v_2 w_2 + \dots + v_d w_d$.

Da B symmetrisch ist, gilt

$$M_{jk} = B(\vec{e}_j, \vec{e}_k) = B(\vec{e}_k, \vec{e}_j) = M_{kj} \Rightarrow \hat{M} = \hat{M}^T \quad (2.9.47)$$

und damit

$$B(\vec{v}, \vec{w}) = \underline{v}^T \hat{M} \underline{w} \underline{v} \cdot (\hat{M} \underline{w}) = \underline{v}^T \hat{M}^T \underline{w} = (\hat{M} \underline{v})^T \underline{w} = (\hat{M} \underline{v}) \cdot \underline{w}. \quad (2.9.48)$$

Um die obige Behauptung über die Eigenvektoren und Eigenwerte von $\hat{M}$ zu beweisen, zeigen wir zunächst, dass die Eigenvektoren zu verschiedenen Eigenwerten notwendig orthogonal sind. Seien also $\underline{u}_1$ und $\underline{u}_2$ Eigenvektoren zu den *voneinander verschiedenen* Eigenwerten λ_1 und λ_2 . Dann folgt wegen (2.9.48) einerseits

$$B(\vec{u}_1, \vec{u}_2) = \underline{u}_1 \cdot (\hat{M} \underline{u}_2) = \underline{u}_1 \cdot (\lambda_2 \underline{u}_2) = \lambda_2 \underline{u}_1 \cdot \underline{u}_2. \quad (2.9.49)$$

Andererseits ist wegen (2.9.48)

$$B(\vec{u}_1, \vec{u}_2) = B(\vec{u}_2, \vec{u}_1) = \underline{u}_2 \cdot (\hat{M} \underline{u}_1) = \lambda_1 \underline{u}_2 \cdot \underline{u}_1 = \lambda_1 \underline{u}_1 \cdot \underline{u}_2. \quad (2.9.50)$$

Ziehen wir nun diese Gleichung von (2.9.49) ab, folgt

$$(\lambda_2 - \lambda_1) \underline{u}_1 \cdot \underline{u}_2 = 0. \quad (2.9.51)$$

Da voraussetzungsgemäß $\lambda_2 \neq \lambda_1$ ist, können wir durch $(\lambda_2 - \lambda_1) \neq 0$ teilen, und es folgt $\vec{u}_1 \cdot \vec{u}_2 = 0$, d.h. $\underline{u}_1$ und $\underline{u}_2$ sind zueinander orthogonal. Wir können ohne Beschränkung der Allgemeinheit annehmen, dass $|\underline{u}_1| = |\underline{u}_2| = 1$ ist.

Sind nun alle Nullstellen des charakteristischen Polynoms reell und voneinander verschieden, ist damit die Behauptung schon gezeigt, denn offenbar gibt es dann genau d zueinander orthogonale normierte Eigenvektoren $\underline{u}_k$ ($k \in \{1, \dots, d\}$), und diese bilden ein Orthonormalsystem von Eigenvektoren. Für die Basistransformation von den kanonischen Einheitsvektoren $\underline{e}_j$ zu diesem Orthonormalsystem von Eigenvektoren gilt

$$\vec{u}_j = \sum_{k=1}^d U_{kj} \vec{e}_k \Rightarrow \hat{U} = (\underline{u}_1, \underline{u}_2, \dots, \underline{u}_d), \quad (2.9.52)$$

d.h. die Transformationsmatrix wird durch die Spaltenvektoren aus den Komponenten der Eigenvektoren von $\hat{M}$ bzgl. der kanonischen Basisvektoren $\underline{e}_j$, gebildet.

Da nun sowohl die $\vec{e}_k$ als auch die $\vec{u}_j$ kartesische Basissysteme sind, ist $\hat{U}$ demnach eine **orthogonale Transformationsmatrix**, d.h.

$$\hat{U}^{-1} \hat{M} \hat{U} = \hat{U}^T \hat{M} \hat{U} = \text{diag}(\lambda_1, \dots, \lambda_d) \quad (2.9.53)$$

eine Diagonalmatrix.

Zu der Transformationsformel gelangt man wie folgt: Es gilt

$$\begin{aligned} B(\vec{v}, \vec{w}) &= \sum_{j,k} B(v_j \vec{e}_j, w_k \vec{e}_k) = \sum_{j,k} v_j w_k B(\vec{e}_j, \vec{e}_k) = \sum_{j,k} M_{jk} v_j w_k = \underline{v}^T \hat{M} \underline{w} \\ &= \sum_{j,k} B(v'_j \vec{u}_j, w'_k \vec{u}_k) = \sum_{j,k} v'_j w'_k B(\vec{u}_j, \vec{u}_k) = \sum_{j,k} M'_{jk} v'_j w'_k = \underline{v}'^T \hat{M}' \underline{w}'. \end{aligned} \quad (2.9.54)$$

Weiter ist für beliebige Vektoren $\vec{x}$

$$\vec{x} = \sum_j x'_j \vec{u}_j = \sum_{j,k} \vec{e}_k U_{kj} x'_j \Rightarrow \underline{x} = \hat{U} \underline{x}' \Leftrightarrow \underline{x}' = \hat{U}^{-1} \underline{x} = \hat{U}^T \underline{x}. \quad (2.9.55)$$

Damit folgt aus (2.9.54)

$$B(\vec{v}, \vec{w}) = \underline{v}'^T \hat{M}' \underline{w} = \underline{v}^T \hat{M} \hat{w} = \underline{v}'^T \hat{U}^T \hat{M} \hat{U} \hat{w}' = \underline{v}'^T (\hat{U}^T \hat{M} \hat{U}) \hat{w}'. \quad (2.9.56)$$

Da diese Gleichung für alle $\vec{v}$ und $\vec{w}$ gilt, folgt (2.9.53).

Wir müssen nun noch zeigen, dass alle Eigenwerte reell sind und auch im Fall, dass nicht alle Eigenwerte verschieden sind, ein vollständiges Orthonormalsystem von Eigenvektoren existiert. Dazu definieren wir die **quadratische Form**

$$Q(\vec{v}) = B(\vec{v}, \vec{v}) = \underline{v} \cdot (\hat{M} \underline{v}). \quad (2.9.57)$$

Wir betrachten nun alle Vektoren $\underline{v}$ auf der Kugelschale $K : \underline{v} \cdot \underline{v} = 1$. Dies ist offenbar eine kompakte Menge im $\mathbb{R}^d$. Da weiter $Q(\vec{v})$ eine stetige Funktion der Komponenten v_k ist, muss es unter dieser Einschränkung einen Vektor $\vec{u}_1$ mit $|\vec{u}_1| = 1$ geben, für den $Q(\vec{v})|_{\vec{v} \in K}$ maximal wird.

Offensichtlich ist nun die quadratische Form auch stetig nach allen Komponenten von $\vec{v}$ differenzierbar. Um das Maximum unter allen Vektoren mit $\vec{v}^2 = 1$ zu finden, führen wir den entsprechenden Lagrange-Parameter λ_1 ein und betrachten die Funktion

$$f(\vec{v}, \lambda) = Q(\vec{v}) - \lambda_1(\vec{v}^2 - 1). \quad (2.9.58)$$

Es ist klar, dass für das Maximum dieser Funktion²

$$\left. \frac{\partial f}{\partial v_j} \right|_{\vec{v}=\vec{u}_1} = 2 \sum_k M_{jk} u_{1k} - 2\lambda_1 u_{1j} \stackrel{!}{=} 0, \quad (2.9.59)$$

$$\frac{\partial}{\partial \lambda_1} f(\vec{u}_1, \lambda_1) = 0 \Rightarrow \vec{u}_1^2 - 1 = 0 \quad (2.9.60)$$

gelten muss. Die Gl. (2.9.59) können wir nun in Matrix-Vektorschreibweise als die Eigenwertgleichung

$$\hat{M} \underline{u}_1 = \lambda_1 \underline{u}_1 \quad (2.9.61)$$

²Das Symbol $\partial f(\vec{v}, \lambda_1) / \partial v_j$ bezeichnet die **partielle Ableitung** nach v_j . Dabei soll definitionsgemäß die Ableitung nach v_j gebildet werden, wobei alle anderen v_k (mit $k \neq j$) und λ_1 als Konstanten betrachtet werden. Entsprechend ist bei der partiellen Ableitung $\partial f(\vec{v}, \lambda_1) / \partial \lambda_1$ beim Ableiten $\vec{v} = \text{const}$ zu denken.

lesen. Da nach der obigen Überlegung unter der Nebenbedingung $\vec{v}^2 = 1$ die quadratische Form (2.9.57) ein Maximum besitzen muss, besitzt $\hat{M}$ demnach einen Eigenvektor $\underline{u}_1 \in \mathbb{R}^d$ mit $\underline{u}_1^2 = 1$ mit einem Eigenwert $\lambda_1 \in \mathbb{R}$.

Nun können wir eine Orthonormalbasis $\vec{e}'_k$ bilden, so dass $\vec{e}'_1 = \vec{u}_1$ ist. Bzgl. dieser neuen Basis besitzt offenbar die entsprechend transformierte symmetrische Matrix $\hat{M}'$ wegen $\hat{M}'\vec{u}'_1 = \hat{M}\vec{u}_1 = \lambda_1\vec{u}_1$ die Gestalt

$$\hat{M}' = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & M'_{22} & \cdots & M'_{2d} \\ 0 & \vdots & \ddots & \vdots \\ 0 & M'_{d2} & \cdots & M'_{dd} \end{pmatrix} = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \hat{M}'^{(d-1)} \end{pmatrix}. \quad (2.9.62)$$

Betrachten wir nun die quadratische Form

$$Q_2(v'_2, \dots, v'_d) = \begin{pmatrix} 0 \\ v'_2 \\ \vdots \\ v'_d \end{pmatrix} \cdot \hat{M}' \begin{pmatrix} 0 \\ v'_2 \\ \vdots \\ v'_d \end{pmatrix} = \begin{pmatrix} v'_2 \\ \vdots \\ v'_d \end{pmatrix} \cdot \hat{M}'^{(d-1)} \begin{pmatrix} v'_2 \\ \vdots \\ v'_d \end{pmatrix}, \quad (2.9.63)$$

können wir auf die quadratische Form in dem $(d-1)$ -dimensionalen Untervektorraum die gleichen Argumente bzgl. der Existenz eines Maximums unter der Einschränkung, dass $(0, v'_2, v'_3, \dots, v'_d)$ ein Einheitsvektor sein soll, wie oben für die Matrix $\hat{M}$ anwenden, d.h. wir erhalten einen weiteren reellen Eigenwert λ_2 und einen zu $\underline{u}_1 = (1, \dots, 0)$ orthogonalen Einheitseigenvektor $\underline{u}_2$ zu diesem Eigenwert.

Dieses Verfahren können wir nun weiterführen, bis wir insgesamt d reelle Eigenwerte λ_j und dazugehörige orthonormierte Eigenvektoren $\vec{u}_1$ gefunden haben. Damit ist also die oben aufgestellte Behauptung bewiesen:

Satz von der Hauptachsentransformation

Eine symmetrische Matrix $\hat{M} \in \mathbb{R}^{d \times d}$ besitzt stets ein Orthonormalsystem von Eigenvektoren $\underline{u}_j \in \mathbb{R}^d$ ($j \in \{1, 2, \dots, d\}$) zu reellen Eigenwerten $\lambda_j \in \mathbb{R}$. Demnach ist also jede symmetrische reelle Matrix mittels einer orthogonalen Transformation diagonalisierbar, d.h. es gibt eine Transformationsmatrix $\hat{U} \in O(d)$ mit

$$\hat{M}' = \hat{U} \hat{M} \hat{U} = \text{diag}(\lambda_1, \dots, \lambda_d). \quad (2.9.64)$$

Gemäß (2.9.52) ist dabei $\hat{U} = (\underline{u}_1, \dots, \underline{u}_d)$, d.h. die Spalten von $\hat{U}$ sind durch die orthonormalen Spaltenvektoren $\underline{u}_j$ bestimmt. Man nennt die entsprechende Transformation auch eine **Hauptachsentransformation** der symmetrischen Matrix $\hat{M}$ bzw. der symmetrischen Bilinearform (2.9.46).

Es ist klar, dass man die Hauptachsentransformation in konkreten Fällen nicht durch die Lösung der Extremwertaufgabe für die quadratische Form findet sondern, wie im vorigen Abschnitt für allgemeine Matrizen gezeigt, indem man die Eigenwerte als die Nullstellen des charakteristischen Polynoms von $\hat{M}$ findet und dann für jeden Eigenwert die dazugehörigen Eigenvektoren berechnet. Unser obiger Beweis besagt dann, dass wir stets reelle Eigenwerte haben und die Eigenvektoren als normierte zueinander orthogonale Eigenvektoren gewählt werden können, d.h. zusammengenommen eine Orthonormalbasis des $\mathbb{R}^d$ bilden.

Beispiel: Wir wollen das Eigenwertproblem für die Matrix

$$\hat{M} = \begin{pmatrix} 2 & 1 & 1 \\ 1 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix} \quad (2.9.65)$$

lösen. Da die Matrix offensichtlich symmetrisch ist, muss es nach den obigen allgemeinen Überlegungen eine orthogonale Transformation geben, die die Matrix diagonalisiert. Um das konkret zu sehen, berechnen wir zunächst die (notwendig reellen!) Eigenwerte als die Nullstellen des charakteristischen Polynoms

$$P_{\hat{M}}(\lambda) = \det(\hat{M} - \lambda \mathbb{1}) = -\lambda^3 + 6\lambda^2 - 9\lambda + 4. \quad (2.9.66)$$

Durch Probieren findet man, dass $\lambda_1 = 1$ ein Eigenwert ist. Durch Polynomdivision durch den entsprechenden Linearfaktor $(\lambda - 1)$ und einige Umformungen finden wir (*Nachrechnen!*)

$$P_{\hat{M}}(\lambda) = -(\lambda - 1)^2(\lambda - 4). \quad (2.9.67)$$

Das bedeutet, dass das Polynom eine doppelte Nullstelle $\lambda_1 = \lambda_2 = 1$ und eine einfache Nullstelle $\lambda_3 = 4$ besitzt. Wir wissen, dass die Matrix mittels einer orthogonalen Transformation diagonalisierbar ist, d.h. zum Eigenwert 1 müssen wir zwei zueinander orthogonale Einheitseigenvektoren finden. Dazu müssen wir das homogene lineare Gleichungssystem

$$(\hat{M} - \mathbb{1})\vec{u} = 0 \Rightarrow \hat{M} = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ u_3 \end{pmatrix} = 0 \quad (2.9.68)$$

lösen. Offenbar ergibt sich in der Tat für alle drei Komponenten dieselbe Gleichung

$$u_1 + u_2 + u_3 = 0 \Rightarrow u_3 = -(u_1 + u_2). \quad (2.9.69)$$

Eine Lösung ist also

$$\vec{u}_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}. \quad (2.9.70)$$

Dabei habe wir den Vektor auch gleich normiert. Als zweite Lösung muss es einen dazu orthogonalen Vektor geben, also

$$\vec{u}_2 = \begin{pmatrix} v_1 \\ v_2 \\ -(v_1 + v_2) \end{pmatrix} \quad \text{mit} \quad \begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix} \cdot \begin{pmatrix} v_1 \\ v_2 \\ -(v_1 + v_2) \end{pmatrix} = 2v_1 + v_2 = 0. \quad (2.9.71)$$

Damit ist die (bis auf das Vorzeichen eindeutige) Lösung für $\vec{u}_2$

$$\vec{u}_2 = \frac{1}{\sqrt{6}} \begin{pmatrix} -1 \\ 2 \\ -1 \end{pmatrix}. \quad (2.9.72)$$

Es ist dann klar, dass der verbliebene Eigenvektor zum Eigenwert $\lambda_3 = 4$ orthogonal zu beiden Vektoren $\vec{u}_1$ und $\vec{u}_2$ sein muss. Um ein *rechtshändiges* Orthonormalsystem aus Eigenvektoren zu erhalten, setzen wir daher

$$\vec{u}_3 = \vec{u}_1 \times \vec{u}_2 = \frac{1}{\sqrt{3}} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}. \quad (2.9.73)$$

In der Tat zeigt man durch direkte Rechnung, dass $\hat{M}\vec{u}_3 = 4\vec{u}_3$ ist. Die Transformationsmatrix ist durch

$$\hat{U} = (\vec{u}_1, \vec{u}_2, \vec{u}_3) = \begin{pmatrix} 1/\sqrt{2} & -1/\sqrt{6} & 1/\sqrt{3} \\ 0 & 2/\sqrt{6} & 1/\sqrt{3} \\ -1/\sqrt{2} & -1/\sqrt{6} & 1/\sqrt{3} \end{pmatrix} \quad (2.9.74)$$

gegeben. Man rechnet nach, dass damit in der Tat

$$\hat{M}' = \hat{U}^T \hat{M} \hat{U} = \text{diag}(1, 1, 4) = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 4 \end{pmatrix} \quad (2.9.75)$$

ist, also mit der orthogonalen Transformationsmatrix (2.9.74) die Matrix $\hat{M}$ gemäß (2.9.53) diagonalisiert wird.

2.10 Vektoren und Tensoren

In diesem Abschnitt wollen wir kurz die allgemeine Algebra von Vektoren und Tensoren besprechen. Dabei können wir einen euklidischen Vektorraum beliebiger Dimension d betrachten, ohne dass dies gegenüber der Beschränkung auf $d = 2$ oder $d = 3$ zu größeren Schwierigkeiten führt. Wir verwenden in diesem Abschnitt stets die **Einsteinsche Summenkonvention** und führen jetzt auch sog. **kovariante und kontravariante Indizes** ein, die wir durch tiefgestellte bzw. hochgestellte Schreibweise unterscheiden. Dieser **Ricci-Kalkül** (Gregorio Ricci-Curbastro, 1853-1925) erleichtert das Rechnen mit Komponenten der diversen algebraischen Objekte bzgl. verschiedener Basen und die Transformationen unter Basiswechseln erheblich.

Wir präzisieren also die bisherige Schreibweise von Vektoren, Basen und Vektorkomponenten wie folgt. Es sei $\vec{b}_j$ mit $j \in \{1, 2, \dots, d\}$ eine beliebige Basis. Die Basisvektoren nummerieren wir also wie bisher mit einem tiefgestellten Index durch. Sei weiter $\vec{b}'_k$ eine beliebige andere Basis. Dann schreiben wir die Transformation zwischen den Basen in der Form

$$\vec{b}_j = T^k_j \vec{b}'_k. \quad (2.10.1)$$

Entsprechend dem allgemeinen Ricci-Kalkül wird stets über doppelt auftretende Indizes summiert, wobei immer einer dieser Indizes oben und der andere unten notiert sein muss. Man sagt, dass sich die Objekte mit unteren Indizes **kovariant transformieren**, also eben wie die Basisvektoren.

Die entsprechende Umkehrtransformation ist durch die inverse Transformationsmatrix U^j_k gegeben:

$$\vec{b}'_k = U^j_k \vec{b}_j. \quad (2.10.2)$$

Mit (2.10.1) folgt

$$\vec{b}'_k = U^j_k T^l_j \vec{b}'_l = T^l_j U^j_k \vec{b}'_l = \delta^l_k \vec{b}'_l. \quad (2.10.3)$$

Dabei ist

$$\delta^l_k = \delta^k_l = \begin{cases} 1 & \text{falls } k = l, \\ 0 & \text{falls } k \neq l \end{cases} \quad (2.10.4)$$

das **Kronecker-Symbol** (Leopold Kronecker, 1823-1891). Da die Basisvektoren linear unabhängig sind, ist die Entwicklung beliebiger Vektoren nach den Basisvektoren eindeutig und daher gemäß (2.10.3)

$$T^l_j U^j_k = \delta^l_k. \quad (2.10.5)$$

Das bedeutet, dass die Matrizen $\hat{T} = (T^l_j)$ und $\hat{U} = (U^j_k)$ zueinander inverse Matrizen sind, d.h. es gilt

$$\hat{T} \hat{U} = \mathbb{1}, \quad (2.10.6)$$

wobei $\mathbb{1} = \text{diag}(1, \dots, 1)$ die **Einheitsmatrix** ist, die nur auf der Diagonalen 1 als Matrixelemente besitzt und alle anderen Matrixelemente verschwinden.

Mit (2.10.1) und (2.10.2) folgt

$$\vec{b}_j = T^k_j \vec{b}'_k = T^k_j U^l_k \vec{b}_l = U^l_k T^k_j \vec{b}_l = \delta^l_j \vec{b}_l \quad (2.10.7)$$

und damit

$$U^l_k T^k_j = \delta^l_j \Rightarrow \hat{U} \hat{T} = \mathbb{1}. \quad (2.10.8)$$

Das bedeutet, dass dass $\hat{T} \hat{U} = \mathbb{1}$ auch $\hat{U} \hat{T} = \mathbb{1}$ ist, obwohl i.a. das Matrixprodukt *nicht kommutativ* ist.

Es ist nun auch einfach, die Transformation der **Vektorkomponenten** zu bestimmen. Es gilt

$$\vec{v} = v'^k \vec{b}'_k = v'^k U^j_k \vec{b}_j = v^j \vec{b}_j, \quad (2.10.9)$$

und wegen der Eindeutigkeit der Zerlegung von Vektoren nach einer Basis folgt

$$v^j = U^j_k v'^k. \quad (2.10.10)$$

Ebenso gilt

$$\vec{v} = v^j \vec{b}_j = v^j T^k_j \vec{b}'_k = v'^k \vec{b}'_k \quad (2.10.11)$$

und damit

$$v'^k = T^k_j v^j. \quad (2.10.12)$$

Man sagt, dass sich die Vektorkomponenten **kontravariant** transformieren bzw. im Vergleich zur Transformationsvorschrift für die Basisvektoren (2.10.2) **kontragredient** zum Transformationsverhalten der Basisvektoren.

Als weitere algebraische Objekte führt man nun die **Linearformen** oder **Tensoren 1. Stufe** ein. Ein solcher Tensor $\underline{L}$ ist eine Abbildung $\underline{L} : V \rightarrow \mathbb{R}$, $\vec{v} \mapsto \underline{L} \vec{v}$ die linear auf den Vektoren operiert, d.h. es gilt für beliebige $\vec{v}_1, \vec{v}_2 \in V$ und $\lambda_1, \lambda_2 \in \mathbb{R}$

$$\underline{L}(\lambda_1 \vec{v}_1 + \lambda_2 \vec{v}_2) = \lambda_1 \underline{L} \vec{v}_1 + \lambda_2 \underline{L} \vec{v}_2. \quad (2.10.13)$$

Wegen der Linearität der Abbildung ist es klar, dass die Linearform bereits vollständig durch die Komponenten

$$L_j = \underline{L} \vec{b}_j \quad (2.10.14)$$

bestimmt ist und sich diese Komponenten kovariant gemäß

$$L'_k = \underline{L} \vec{b}'_k = \underline{L}(U^j_l \vec{b}_j) = U^j_l \underline{L} \vec{b}_j = U^j_l L_j \quad (2.10.15)$$

$$L_j = \underline{L} \vec{b}_j = \underline{L}(T^k_j \vec{b}'_k) = T^k_l \underline{L} \vec{b}'_k = T^k_l L'_k. \quad (2.10.16)$$

transformieren.

Weiter gilt

$$\underline{L} \vec{v} = \underline{L}(v^j \vec{b}_j) = L_j v^j = L'_k v'^k. \quad (2.10.17)$$

Die Gleichheit der beiden Ausdrücke folgt auch direktaus dem Transformationsverhalten der Vektor- bzw. Tensorkomponenten (2.10.12) und (2.10.15):

$$L'_k v'^k = U^j_k L_j T^k_l v'^l = L_j U^j_k T^k_l v'^l = L_j \delta^j_l v'^l = L_j v'^j. \quad (2.10.18)$$

Die Linearformen $\underline{L}$ bilden nun selbst auch einen d -dimensionalen Vektorraum, wenn man die Addition und Multiplikation mit Skalaren durch

$$(\lambda_1 \underline{L}_1 + \lambda_2 \underline{L}_2) \vec{v} = \lambda_1 \underline{L}_1 \vec{v} + \lambda_2 \underline{L}_2 \vec{v} \quad (2.10.19)$$

definiert. Man nennt diesen d -dimensionalen Vektorraum den zu V **dualen Vektorraum** V^* .

Entsprechend kann man zur Basis $\vec{b}_j$ des Vektorraums die **duale Basis** $\underline{b}^j$ einführen. Wir werden gleich sehen, dass die Stellung als oberer Index durch das entsprechende Transformationsverhalten unter Basiswechseln gerechtfertigt ist. Wir definieren diese duale Basis durch die Wirkung auf die Basisvektoren

$$\underline{b}^j \vec{b}_k = \delta_k^j. \quad (2.10.20)$$

Dann ist für eine beliebige Linearform

$$\underline{L} = L_j \underline{b}^j, \quad (2.10.21)$$

denn es gilt

$$L_j \underline{b}^j \vec{b}_k = L_j \delta_k^j = L_k = \underline{L} \vec{b}_k \quad (2.10.22)$$

Da die Linearform durch ihre Wirkung auf die basis $\vec{b}_j$ eindeutig bestimmt ist, gilt tatsächlich (2.10.20).

Bestimmen wir nun die Transformation der dualen Basis. Es gilt

$$\delta_l^k = \underline{b}^{'k} \vec{b}_l' = \underline{b}^{'k} U_l^j \vec{b}_j = U_l^j \underline{b}^{'k} \vec{b}_j. \quad (2.10.23)$$

Da $\hat{T} = \hat{U}^{-1}$ ist, folgt

$$\underline{b}^{'k} \vec{b}_j = T^k_j \quad (2.10.24)$$

und damit unter Verwendung von (2.10.14) und (2.10.20) für $\underline{L} = \underline{b}^{'k}$

$$\underline{b}^{'k} = T^k_j \underline{b}^j, \quad (2.10.25)$$

d.h. die dualen Basisvektoren transformieren sich entsprechend der Indexstellung kontravariant.

Dieselben Betrachtungen lassen sich nun auch für **Multilinearformen** anwenden, d.h. Abbildungen $\underline{T} : V^n \rightarrow \mathbb{R}$, die in allen n Argumenten linear operieren. Für eine gegebene Basis ist ein solcher **Tensor** n -ter Stufe eindeutig durch die entsprechenden Tensorkomponenten

$$T_{j_1 j_2 \dots j_n} = \underline{T}(\vec{b}_{j_1}, \vec{b}_{j_2}, \dots, \vec{b}_{j_n}) \quad (2.10.26)$$

eindeutig bestimmt.

Weiter kann man aus einem Tensor $\underline{T}$ n -ter Stufe und einem Tensor $\underline{U}$ m -ter Stufe das sog. **Kronecker-Produkt**, einen Tensor $(n+m)$ -ter Stufe definieren:

$$\underline{T} \otimes \underline{U}(\vec{v}_1, \dots, \vec{v}_n, \vec{v}_{n+1}, \dots, \vec{v}_{n+k}) = \underline{T}(\vec{v}_1, \dots, \vec{v}_n) \underline{U}(\vec{v}_{n+1}, \dots, \vec{v}_{n+k}). \quad (2.10.27)$$

Man macht sich dann schnell klar, dass sich aus den Tensorkomponenten (2.10.26) der Tensor durch

$$\underline{T} = T_{j_1 \dots j_n} \underline{b}^{j_1} \otimes \dots \otimes \underline{b}^{j_n} \quad (2.10.28)$$

ergibt.

Ebenso transformieren sich die Tensorkomponenten entsprechend ihrer Indexstellung kovariant, d.h.

$$T'_{k_1 \dots k_n} = U^{j_1}_{k_1} \dots U^{j_n}_{k_n} T_{j_1 \dots j_n}. \quad (2.10.29)$$

Weiter können wir auch Linearformen auf dem Dualraum V^* betrachten, also lineare Abbildungen $\tilde{M} : V^* \rightarrow \mathbb{R}$. Diese bilden wieder einen d -dimensionalen Vektorraum V^{**} . Wir können nun diesen **Bidualraum** auf

„kanonische Weise“, also ohne Rückgriff auf eine spezielle Basis, mit dem ursprünglichen Vektorraum V selbst identifizieren. Der entsprechende **Isomorphismus** ist durch die Abbildung $\Psi: V \rightarrow V^{**}$, $\vec{v} \mapsto \tilde{M}_{\vec{v}}$ mit

$$\tilde{M}_{\vec{v}} \underline{L} = \underline{L} \vec{v}. \quad (2.10.30)$$

Im Sinne dieser Identifikation von $\vec{v} \in V$ mit $\tilde{M}_{\vec{v}}$ können wir also auch einfach schreiben

$$\vec{v} \underline{L} := \tilde{M}_{\vec{v}} \underline{L} = \underline{L} \vec{v}. \quad (2.10.31)$$

Weiter können wir damit den Tensorbegriff erweitern zu Tensoren vom Rang (p, q) , als Linearformen der Art $\overrightarrow{T}: V^{*p} \times V^q \rightarrow \mathbb{R}$. Die entsprechenden Komponenten sind durch

$$T_{j_1 \dots j_q}^{k_1 \dots k_p} = \overrightarrow{T}(\underline{b}^{j_1}, \dots, \underline{b}^{j_q}; \vec{b}_{k_1}, \dots, \vec{b}_{k_p}) \quad (2.10.32)$$

definiert, und diese transformieren sich entsprechend ihrer Indexstellung (kovariant für untere und kontravariant für obere Indizes).

Man macht sich auch sofort klar, dass man aus einem Tensor vom Rang (p, q) durch **Kontraktion** einen Tensor vom Rang $(p-1, q-1)$ machen kann, indem man die entsprechenden Komponenten durch

$$U_{j_2 \dots j_q}^{k_1 \dots k_p} = T_{j_1 j_2 \dots j_q}^{j_1 k_2 \dots k_p} \quad (2.10.33)$$

definiert. Man beachte, dass hierbei gemäß der Einsteinschen Summenkonvention über j_1 zu summieren ist. Die bis jetzt besprochenen algebraischen Manipulationen mit Tensoren erfordern offensichtlich nur einen endlichdimensionalen Vektorraum. Für einen Euklidischen Vektorraum haben wir noch das **Skalarprodukt** als zusätzliche algebraische Struktur. Wie oben bereits erörtert, handelt es sich um einen Tensor $\underline{g}$ vom Rang 2 (bzw. $(0, 2)$), der zusätzlich **symmetrisch** ist, d.h. für den $\underline{g}(\vec{v}, \vec{w}) = \underline{g}(\vec{w}, \vec{v})$ gilt. Weiter soll er auch **positiv definit sein**, d.h. es gilt stets $\underline{g}(\vec{v}, \vec{v}) \geq 0$ und aus $\underline{g}(\vec{v}, \vec{v}) = 0$ folgt zwingend $\vec{v} = \vec{0}$. Wir schreiben, wie bisher $\underline{g}(\vec{v}, \vec{w}) = \vec{v} \cdot \vec{w}$.

Mit Hilfe dieses ausgezeichneten Tensors 2. Stufe können wir nun auch kanonisch (d.h. basisunabhängig) Vektoren und Linearformen identifizieren. Der entsprechende Isomorphismus ist durch $\vec{v} \mapsto \underline{L}_{\vec{v}}$ definiert, für den

$$\underline{L}_{\vec{v}} \vec{w} = \vec{v} \cdot \vec{w}. \quad (2.10.34)$$

Im Sinne dieser Identifikation von Vektoren und Linearformen können wir entsprechend Indizes von oben nach unten ziehen:

$$v_j = \underline{L}_{\vec{v}} \vec{b}_j = \vec{v} \cdot \vec{b}_j = v^k \vec{b}_k \cdot \vec{b}_j = g_{kj} v^k = g_{jk} v^j. \quad (2.10.35)$$

Wie wir im vorigen Abschnitt gesehen haben, kann man die Matrix $\overleftarrow{g} = (g_{jk})$ diagonalisieren, d.h. es gibt eine Basis $\underline{e}_j$ von Eigenvektoren dieser Matrix, für die

$$\overleftarrow{g} \underline{e}_j = \lambda_j \underline{e}_j \quad (2.10.36)$$

gilt. Da die entsprechenden Vektoren $\vec{e}_j \neq \vec{0}$ sind, folgt aus der positiven Definitheit des Skalarprodukts

$$0 < \vec{e}_j \cdot \vec{e}_j = \underline{e}_j^T \overleftarrow{g} \underline{e}_j = \lambda_j \underline{e}_j^T \underline{e}_j \Rightarrow \lambda_j > 0. \quad (2.10.37)$$

Demnach ist $\det \overleftarrow{g} = \lambda_1 \cdots \lambda_d > 0$, und damit ist $\overleftarrow{g}$ invertierbar, und wir bezeichnen die Matrixelemente der entsprechenden inversen Matrix mit oberen Indizes:

$$(g^{jk}) = (g^{kj}) = \overleftarrow{g}^{-1}. \quad (2.10.38)$$

Wir können dann den obigen Isomorphismus durch

$$\vec{v} = v^j \vec{b}_j \mapsto \underline{L} \vec{v} = v_j \vec{b}^j = g_{jk} v^k \underline{b}^j \Rightarrow v_j = g_{jk} v^k \quad (2.10.39)$$

bzw. für die Umkehrabbildung

$$\underline{L} = L_j \underline{b}^j \mapsto \vec{v}_L = L^j \vec{b}_j = g^{jk} L_k \vec{b}_k \Rightarrow L^j = g^{jk} L_k. \quad (2.10.40)$$

Man kann also obere (kontravariante) Indizes mit Hilfe der Matrixelemente g_{jk} nach unten ziehen, was kovariante Indizes ergibt bzw. umgekehrt untere (kovariante) Indizes mittels der Matrixelemente g^{jk} nach oben ziehen, wodurch sich kontravariante Indizes ergeben. Die entsprechenden Vektoren bzw. Dualvektoren sind aufgrund der entsprechenden Transformationseigenschaften dieser Objekte unabhängig von der gewählten Basis, d.h. es ist z.B.

$$\underline{L} = L_j \underline{b}^j = L^j \rightarrow b_j \quad (2.10.41)$$

wobei wir die entsprechenden Vektoren und Dualvektoren im Sinne des obigen Isomorphismusses einfach identifizieren.

Kapitel 3

Vektoranalysis

3.1 Kurven in der Ebene

Als erstes betrachten wir Kurven in der Ebene und im Raum, die wir durch beliebige Funktionen parametrisieren können. Diese Funktionen können wir dann ggf. differenzieren und integrieren, so dass die Methoden der Analysis für die Analyse von Kurven bzw. von Bahnkurven von Massenpunkten als Funktionen der Zeit (**Trajektorien**) zur Verfügung stehen. Wir gehen davon aus, dass die Begriffe der Ableitung und des Integrals von Funktionen von der Schule her bekannt sind.

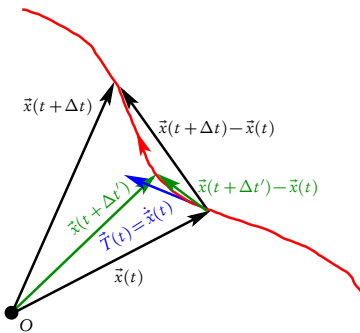
Wir betrachten als erstes Kurven in der Euklidischen Ebene $\mathbb{R}^2$. Nach Festlegung des Koordinatenursprungs können sie als Funktion der Ortsvektoren der Punkte auf der Kurve $\underline{x} : \mathbb{R} \supseteq (t_0, t_1) \rightarrow \mathbb{R}^2$ dargestellt werden. Im Folgenden nehmen wir an, diese Funktion sei im ganzen Intervall (t_0, t_1) stetig nach dem Parameter differenzierbar, d.h. dass beide Vektorkomponenten in allen Punkten $t \in (t_0, t_1)$ Ableitungen besitzen und die Ableitungsfunktion dort stetig ist. Die Ableitung ist dabei durch den Grenzwert

$$\frac{d}{dt} \underline{x}(t) = \dot{\underline{x}} = \lim_{\Delta t \rightarrow 0} \frac{\underline{x}(t + \Delta t) - \underline{x}(t)}{\Delta t} \quad (3.1.1)$$

definiert. Es ist klar, dass wir diesen Vektor wieder in den Raum der geometrischen Vektoren abbilden können:

$$\frac{d}{dt} \vec{x}(t) = \dot{\vec{x}}(t) = \sum_{j=1}^3 \dot{x}_j(t) \vec{e}_j = \frac{d}{dt} \sum_{j=1}^3 x_j(t) \vec{e}_j. \quad (3.1.2)$$

Letzteres gilt, da die Basisvektoren $(\vec{e}_j)$ fest vorgegebene konstante Vektoren sind und folglich nicht von t abhängen.



Jetzt wollen wir die lokalen Eigenschaften einer solchen Kurve durch geometrische vom Koordinatensystem unabhängige Größen charakterisieren. An einem Punkt t sei nun

$$\frac{d\vec{x}}{dt} \neq 0. \quad (3.1.3)$$

Wir nennen den dazugehörigen Punkt $\vec{x}(t)$ auf der Kurve einen **regulären Punkt**.

Offensichtlich besitzt der Vektor $d\vec{x}/dt$ in jedem regulären Punkt der Kurve die geometrische Bedeutung eines **Tangentenvektors**.

Dies zeigt die nebenstehende Skizze: Für endliche Δt ist der Zähler in (3.1.1) offenbar der Verbindungsvektor¹ zwischen den Punkten $\vec{x}(t + \Delta t)$ und $\vec{x}(t)$. Lassen wir Δt kleiner werden, nähert sich die entsprechende Sekante immer mehr der Tangente an die Kurve im Punkt $\vec{x}(t)$, vorausgesetzt die Kurve ist differenzierbar.

Wir definieren nun durch

$$\vec{T} = \left| \frac{d\vec{x}}{dt} \right|^{-1} \frac{d\vec{x}}{dt} \quad (3.1.4)$$

die **Tangenteneinheitsvektoren** an die Kurve, die an jedem regulären Punkt wohldefiniert sind.

Leiten wir die Identität $\vec{T}^2 = 1$ nach dem Parameter ab, erhalten wir

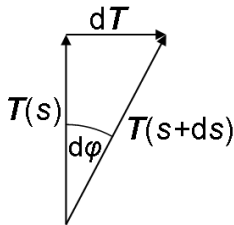
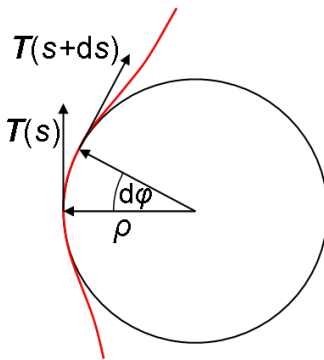
$$\vec{T} \cdot \frac{d\vec{T}}{dt} = 0. \quad (3.1.5)$$

Dies bedeutet, dass entweder $d\vec{T}/dt$ ein auf dem Tangentenvektor senkrechter Vektor oder der Nullvektor ist. Wir betrachten nun den ersteren Fall weiter.

Wir können dann in diesem Punkt durch

$$\vec{N} = \left| \frac{d\vec{T}}{dt} \right|^{-1} \frac{d\vec{T}}{dt} \quad (3.1.6)$$

einen zweiten zu $\vec{T}$ senkrechten Einheitsvektor definieren. Das Paar $\vec{T}$ und $\vec{N}$ bezeichnen wir als das **die Kurve begleitende Zweibein**.



Bislang haben wir mit einem ganz allgemeinen Parameter t zur Beschreibung der Kurve gearbeitet. Dieser Parameter besitzt im allgemeinen keine besondere geometrische Bedeutung. Eine natürlichere Parametrisierung ist durch die **Bogenlänge** der Kurve, gemessen vom Anfangspunkt $\vec{x}(t_0)$ gegeben. Der Zuwachs der Bogenlänge, der sich ergibt, wenn wir um ein infinitesimales Stückchen weiterrücken, ist offensichtlich durch

$$ds = \left| \frac{d\vec{x}}{dt} \right| dt \quad (3.1.7)$$

gegeben. Denken wir uns also die Kurve durch die Bogenlänge s parametrisiert, vereinfacht sich (3.1.4) zu

$$\vec{T} = \frac{d\vec{x}}{ds}. \quad (3.1.8)$$

Eine Gerade ist nun offenbar durch

$$g: \vec{x}(s) = \vec{x}_0 + \vec{T}s \quad \text{mit} \quad \vec{T} = \text{const.} \quad \text{und} \quad |\vec{T}| = 1, \quad (3.1.9)$$

gegeben. Der Einheitstangentenvektor entlang einer Geraden ändert sich natürlich nicht. Folglich charakterisiert die Größe

$$\kappa = \left| \frac{d\vec{T}}{ds} \right| \quad (3.1.10)$$

Abbildung 3.1: Zur Konstruktion des Krümmungskreises an eine vorgegebene Kurve (Abbildung aus der Wikipedia).

¹repräsentiert durch die Komponenten bzgl. der kartesischen Basis

3.1. Kurven in der Ebene

die Abweichung der Kurve von einer Geraden, also die **Krümmung** der Kurve durch rein geometrische von der Wahl der Parametrisierung der Kurve unabhängige Größen. Die Krümmung κ lässt sich nun noch geometrisch interpretieren. Wie wir oben gesehen haben, steht die infinitesimale Änderung des Tangentenvektors $d\vec{T} = ds d\vec{T}/ds$ senkrecht auf $\vec{T}$. Für den Winkel $d\varphi$ zwischen $\vec{T} + d\vec{T}$ und $\vec{T}$ gilt also (s. Abb. 3.1)

$$\sin(d\varphi) \simeq d\varphi = ds \left| \frac{d\vec{T}}{ds} \right| = ds \kappa. \quad (3.1.11)$$

Die Bogenlänge ist also $ds = d\varphi/\kappa$, d.h. $\rho = 1/\kappa$ ist der Radius des sich an die Kurve im betrachteten Punkt anschmiegenden Kreises, den man in diesem Zusammenhang als **Krümmungskreis** bezeichnet. Die Größe ρ heißt **Krümmungsradius**. Der Ortsvektor des Mittelpunktes des Krümmungskreises an die Kurve in dem betrachteten Punkt liegt demnach bei

$$\vec{K} = \vec{x} + \rho \vec{N}. \quad (3.1.12)$$

Die Menge der Krümmungskreismittelpunkte bildet ihrerseits wieder eine Kurve, die sogenannte **Evolute** der Ausgangskurve.

Wir geben noch die expliziten Ausdrücke für diese Größen für eine beliebige Parametrisierung der Kurve, ausgedrückt durch die Komponenten bzgl. eines kartesischen Koordinatensystems an. Zunächst ist

$$\frac{ds}{dt} = \left| \frac{d\vec{x}}{dt} \right| = |\dot{\vec{x}}| = \sqrt{\dot{x}_1^2 + \dot{x}_2^2}. \quad (3.1.13)$$

Weiter ist wegen (3.1.4) in dieser Schreibweise

$$\vec{T} = \frac{d\vec{x}}{ds} = \frac{\dot{\vec{x}}}{|\dot{\vec{x}}|} = \frac{1}{\sqrt{\dot{x}_1^2 + \dot{x}_2^2}} \begin{pmatrix} \dot{x}_1 \\ \dot{x}_2 \end{pmatrix}. \quad (3.1.14)$$

Wir benötigen noch

$$\frac{d}{dt} |\dot{\vec{x}}| = \frac{d}{dt} \sqrt{\dot{\vec{x}}^2} = \frac{\ddot{\vec{x}} \cdot \dot{\vec{x}}}{|\dot{\vec{x}}|} = \frac{\dot{x}_1 \ddot{x}_1 + \dot{x}_2 \ddot{x}_2}{\sqrt{\dot{x}_1^2 + \dot{x}_2^2}}, \quad (3.1.15)$$

und

$$\frac{d\vec{T}}{ds} = \frac{\dot{\vec{T}}}{|\dot{\vec{T}}|} = \frac{\ddot{\vec{x}} |\dot{\vec{x}}|^2 - \dot{\vec{x}} (\ddot{\vec{x}} \cdot \dot{\vec{x}})}{|\dot{\vec{x}}|^4} = \frac{\dot{x}_1 \ddot{x}_2 - \dot{x}_2 \ddot{x}_1}{(\dot{x}_1^2 + \dot{x}_2^2)^2} \begin{pmatrix} -\dot{x}_2 \\ \dot{x}_1 \end{pmatrix}, \quad (3.1.16)$$

woraus nach einiger Rechnung aus (3.1.10)

$$\kappa = \frac{|\dot{x}_1 \ddot{x}_2 - \dot{x}_2 \ddot{x}_1|}{\sqrt{\dot{x}_1^2 + \dot{x}_2^2}^3} \quad (3.1.17)$$

folgt. Normierung von (3.1.16) liefert gemäß (3.1.6) schließlich für den Normalenvektor

$$\vec{N} = \frac{\text{sign}(\dot{x}_1 \ddot{x}_2 - \dot{x}_2 \ddot{x}_1)}{\sqrt{\dot{x}_1^2 + \dot{x}_2^2}} \begin{pmatrix} -\dot{x}_2 \\ \dot{x}_1 \end{pmatrix}. \quad (3.1.18)$$

Für die Evolute erhalten wir wegen (3.1.12) und (3.1.18)

$$\vec{K} = \vec{x} + \frac{\dot{x}_1^2 + \dot{x}_2^2}{\dot{x}_1 \ddot{x}_2 - \dot{x}_2 \ddot{x}_1} \begin{pmatrix} -\dot{x}_2 \\ \dot{x}_1 \end{pmatrix}. \quad (3.1.19)$$

Beispiel: Für einen Kreis mit Radius r um den Ursprung des Koordinatensystems können wir die Parametrisierung

$$\vec{x} = r \begin{pmatrix} \cos t \\ \sin t \end{pmatrix}, \quad t \in [0, 2\pi) \quad (3.1.20)$$

wählen. Hier ist eine Parametrisierung mit der Bogenlänge einfach, denn es ist offenbar

$$s(t) = \int_0^t dt' |\dot{\vec{x}}(t')| = r t. \quad (3.1.21)$$

Es ist also

$$\vec{x}(s) = r \begin{pmatrix} \cos(s/r) \\ \sin(s/r) \end{pmatrix}, \quad s \in [0, 2\pi r). \quad (3.1.22)$$

Der Tangenteneinheitsvektor an jedem Punkt ist also durch

$$\vec{T} = \frac{d\vec{x}}{ds} = \begin{pmatrix} -\sin(s/r) \\ \cos(s/r) \end{pmatrix} \quad (3.1.23)$$

gegeben und

$$\frac{d\vec{T}}{ds} = -\frac{1}{r} \begin{pmatrix} \cos(s/r) \\ \sin(s/r) \end{pmatrix}, \quad \vec{N} = -\begin{pmatrix} \cos(s/r) \\ \sin(s/r) \end{pmatrix} \quad (3.1.24)$$

Die Krümmung ist wegen (3.1.10)

$$\kappa = \frac{1}{\rho} = \left| \frac{d\vec{T}}{ds} \right| = \frac{1}{r}. \quad (3.1.25)$$

Das ist verständlich, denn der Schmiegkreis an den Kreis ist dieser Kreis selbst. Entsprechend ist die Evolute eines Kreises einfach ein einzelner Punkt, eben der Mittelpunkt des Kreises. In der Tat folgt aus (3.1.12) und (3.1.25) sofort $\vec{K} = 0 = \text{const.}$

3.2 Kegelschnitte

Kegelschnitte sind Kurven, die sich als Schnittmengen eines geraden Doppelkreiskegels mit einer Ebene ergeben.

3.2.1 Definition der Kegelschnitte

Für unsere Zwecke ist nicht die Definition der Kegelschnitte als die Schnittkurven einer Ebene mit einem Kreiskegel bequem, sondern zunächst einmal die **Ortsdefinitionen**. Wir unterscheiden drei Grundtypen:

- **Ellipse:** Für eine Ellipse sind zwei Punkte in einer Ebene vorgegeben, die **Brennpunkte** oder **Foci** F_1 und F_2 . Die Ellipse ist dann diejenige Kurve in dieser Ebene, für die jeder Punkt P auf der Ellipse die Gleichung

$$|F_1P| + |F_2P| = 2a = \text{const} \quad (3.2.1)$$

erfüllt, d.h. für jeden Punkt P der Ellipse ist die **Summe der Abstände** von den beiden Brennpunkten gleich $2a$.

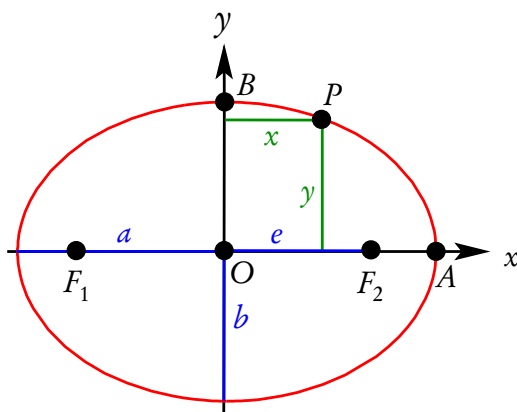
Daraus ergibt sich auch die praktische Möglichkeit eine Ellipse zu zeichnen, die sog. Gärtnerkonstruktion. Man befestigt die Enden einer Schnur der Länge $2a$ in den beiden Brennpunkten und zeichnet dann die Kurve, die sich ergibt, wenn man den Stift entlang des straff gespannten Seils führt. Der Name „Gärtnerkonstruktion“ rührt daher, dass man dieses Verfahren zum Begrenzen schöner ellipsenförmiger Beete im Garten anwenden kann.

- **Hyperbel:** Auch für eine Hyperbel sind zwei Brennpunkte F_1 und F_2 vorgegeben. Hier ist die **Differenz der Abstände der Punkte** auf der Hyperbel betragsmäßig konstant, d.h.

$$||F_1P| - |F_2P|| = 2a. \quad (3.2.2)$$

- **Parabel:** Für eine Parabel ist ein Brennpunkt F vorgegeben und eine Gerade, die sog. **Leitgerade** in einer Ebene. Die Parabel ist dann durch diejenigen Punkte in der Ebene definiert, für die die Summe der Abstände von der Geraden und dem Brennpunkt konstant ist.

3.2.2 Ellipse



Beginnen wir mit der Ellipse und betrachten die nebenstehende Skizze. Zuerst leiten wir einige geometrische Eigenschaften her. Betrachten wir dazu den Punkt A auf der großen Halbachse der Ellipse. Aus der Definition der Ellipse folgt dann

$$|F_1A| + |F_2A| = (e + a) + (a - e) = 2a. \quad (3.2.3)$$

Damit ist die entlang der Ellipse konstante Summe $2a$. Wenden wir den Satz des Pythagoras auf das rechtwinklige Dreieck OF_2B an, erhalten wir

$$e^2 + b^2 = a^2, \quad (3.2.4)$$

denn die Hypotenuse dieses Dreiecks ist aus Symmetriegründen $(|F_1B| + |F_2B|)/2 = 2a/2 = a$. Nun können wir eine Gleichung für die Koordinaten (x, y) des Punktes P herleiten. Es ist nämlich, wieder nach dem Satz des Pythagoras

$$|F_1P| = \sqrt{(e + x)^2 + y^2}, \quad |F_2P| = \sqrt{(e - x)^2 + y^2}. \quad (3.2.5)$$

Damit folgt

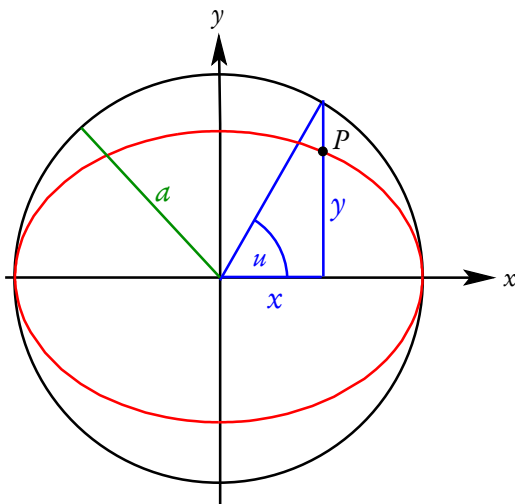
$$\sqrt{(e + x)^2 + y^2} + \sqrt{(e - x)^2 + y^2} = 2a. \quad (3.2.6)$$

Man kann nach einer einfachen aber etwas umfangreichen Rechnung (*Übungsaufgabe!*) zeigen, dass daraus

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1 \quad (3.2.7)$$

folgt.

3. Vektoranalysis



Eine andere nützliche Darstellung ergibt sich mit Polarkoordinaten mit dem Mittelpunkt der Ellipse als Ursprung. Aus der nebenstehenden Skizze lesen wir zunächst ab

$$x = a \cos u \quad (3.2.8)$$

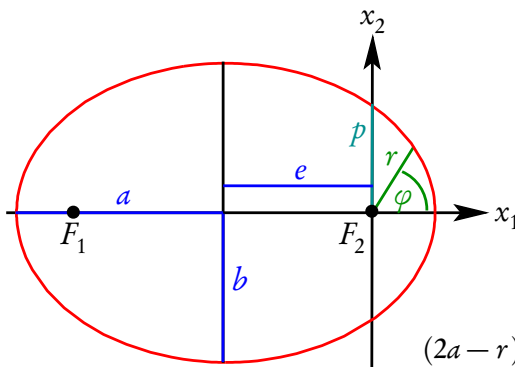
Setzen wir dies in (3.2.7) ein, erhalten wir

$$\cos^2 u + \frac{y^2}{b^2} = 1 \Rightarrow \frac{y^2}{b^2} = 1 - \cos^2 u = \sin^2 u. \quad (3.2.9)$$

Man macht sich schnell klar, dass man mit

$$x = a \cos u, \quad y = b \sin u, \quad u \in [0, 2\pi) \quad (3.2.10)$$

eine vollständige Parametrisierung der Ellipse erhält.



Schließlich benötigen wir für das Verständnis des Kepler-Problems in der Mechanik noch die Ellipsengleichung in Polarkoordinaten mit dem Ursprung in einem Brennpunkt (s. nebenstehende Skizze). Dies lesen wir wieder direkt aus der geometrischen Definition (3.2.3) ab:

$$r + \sqrt{(2e + r \cos \varphi)^2 + r^2 \sin^2 \varphi} = 2a. \quad (3.2.11)$$

Dies formen wir zunächst um in

$$\begin{aligned} (2a - r)^2 &= (2e + r \cos \varphi)^2 + r^2 \sin^2 \varphi \\ &= 4e^2 + 4er \cos \varphi + r^2. \end{aligned} \quad (3.2.12)$$

Mit (3.2.4) und mit den durch

$$p = \frac{b^2}{a}, \quad \epsilon = \frac{e}{a} \quad (3.2.13)$$

definierten Parametern p (**Halbparameter**) und ϵ (**numerische Exzentrizität**) finden wir schließlich

$$r = \frac{p}{1 + \epsilon \cos \varphi}. \quad (3.2.14)$$

Dies erklärt auch die geometrische Bedeutung des in der nebenstehenden Skizze eingezeichneten Halbparameters, denn offenbar ist $p = r(\varphi = \pi/2)$.

Nun berechnen wir noch die Fläche der Ellipse. Dazu berechnen wir die Fläche unter der Kurve

$$y = b \sqrt{1 - \frac{x^2}{a^2}}, \quad (3.2.15)$$

was dem oberen Teil der durch (3.2.7) beschriebenen Ellipse, also der Lösung dieser Gleichung mit $y \geq 0$ entspricht. Es bietet sich auch gleich die Substitution $x = a \cos u$ mit $u \in [0, \pi]$ an. Dann erhalten wir für die Fläche der Ellipse

$$A = 2b \int_{-a}^a dx \sqrt{1 - \frac{x^2}{a^2}} = 2b \int_0^\pi du a \sin u \sqrt{1 - \cos^2 u} = 2ab \int_0^\pi du \sin^2 u. \quad (3.2.16)$$

3.2. Kegelschnitte

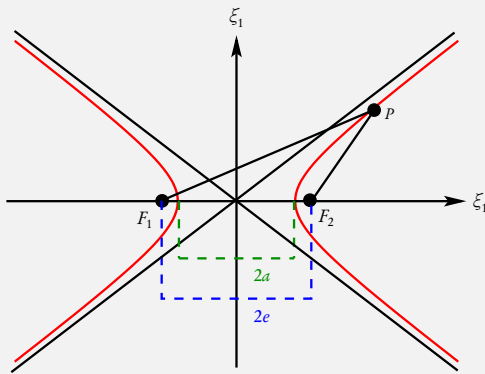
Um das Integral zu berechnen, verwenden wir die Doppelwinkelformel für den Kosinus

$$\cos(2u) = \cos^2 u - \sin^2 u = 1 - 2\sin^2 u \Rightarrow \sin^2 u = \frac{1}{2}[1 - \cos(2u)]. \quad (3.2.17)$$

Damit folgt aus (3.2.16)

$$A = ab \int_0^\pi [1 - \cos(2u)] = ab \left[u - \frac{1}{2} \sin(2u) \right]_{u=0}^{u=\pi} = \pi ab. \quad (3.2.18)$$

3.2.3 Hyperbel



Betrachten wir Punkt A der nebenstehend abgebildeten Hyperbel, finden wir aus der oben angegebenen Ortsdefinition, dass in der Tat

$$||F_1 A| - |F_2 A|| = 2a \quad (3.2.19)$$

ist. Mit dem Satz des Pythagoras lesen wir für die Ortskoordinaten

$$\sqrt{(x+e)^2 + y^2} - \sqrt{(x-e)^2 + y^2} = \pm 2a \quad (3.2.20)$$

Nach einigen einfachen aber umfangreichen Umformungen, erhält man

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1, \quad b = \sqrt{e^2 - a^2}. \quad (3.2.21)$$

Dies beschreibt beide Äste der Hyperbel.

Eine nützliche Parameterdarstellung erhält man mit Hilfe der **Hyperbelfunktionen** $\cosh$ und $\sinh$ (Kosinus bzw. Sinus hyperbolicus), deren Name von dieser geometrischen Bedeutung herrührt, und die durch

$$\cosh u = \frac{1}{2}[\exp(u) + \exp(-u)], \quad \sinh u = \frac{1}{2}[\exp(u) - \exp(-u)]. \quad (3.2.22)$$

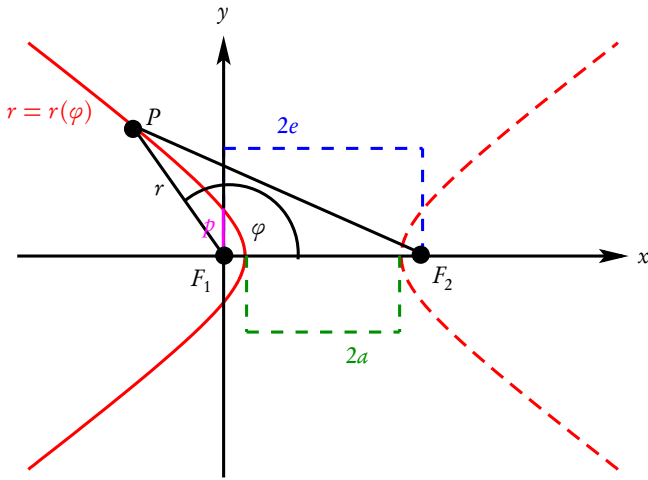
Wir erhalten dann den rechten (oberes Vorzeichen) bzw. linken (unteres Vorzeichen) Ast der Hyperbel durch

$$x = \pm a \cosh u, \quad y = b \sinh u, \quad u \in \mathbb{R}, \quad (3.2.23)$$

denn man beweist leicht durch direktes Nachrechnen mit (3.2.22) (*Übung*), dass stets

$$\cosh^2 u - \sinh^2 u = 1 \quad (3.2.24)$$

gilt.



Für die Bahnbewegung beim Kepler-Problem benötigen wir wieder die Beschreibung des einen Hyperbelastes in Polarkoordinaten mit dem einen Brennpunkt im Koordinatenursprung. Da die Astronomen den sonnennächsten Punkt als Bezugspunkt verwenden, d.h. das Perihel bzw. Periastron entspricht $\varphi = 0$, handelt es sich um den linken Ast. Mit Hilfe des Kosinus-Satzes für das Dreieck F_1F_2P lesen wir ab

$$|PF_2|^2 = (2e)^2 + r^2 - 4er \cos \varphi. \quad (3.2.25)$$

Mit der Definition der Hyperbel erhalten wir damit

$$|PF_2| - |PF_1| = \sqrt{(2e)^2 + r^2 - 4er \cos \varphi} - r = 2a. \quad (3.2.26)$$

Dies wollen wir nach r auflösen. Dazu rechnen wir

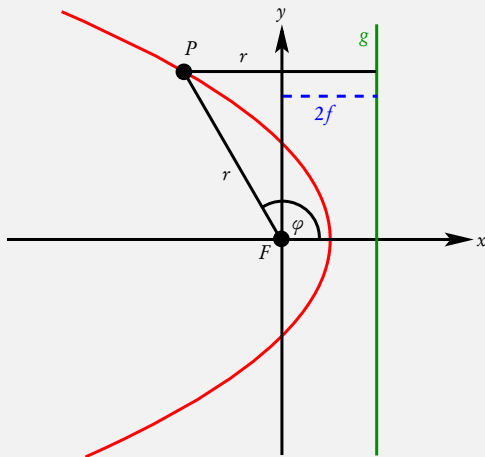
$$(2a + r)^2 = (2e)^2 + r^2 - 4er \cos \varphi \quad (3.2.27)$$

Mit der binomischen Formel auf der linken Seite folgt nach einigen einfachen Umformungen

$$(2a)^2 + 4ar + r^2 = (2e)^2 + r^2 - 4er \cos \varphi \Rightarrow r(\varphi) = \frac{p}{1 + \epsilon \cos \varphi}. \quad (3.2.28)$$

Dabei ist wieder $p = b^2/a$ und $\epsilon = e/a$, d.h. die Formeln sind formal identisch mit denen bei der Ellipse. Der entscheidende Unterschied ist, dass bei der Hyperbel $e > a$ ist, d.h. $\epsilon > 1$. Da $r \geq 0$ ist, ist der geometrisch sinnvolle Winkelbereich eingeschränkt. Wählen wir als Bereich $\varphi \in (-\pi, \pi)$, folgt aus $1 + \epsilon \cos \varphi \geq 0$, dass $\cos \varphi \geq -1/\epsilon$ sein muss, d.h. $|\varphi| \leq \arccos(-1/\epsilon)$, d.h. der geometrisch sinnvolle Bereich ist $\varphi \in (-\arccos(-1/\epsilon), \arccos(-1/\epsilon))$. Bei der Annäherung von φ an die entsprechenden Werte $\pm \arccos(-1/\epsilon)$ wird $r \rightarrow \infty$, d.h. die Hyperbelbahn ist unendlich ausgedehnt. Beim Kepler-Problem bedeutet dies, dass sich der Himmelskörper aus dem Unendlichen an die Sonne annähert und auch wieder im Unendlichen verschwindet, d.h. es liegt keine gebundene Bahn wie bei der Ellipse vor.

3.2.4 Parabel



Bei der Parabel beschränken wir uns auf die Herleitung der für das Kepler-Problem benötigten Gleichung in Polarkoordinaten mit dem Brennpunkt im Koordinatenursprung. Aus der Definition der Parabel lesen wir ab

$$r = 2f - r \cos \varphi \Rightarrow r(\varphi) = \frac{2f}{1 + \cos \varphi} = \frac{p}{1 + \cos \varphi}, \quad (3.2.29)$$

d.h. auch hier haben wir mit dem Parameter $p = 2f$ dieselbe Form für die Gleichung wie bei Ellipse und Hyperbel, wobei aber $\epsilon = 1$ ist. Der geometrische Winkelbereich ist offenbar $\varphi \in (-\pi, \pi)$, wobei für $\varphi \rightarrow \pm\pi$ offenbar $r \rightarrow \infty$ ist. Auch die Parabelbahn entspricht also beim Kepler-Problem einer ungebundenen Bewegung.

3.3 Raumkurven und Fresnetsche Formeln

Die Verallgemeinerung der Betrachtungen zu Kurven im Euklidischen $\mathbb{R}^3$ ist nicht besonders schwierig. Auch hier definieren wir eine Kurve durch eine Parameterdarstellung $\vec{x} : \mathbb{R} \supseteq (t_0, t_1) \rightarrow \mathbb{R}^3$ und ebenso wie für die ebenen Kurven ist ein intrinsischer Parameter durch die Bogenlänge der Kurve, gezählt vom Anfangspunkt $\vec{x}(t_0)$ eine geometrische der Kurve immanenter (d.h. von der Wahl der Parametrisierung unabhängige) Größe. Betrachten wir also zunächst die Kurve wieder in dieser Parametrisierung.

Dann ist der **Einheitstangentenvektor** durch

$$\vec{T} = \frac{d\vec{x}}{ds} = \frac{\dot{\vec{x}}}{|\dot{\vec{x}}|} \quad (3.3.1)$$

definiert, und der **Hauptnormaleneinheitsvektor** ist durch

$$\vec{N} = \frac{1}{\kappa} \frac{d\vec{T}}{ds}, \quad \kappa = \frac{1}{\rho} = \left| \frac{d\vec{T}}{ds} \right| \quad (3.3.2)$$

gegeben. Dabei ist ρ wieder der **Krümmungsradius** und $\kappa = 1/\rho$ die **Krümmung der Kurve** an dem betrachteten Punkt.

Wir ergänzen nun $\vec{T}$ und $\vec{N}$ durch die Definition des **Binormaleneinheitsvektors** vermöge

$$\vec{B} = \vec{T} \times \vec{N}. \quad (3.3.3)$$

Damit ist durch $(\vec{T}, \vec{N}, \vec{B})$ eine allein durch die Geometrie der Kurve bestimmte (also von der Wahl der Parametrisierung unabhängige) rechtshändige Orthonormalbasis, das die Kurve **begleitende Dreibein**, definiert.

Gegenüber den ebenen Kurven ist im Raum die Frage sinnvoll, ob die Kurve im Raum in einer Ebene bleibt oder nicht. Der Tangentenvektor $\vec{T}$ und die Hauptnormale $\vec{N}$ spannen die sogenannte **Schmiegeebene** auf. Ändert sich diese Ebene entlang der Kurve nicht, handelt es sich definitionsgemäß um eine **ebene Kurve im Raum**. Da der Binormalenvektor stets senkrecht auf der Schmiegeebene steht, ist dies genau dann der Fall, wenn $d\vec{B}/ds = 0$ ist.

Die geometrische Bedeutung der Schmiegeebene wird klar, wenn man überlegt, wie sich die Kurve in der Umgebung eines Punktes $\vec{x}_0 = \vec{x}(s_0)$ verhält. Dazu entwickeln wir die Parametrisierung der Kurve bis zu zweiter Ordnung um s_0 :

$$\vec{x}(s_0 + \Delta s) = \vec{x}(s_0) + \Delta s \frac{d}{ds} \vec{x}(s_0) + \frac{\Delta s^2}{2} \frac{d^2}{ds^2} \vec{x}(s_0) + \mathcal{O}(\Delta s^3). \quad (3.3.4)$$

Nun ist gemäß (3.3.1) und (3.3.2)

$$\frac{d}{ds} \vec{x}(s_0) = \vec{T}(s_0), \quad \frac{d^2}{ds^2} \vec{x}(s_0) = \frac{d}{ds} \vec{T}(s_0) = \frac{1}{\rho} \vec{N}(t_0). \quad (3.3.5)$$

Dies in (3.3.4) eingesetzt liefert

$$\vec{x}(s_0 + \Delta s) = \vec{x}(s_0) + \Delta s \vec{T}(s_0) + \frac{\Delta s^2}{2} \frac{1}{\rho} \vec{N}(t_0) + \mathcal{O}(\Delta s^3). \quad (3.3.6)$$

3. Vektoranalysis

In dieser Ordnung der Näherung bleibt also die Kurve in der durch $\vec{T}(s_0)$ und $\vec{N}(s_0)$ aufgespannten Ebene durch $\vec{x}(s_0)$, was die Bezeichnung Schmiegeebene rechtfertigt. Der Binormalenvektor ist dann der Normaleneinheitsvektor auf diese Ebene, wobei die Orientierung definitionsgemäß wie durch (3.3.3) angegeben gewählt wurde. Die Kurve bleibt demnach in der Tat in einer Ebene, wenn sich $\vec{B}$ entlang der Kurve nicht ändert.

Wir suchen nun ein Maß für die Abweichung der Kurve von einer ebenen Kurve, also dafür, um wieviel sich die Kurve beim Fortschreiten um eine Bogenlängenänderung ds aus der Schmiegeebene herauswindet. Dafür bietet sich offensichtlich der Vektor $d\vec{B}/ds$ an. Da das begleitende Dreiein eine Orthonormalbasis bildet, muss sich allerdings dieser Vektor als Linearkombination dieser Basisvektoren ausdrücken lassen. Wir wollen nun zeigen, dass er parallel zum Hauptnormalenvektor $\vec{N}$ ist. Dazu differenzieren wir die Gleichung $\vec{T} \cdot \vec{B} = 0 = \text{const}$ nach der Bogenlänge:

$$\frac{d\vec{T}}{ds} \cdot \vec{B} + \vec{T} \cdot \frac{d\vec{B}}{ds} \stackrel{(3.3.2)}{=} \chi \vec{N} \cdot \vec{B} + \vec{T} \cdot \frac{d\vec{B}}{ds} \stackrel{(3.3.3)}{=} \vec{T} \cdot \frac{d\vec{B}}{ds} = 0. \quad (3.3.7)$$

Dies bedeutet, dass $d\vec{B}/ds$ senkrecht auf $\vec{T}$ steht. Ebenso folgt aus $\vec{B}^2 = 1 = \text{const}$, dass dieser Vektor auch senkrecht auf $\vec{B}$ steht.

Damit muss aber

$$\frac{d\vec{B}}{ds} = -\chi \vec{N}, \quad \chi = -\vec{N} \cdot \frac{d\vec{B}}{ds} \quad (3.3.8)$$

sein. Die Größe χ heißt **Windung oder Torsion**.

Sie ist eine vorzeichenbehaftete Größe. Die geometrische Bedeutung wird klar, wenn wir (3.3.3) beachten. Aus der Formel für das doppelte Kreuzprodukt (2.1.53) folgt

$$-\vec{N} = \vec{T} \times \vec{B}, \quad (3.3.9)$$

so dass wir (3.3.8) auch in der Form

$$\frac{d\vec{B}}{ds} = \chi \vec{T} \times \vec{B} \quad (3.3.10)$$

schreiben können. Ist also χ positiv, so ist die Änderung von $\vec{B}$ in diesem Punkt bei infinitesimalem Fortschreiten entlang der Kurve um die Länge ds durch eine infinitesimale Drehung um den Winkel χds um die durch $\vec{T}$ definierte Drehachse im Sinne der Rechten-Hand-Regel gegeben.

Die Größe χ gibt also die lokale Ganghöhe der **Schraube** an, und zwar im Sinne einer **Rechtsschraube** falls $\chi > 0$ und im Sinne einer **Linksschraube** falls $\chi < 0$ ist.

Bevor wir ein charakteristisches Beispiel für diese Begriffsbildungen durchrechnen, wollen wir noch die **Frenet-Serret-Formeln** herleiten. Diese geben die Komponenten der Ableitungen des begleitenden Dreieins nach der Bogenlänge der Kurve bzgl. der durch das begleitende Dreiein gegebenen Basis an.

Aus (3.3.2) und (3.3.8) folgen bereits zwei der drei gesuchten Formeln

$$\frac{d\vec{T}}{ds} = \kappa \vec{N} \Rightarrow \vec{T} \cdot \frac{d\vec{T}}{ds} = 0, \quad \vec{N} \cdot \frac{d\vec{T}}{ds} = \kappa, \quad \vec{B} \cdot \frac{d\vec{T}}{ds} = 0, \quad (3.3.11)$$

$$\frac{d\vec{B}}{ds} = -\chi \vec{N} \Rightarrow \vec{T} \cdot \frac{d\vec{B}}{ds} = 0, \quad \vec{N} \cdot \frac{d\vec{B}}{ds} = -\chi, \quad \vec{B} \cdot \frac{d\vec{B}}{ds} = 0. \quad (3.3.12)$$

Daraus ergibt sich mit (3.3.9) aber auch sofort für die verbleibende Ableitung

$$\frac{d\vec{N}}{ds} = \frac{d}{ds}(\vec{B} \times \vec{T}) = \frac{d\vec{B}}{ds} \times \vec{T} + \vec{B} \times \frac{d\vec{T}}{ds} = -\chi \vec{N} \times \vec{T} + \kappa \vec{B} \times \vec{N} = \chi \vec{B} - \kappa \vec{T}, \quad (3.3.13)$$

wobei wir benutzt haben, dass $(\vec{T}, \vec{N}, \vec{B})$ eine rechtshändige Orthonormalbasis bilden.

Aus (3.3.13) folgt also

$$\vec{T} \cdot \frac{d\vec{N}}{ds} = -\kappa, \quad \vec{N} \cdot \frac{d\vec{N}}{ds} = 0, \quad \vec{B} \cdot \frac{d\vec{N}}{ds} = \chi. \quad (3.3.14)$$

Oft ist eine Parametrisierung der Kurve durch die Bogenlänge unbequem, so dass wir das begleitende Dreibein und die Krümmung und Torsion noch für eine beliebige Parametrisierung umschreiben wollen. Dazu benötigen wir nur

$$\frac{ds}{dt} = \left| \dot{\vec{x}} \right|, \quad (3.3.15)$$

wobei der Punkt über einem Symbol wieder die Ableitung nach dem beliebigen Parameter t bedeutet. Daraus folgt

$$\vec{T} = \frac{d\vec{x}}{ds} = \frac{\dot{\vec{x}}}{\left| \dot{\vec{x}} \right|}. \quad (3.3.16)$$

Weiter ist wegen $\left| \dot{\vec{x}} \right| = \sqrt{\dot{\vec{x}}^2}$

$$\frac{d}{dt} \left| \dot{\vec{x}} \right| = \frac{\ddot{\vec{x}} \cdot \dot{\vec{x}}}{\left| \dot{\vec{x}} \right|}. \quad (3.3.17)$$

Daraus folgt

$$\dot{\vec{T}} = \frac{\ddot{\vec{x}} \dot{\vec{x}}^2 - \dot{\vec{x}} (\ddot{\vec{x}} \cdot \dot{\vec{x}})}{\left| \dot{\vec{x}} \right|^3} \stackrel{(2.1.53)}{=} \frac{\dot{\vec{x}} \times (\ddot{\vec{x}} \times \dot{\vec{x}})}{\left| \dot{\vec{x}} \right|^3}. \quad (3.3.18)$$

Für die Normierung dieses Vektors verwenden wir

$$\left| \dot{\vec{T}} \right| = \frac{\left| \dot{\vec{x}} \right| \left| \ddot{\vec{x}} \times \dot{\vec{x}} \right|}{\left| \dot{\vec{x}} \right|^3} = \frac{\left| \ddot{\vec{x}} \times \dot{\vec{x}} \right|}{\left| \dot{\vec{x}} \right|^2}. \quad (3.3.19)$$

Der Hauptnormalenvektor ist also durch

$$\vec{N} = \frac{\dot{\vec{T}}}{\left| \dot{\vec{T}} \right|} = \frac{\dot{\vec{x}} \times (\ddot{\vec{x}} \times \dot{\vec{x}})}{\left| \dot{\vec{x}} \right| \left| \ddot{\vec{x}} \times \dot{\vec{x}} \right|} \quad (3.3.20)$$

3. Vektoranalysis

gegeben und der Binormalenvektor durch

$$\vec{B} = \vec{T} \times \vec{N} = \frac{\dot{\vec{x}} \times \ddot{\vec{x}}}{|\dot{\vec{x}} \times \ddot{\vec{x}}|}. \quad (3.3.21)$$

Aus (3.3.11) ergibt sich für die Krümmung

$$\kappa = \vec{N} \cdot \frac{d\vec{T}}{ds} = \frac{\vec{N} \cdot \dot{\vec{T}}}{|\dot{\vec{x}}|} \stackrel{(3.3.18, 3.3.20)}{=} \frac{|\dot{\vec{x}} \times (\ddot{\vec{x}} \times \dot{\vec{x}})|^2}{|\dot{\vec{x}}|^5 |\dot{\vec{x}} \times \ddot{\vec{x}}|} = \frac{|\ddot{\vec{x}} \times \dot{\vec{x}}|}{|\dot{\vec{x}}|^3}. \quad (3.3.22)$$

Zur Berechnung der Torsion gehen wir von (3.3.12) aus:

$$\chi = -\vec{N} \cdot \frac{d\vec{B}}{ds} = -\frac{\vec{N} \cdot \dot{\vec{B}}}{|\dot{\vec{x}}|}. \quad (3.3.23)$$

Die Ableitung von $\vec{B}$ ist

$$\dot{\vec{B}} = \frac{\dot{\vec{x}} \times \ddot{\vec{x}}}{|\dot{\vec{x}} \times \ddot{\vec{x}}|} - \frac{\vec{B}}{|\dot{\vec{x}} \times \ddot{\vec{x}}|} \frac{d}{dt} |\dot{\vec{x}} \times \ddot{\vec{x}}|. \quad (3.3.24)$$

Multiplikation mit $\vec{N}$ liefert schließlich

$$\chi = \frac{\ddot{\vec{x}} \cdot (\dot{\vec{x}} \times \ddot{\vec{x}})}{|\dot{\vec{x}} \times \ddot{\vec{x}}|}. \quad (3.3.25)$$

Als **Beispiel** für diese geometrischen Betrachtungen, das besonders anschaulich ist, betrachten wir die Schraubenlinie

$$\vec{x}(t) = \begin{pmatrix} r \cos t \\ r \sin t \\ at \end{pmatrix}, \quad t \in [0, \infty), \quad r > 0, \quad a \in \mathbb{R}. \quad (3.3.26)$$

Wir haben

$$|\dot{\vec{x}}| = \sqrt{r^2 + a^2} \Rightarrow s = t \sqrt{r^2 + a^2}. \quad (3.3.27)$$

Wir können also in diesem Fall sehr einfach zur Parametrisierung mit der Bogenlänge übergehen, was die weiteren Rechnungen etwas erleichtert:

$$\vec{x}(s) = \begin{pmatrix} r \cos\left(\frac{s}{\sqrt{r^2 + a^2}}\right) \\ r \sin\left(\frac{s}{\sqrt{r^2 + a^2}}\right) \\ \frac{as}{\sqrt{r^2 + a^2}} \end{pmatrix}. \quad (3.3.28)$$

Daraus ergibt sich mit (3.3.1-3.3.3) für das begleitende Dreibein

$$\vec{T} = \frac{1}{\sqrt{r^2 + a^2}} \begin{pmatrix} -r \sin\left(\frac{s}{\sqrt{r^2 + a^2}}\right) \\ r \cos\left(\frac{s}{\sqrt{r^2 + a^2}}\right) \\ a \end{pmatrix}, \quad (3.3.29)$$

$$\vec{N} = \begin{pmatrix} -\cos\left(\frac{s}{\sqrt{r^2 + a^2}}\right) \\ -\sin\left(\frac{s}{\sqrt{r^2 + a^2}}\right) \\ 0 \end{pmatrix}, \quad (3.3.30)$$

$$\vec{B} = \frac{1}{\sqrt{r^2 + a^2}} \begin{pmatrix} a \sin\left(\frac{s}{\sqrt{r^2 + a^2}}\right) \\ -a \cos\left(\frac{s}{\sqrt{r^2 + a^2}}\right) \\ r \end{pmatrix}, \quad (3.3.31)$$

Krümmung und Torsion der Kurve berechnen wir am bequemsten mit Hilfe der Gleichungen (3.3.2) bzw. (3.3.8):

$$\kappa = \frac{r}{r^2 + a^2}, \quad \chi = \frac{a}{r^2 + a^2}. \quad (3.3.32)$$

Die Kurve ist also für $a > 0$ eine Rechts-, $a < 0$ eine Linksschraube und für $a = 0$ eben. In der Tat ergibt sich für $a = 0$ ein Kreis in der 12-Ebene mit Radius r um den Ursprung.

3.3.1 Anwendung auf die Bewegung eines Punktteilchens

Eine grundlegende Fragestellung der **klassischen Mechanik** ist die Beschreibung der Bewegung von **Punktteilchen** unter Einfluss von **Kräften**. Dabei versteht man unter einem Punktteilchen einen makroskopischen Körper *endlicher Ausdehnung*, dessen Ausdehnung allerdings für die betrachtete Bewegung irrelevant ist. Das Paradebeispiel für den Erfolg einer solchen **effektiven Beschreibung** von Körpern als Punktteilchen ist die **Himmelsmechanik**, die die eigentliche Motivation für die Entwicklung der klassischen Mechanik durch **Newton** war. Die Ausdehnung der Sonne und der Planeten in unserem Sonnensystem kann in der Tat für die Berechnung der Bewegung dieser Himmelskörper unter Einwirkung ihrer gegenseitigen **Gravitation** in sehr guter Näherung vernachlässigt werden.

Hier betrachten wir die einfachste Fragestellung dieser Art: Ein als **Massenpunkt** idealisierter Körper bewege sich unter Einfluss irgendwelcher vorgegebener Kräfte. Z.B. kann man bei der Planetenbewegung in einer ersten Näherung die Sonne als feststehenden Körper betrachten da ihre Masse M sehr groß gegenüber der Masse m des Planeten ist. Wir wählen ein kartesisches Koordinatensystem mit dem Ursprung im Schwerpunkt der Sonne. Dann ist die Kraft auf den Planeten durch das **Newtonsche Gravitationsgesetz**

$$\vec{F}(\vec{x}) = -\gamma \frac{mM}{r^2} \frac{\vec{x}}{r} \quad (3.3.33)$$

gegeben. Dabei ist $\vec{x}$ der **Ortsvektor** des Planeten, also der Vektor vom Schwerpunkt der Sonne zum Ort des Planeten, $r = |\vec{x}|$ sein Betrag und γ die **Newtonsche Gravitationskonstante**. Die Gleichung (3.3.33) besagt, dass die Kraft stets vom Planeten in Richtung zur Sonne weist und proportional zur Masse des Planeten und der Sonne ist und mit dem **Quadrat des Abstands** abfällt. Aufgabe der **Dynamik** ist es nun, die **Position** des Planeten als Funktion der **Zeit** aus diesem Kraftgesetz zu bestimmen. Dazu verwenden wir die **Newtonsche Bewegungsgleichung**

$$m\ddot{\vec{x}} = \vec{F}(\vec{x}). \quad (3.3.34)$$

Dabei bedeutet $\ddot{\vec{x}}$ die zweite Ableitung des Ortsvektors des Planeten nach der Zeit, also seine **Beschleunigung**. Es ist klar, dass die Lösung dieses Problems nicht einfach ist, denn die unbekannte Funktion $\vec{x}(t)$ kommt

mitsamt der zweiten Ableitung in der Gleichung vor, und wir suchen möglichst alle Funktionen, die diese Gleichung erfüllen, um zu sehen, welche Bahnen es bei dieser Bewegung gibt. Diese Aufgabenstellung ist typisch für die gesamte Physik. Es handelt sich um **Differentialgleichungen**, und ein wesentliches Ziel dieser Vorlesung ist es, Lösungsstrategien für diese Gleichungen zu entwickeln.

Hier betrachten wir die einfachere Aufgabe, die allgemeine **Kinematik** solcher Bewegungen von Punktteilchen zu verstehen. Wir nehmen also an, wir hätten eine Lösung $\vec{x}(t)$ gefunden und untersuchen die geometrische und physikalische Bedeutung einer solchen Lösung. Zunächst einmal ist klar, dass $\vec{x}(t)$ die **Bahnkurve** des Massenpunktes beschreibt, wobei die Zeit als Parameter für diese Bahnkurve dient. Es sind nun alle in diesem Abschnitt entwickelten geometrischen Analysemethoden anwendbar. Betrachten wir als erstes die **Momentangeschwindigkeit** des Teilchens. Sie ist durch die **Zeitableitung**

$$\vec{v}(t) = \dot{\vec{x}}(t) = \frac{d}{dt} \vec{x}(t). \quad (3.3.35)$$

Die geometrische Bedeutung wird dabei klar, wenn wir uns die Kurve als Funktion der Bogenlänge, wie oben beschrieben, parametrisiert denken. Wir führen also zunächst die neue Funktion $\vec{x}(s)$ ein. Dabei ist die Bogenlänge durch das Differential

$$ds = dt |\dot{\vec{x}}(t)| \quad (3.3.36)$$

gegeben, und wir erhalten die Bogenlänge, von einem Anfangszeitpunkt t_0 aus gerechnet durch **Integration**

$$s(t) = \int_{t_0}^t dt' |\dot{\vec{x}}(t')|. \quad (3.3.37)$$

In kartesischen Koordinaten ausgeschrieben lautet das Integral

$$s(t) = \int_{t_0}^t dt' \sqrt{\dot{x}_1^2(t') + \dot{x}_2^2(t') + \dot{x}_3^2(t')}. \quad (3.3.38)$$

Es ist klar, dass (außer für den trivialen Fall, dass das Teilchen über ein endliches Zeitintervall ruht) die Bogenlänge eine monoton wachsende Funktion der Zeit ist, denn der Integrand ist stets positiv.

Nun folgt aus dem Hauptsatz der Differential- und Integralrechnung, dass

$$\dot{s}(t) = \frac{d}{dt} s(t) = |\dot{\vec{x}}(t)| = |\vec{v}(t)| = v(t). \quad (3.3.39)$$

Andererseits ergibt die Kettenregel

$$\vec{v}(t) = \frac{d}{dt} \vec{x}[s(t)] = \dot{s}(t) \frac{d}{ds} \vec{x}[s(t)] = v(t) \vec{T}[s(t)]. \quad (3.3.40)$$

Dabei haben wir die Definition des Einheitstangentenvektors an die Kurve (3.3.1) verwendet. Dies macht den Geschwindigkeitsvektor intuitiv gut verständlich: Sein Betrag gibt die momentane pro Zeiteinheit zurückgelegte Strecke ds/dt an, und seine Richtung ist tangential zur Bahnkurve.

Kommen wir nun zur Beschleunigung. Sie ist die zeitliche Änderung der Geschwindigkeit pro Zeiteinheit, also

$$\vec{a}(t) = \dot{\vec{v}}(t) = \dot{v}(t) \vec{T}[s(t)] + v^2(t) \frac{d}{ds} \vec{T}[s(t)], \quad (3.3.41)$$

wobei wir von der Produkt- und Kettenregel Gebrauch gemacht haben.

Wegen (3.3.2) können wir dies in der Form

$$\vec{a}(t) = \dot{v}(t)\vec{T}[s(t)] + \rho(t)v^2(t)\vec{N}[s(t)], \quad (3.3.42)$$

wobei $\rho(t)$ der Krümmungsradius in dem momentanen Punkt der Bahn ist. Die Beschleunigung setzt sich demnach aus zwei Teilen zusammen: Der erste Term ist die **Tangentialbeschleunigung**, die sich aus der zeitlichen Änderung des Betrages der Geschwindigkeit ergibt und in tangentielle Richtung weist. Dabei kann freilich auch $\dot{v}(t) < 0$ sein. In dem Fall wird das Teilchen „abgebremst“. Auch dies ist physikalisch gesehen eine Beschleunigung. Der zweite Term in (3.3.42) heißt **Zentripetalbeschleunigung**. Er tritt nur auf, wenn sich auch die *Richtung* der Geschwindigkeit ändert, also die Bahn in dem betreffenden Punkt von einer Geraden abweicht. Dabei ist $\vec{N}[s(t)]$ ist der Hauptnormalenvektor an die Bahnkurve in dem betreffenden Punkt und stets $\perp \vec{T}[s(t)]$.

Betrachten wir als einfaches Beispiel die Bewegung auf einer Kreisbahn mit Radius R um den Ursprung in der 12-Ebene des kartesischen Koordinatensystems. Wir können sie wie folgt parametrisieren

$$\vec{r}(t) = R[\cos \varphi(t)\vec{e}_1 + \sin \varphi(t)\vec{e}_2]. \quad (3.3.43)$$

Dabei ist $\varphi = \varphi(t)$ der Winkel des Ortsvektors zur 1-Achse des Koordinatensystems. Wir finden nun nacheinander die Geschwindigkeit und die Beschleunigung durch Ableiten nach der Zeit. Unter Verwendung der Kettenregel ergibt sich

$$\vec{v}(t) = \dot{\vec{r}}(t) = R\dot{\varphi}[-\sin \varphi(t)\vec{e}_1 + \cos \varphi(t)\vec{e}_2]. \quad (3.3.44)$$

Wir nehmen im folgenden der Einfachheit halber an, dass zu dem betrachteten Zeitpunkt $\dot{\varphi}(t) > 0$ ist. Dann ist der Tangenteneinheitsvektor gerade durch die eckige Klammer gegeben, denn es ist

$$\vec{T}(t) = \frac{\vec{v}(t)}{|\vec{v}(t)|} = \frac{\vec{v}(t)}{R\dot{\varphi}} = -\sin \varphi(t)\vec{e}_1 + \cos \varphi(t)\vec{e}_2. \quad (3.3.45)$$

Die Beschleunigung ergibt sich wiederum durch Ableiten von (3.3.44) nach der Zeit

$$\vec{a}(t) = \dot{\vec{v}}(t) = R\ddot{\varphi}\vec{T}(t) - R\dot{\varphi}^2[\cos \varphi(t)\vec{e}_1 + \sin \varphi(t)\vec{e}_2]. \quad (3.3.46)$$

Andererseits ist

$$\frac{d\vec{T}}{ds} = \frac{d\vec{T}}{dt} \frac{1}{ds/dt} = -\dot{\varphi}[\cos \varphi(t)\vec{e}_1 + \sin \varphi(t)\vec{e}_2] \frac{1}{R\dot{\varphi}} = -\frac{1}{R}[\cos \varphi(t)\vec{e}_1 + \sin \varphi(t)\vec{e}_2] \quad (3.3.47)$$

Damit ist gemäß (3.3.2)

$$\vec{N}(t) = -\cos \varphi(t)\vec{e}_1 - \sin \varphi(t)\vec{e}_2 \quad \text{und} \quad \kappa = \frac{1}{\rho} = \left| \frac{d\vec{T}}{ds} \right| = \frac{1}{R} \Rightarrow \rho = R. \quad (3.3.48)$$

Wie wir sehen, gilt also in der Tat (3.3.42).

Dieses einfache Beispiel liefert sogleich auch die anschauliche Bedeutung von (3.3.42): Lokal können wir die Bewegung des Massenpunktes entlang seiner Bahn als Bewegung auf dem Schmiegekreis in dem gerade durchlaufenen Punkt beschreiben:

Die Zentripetalbeschleunigung ist auf den Mittelpunkt des Schmieglekreises hingerrichtet und zieht den Massenpunkt entsprechend der Krümmung der Kurve in die entsprechende Richtung. Der Betrag der Zentripetalbeschleunigung ist $\rho\dot{\varphi}^2$, wobei $d\varphi = dt\dot{\varphi}$ der während des kleinen Zeitschritts dt durchlaufene Winkel ist (s. Skizze). Man nennt $\dot{\varphi}$ auch die **momentane Winkelgeschwindigkeit** Bewegung, entsprechend der Bedeutung in dem eben betrachteten Beispiel einer Kreisbewegung. Weiter zeigt der Vergleich von (3.3.42) mit (3.3.46), dass die Tangentialbeschleunigung in diesem Bild einer momentanen Bewegung auf dem Schmieglekreis durch die **momentane Winkelbeschleunigung** $\ddot{\varphi}$ verursacht wird.

3.4 Skalare Felder, Gradient und Richtungsableitung

Im Folgenden betrachten wir **Felder**. Sie sind von eminenter Bedeutung für die gesamte Physik. Wir beginnen in diesem Abschnitt mit **skalaren Feldern**. Sie ordnen jedem Punkt im Raum, gegeben durch den **Ortsvektor** $\vec{x}$ bzgl. einem willkürlich vorgegebenen Punkt, dem Ursprung des Koordinatensystems, eine **reelle Größe** zu: $\Phi: E^3 \rightarrow \mathbb{R}, \vec{x} \mapsto \Phi(\vec{x})$.

Ein *Beispiel* ist die **Temperatur**: Wir können an jedem Ort im Raum die Temperatur messen und dadurch das **Temperaturfeld** definieren. Genauso ist die Dichte eines Mediums ein skalares Feld: An jedem Punkt des Raumes können wir die Masse Δm der Materie, in einem kleinen Volumenelement ΔV um diesen Punkt bestimmen. Dann ist die mittlere Dichte in diesem Volumenelement $\Delta m/\Delta V$. Im Limes $\Delta V \rightarrow 0$ erhalten wir die Dichte $\rho(\vec{x})$. In einem infinitesimal kleinen Volumenelement dV um den Punkt $\vec{x}$ ist dann Materie von der Masse $dm = dV\rho(\vec{x})$ enthalten. Wir werden im Folgenden oft physikalische Beispiele aus der **Kontinuumsmechanik** heranziehen, um die mathematischen Betrachtungen zu veranschaulichen.

Führen wir wieder eine beliebige rechtshändige Orthonormalbasis $(\vec{e}_j)$ ein, können wir das Feld als Funktion der drei Koordinaten bzw. als Funktion des entsprechenden Spaltenvektors $\underline{x} \in \mathbb{R}^3$ ansehen. Dies impliziert das Transformationsverhalten des Feldes unter orthogonalen Transformationen. Bezeichnen wir nämlich mit $\tilde{\Phi}$ das Skalarfeld als Funktion des Komponentenvektors $\underline{x}$ bzgl. der Orthonormalbasis $(\vec{e}_j)$ und mit $\tilde{\Phi}'$ als Funktion des Komponentenvektors $\underline{x}'$ bzgl. der Orthonormalbasis $(\vec{e}'_j)^2$, so muss definitionsgemäß

$$\tilde{\Phi}(\underline{x}) = \tilde{\Phi}'(\underline{x}') = \Phi(\vec{x}) \quad (3.4.1)$$

gelten, denn es ist klar, dass z.B. die Dichte eines Gases nicht von der Wahl des Koordinatensystems abhängt sondern lediglich von dem Punkt im Raum, an dem wir diese Dichte messen. Betrachten wir die Basistransformation von einer zur anderen Orthonormalbasis gemäß (3.4.4), folgt wegen (3.4.5)

$$\underline{x}' = \hat{T}\underline{x}, \quad \tilde{\Phi}'(\underline{x}') = \tilde{\Phi}(\underline{x}) = \tilde{\Phi}(\hat{T}^{-1}\underline{x}'). \quad (3.4.2)$$

Wir bezeichnen alle Funktionen $\mathbb{R}^3 \rightarrow \mathbb{R}$, die sich bzgl. orthogonaler Transformationen ihrer Argumente so verhalten wie in (3.4.2) angegeben als **skalare Felder**. Sie beschreiben Größen als Funktionen vom Ort.

Wie für Funktionen einer reellen Veränderlichen können wir auch für Funktionen, die von mehr als einer Variablen abhängen, **Ableitungen** definieren, die die Änderungen der Funktionen bei kleinen Änderungen ihrer Argumente beschreiben. Dazu definieren wir zunächst die **partiellen Ableitungen** der Funktion $\tilde{\Phi}(\underline{x}) \equiv$

²Gewöhnlich nehmen Physiker die subtile Unterscheidung in der Bezeichnung der Felder nicht vor und schreiben einfach $\Phi(\underline{x})$, d.h. sie verwenden dasselbe Symbol für die Funktion $\Phi: E^3 \rightarrow \mathbb{R}$ und die Funktion $\tilde{\Phi}: \mathbb{R}^3 \rightarrow \mathbb{R}$. Es ist aber gerade zu Beginn der Beschäftigung mit Vektoranalysis wichtig, diese Unterscheidung in Erinnerung zu behalten. Freilich definieren sich bei vorgegebener Orthonormalbasis $(\vec{e}_j)$, auf die sich die Komponenten $\underline{x}$ von $\vec{x}$ beziehen, diese beiden Funktionen umkehrbar eindeutig wechselseitig.

$\tilde{\Phi}(x_1, x_2, x_3)$. Dabei ist die partielle Ableitung nach der Variablen x_i durch den Limes des entsprechenden Differenzenquotienten bei Änderungen nur dieser einen Variablen x_i definiert, während die anderen beiden Variablen konstant gehalten werden. Für die partielle Ableitung nach x_1 bedeutet dies

$$\frac{\partial}{\partial x_1} \tilde{\Phi}(x_1, x_2, x_3) = \partial_1 \tilde{\Phi}(x_1, x_2, x_3) = \lim_{\Delta x_1 \rightarrow 0} \frac{\tilde{\Phi}(x_1 + \Delta x_1, x_2, x_3) - \tilde{\Phi}(x_1, x_2, x_3)}{\Delta x_1} \quad (3.4.3)$$

und entsprechend analog für die partiellen Ableitungen ∂_2 und ∂_3 nach x_2 bzw. x_3 .

Falls diese partiellen Ableitungen in einem gegebenen Punkt $\underline{x} = (x_1, x_2, x_3)^T$ existieren, heißt die Funktion $\tilde{\Phi}$ dort **partiell differenzierbar**. Ist sie in jedem Punkt eines bestimmten offenen Gebiets $G \subseteq \mathbb{R}^3$ partiell differenzierbar, heißt sie in diesem Gebiet **partiell differenzierbar**. Falls die partiellen Ableitungen in diesem Fall allesamt stetig sind, heißt die Funktion dort **stetig partiell differenzierbar**. In der Physik gehen wir davon aus, dass die meisten dort relevanten Funktionen bzw. Felder **lokal stetig differenzierbar** sind, d.h. bis auf einige **singuläre Punkte** existiert um jeden Punkt eines bestimmten Gebietes eine Kugel (bzw. **Umgebung**), in der die Funktion stetig differenzierbar ist. Manchmal bezeichnet man solche Funktionen auch als **lokal glatt**.

Es gelten für partielle Ableitungen freilich alle Rechenregeln wie für Ableitungen von Funktionen einer unabhängigen Variablen. Besonders wichtig ist, dass die partiellen Ableitungen **lineare Operatoren** sind, d.h. sind $\tilde{\Phi}_1$ und $\tilde{\Phi}_2$ lokal glatt, so gilt für beliebige $\lambda_1, \lambda_2 \in \mathbb{R}$ in jedem regulären Punkt der Funktion

$$\partial_j (\lambda_1 \tilde{\Phi}_1 + \lambda_2 \tilde{\Phi}_2) = \lambda_1 \partial_j \tilde{\Phi}_1 + \lambda_2 \partial_j \tilde{\Phi}_2. \quad (3.4.4)$$

Weitere wichtige Ableitungsregeln sind die **Produkt- und Quotientenregel**

$$\begin{aligned} \partial_j (\tilde{\Phi}_1 \tilde{\Phi}_2) &= (\partial_j \tilde{\Phi}_1) \tilde{\Phi}_2 + \tilde{\Phi}_1 (\partial_j \tilde{\Phi}_2), \\ \partial_j \left(\frac{\tilde{\Phi}_1}{\tilde{\Phi}_2} \right) &= \frac{(\partial_j \tilde{\Phi}_1) \tilde{\Phi}_2 - \tilde{\Phi}_1 \partial_j \tilde{\Phi}_2}{\tilde{\Phi}_2^2}, \end{aligned} \quad (3.4.5)$$

wobei die Quotientenregeln natürlich nur an Stellen sinnvoll ist, wo $\tilde{\Phi}_2 \neq 0$ ist.

Einer Erweiterung erfährt allerdings die Kettenregel. Dazu betrachten wir eine beliebige differenzierbare Kurve $\underline{x}(t)$, die durch ein Gebiet verläuft, wo das Skalarfeld $\tilde{\Phi}(\underline{x})$ stetig differenzierbar ist.

Dann können wir fragen, wie sich dieses Feld entlang dieser Bahnkurve ändert und die **totale Ableitung** bilden. Für diese gilt

$$\frac{d}{dt} \tilde{\Phi}[\underline{x}(t)] = \sum_{j=1}^3 \dot{x}_j \partial_j \tilde{\Phi}[\underline{x}(t)]. \quad (3.4.6)$$

Um diese Behauptung zu beweisen, betrachten wir den Differenzenquotienten

$$\frac{\Delta \tilde{\Phi}}{\Delta t} = \frac{\tilde{\Phi}[\underline{x}(t + \Delta t)] - \tilde{\Phi}[\underline{x}(t)]}{\Delta t}. \quad (3.4.7)$$

Setzen wir nun $\Delta \underline{x} = \underline{x}(t + \Delta t) - \underline{x}(t)$, ergibt sich

$$\frac{\Delta \tilde{\Phi}}{\Delta t} = \frac{\tilde{\Phi}(\underline{x} + \Delta \underline{x}) - \tilde{\Phi}(\underline{x})}{\Delta t}. \quad (3.4.8)$$

3. Vektoranalysis

Wir addieren und subtrahieren nun im Zähler dieses Ausdrucks geeignete Zusatzterme, die Differenzen ergeben, so dass sich im Argument jeweils immer nur eine der Variablen x_j ändert:

$$\begin{aligned} \frac{\Delta \tilde{\Phi}}{\Delta t} &= \frac{\tilde{\Phi}(x_1 + \Delta x_1, x_2 + \Delta x_2, x_3 + \Delta x_3) - \tilde{\Phi}(x_1, x_2 + \Delta x_2, x_3 + \Delta x_3)}{\Delta t} \\ &\quad + \frac{\tilde{\Phi}(x_1, x_2 + \Delta x_2, x_3 + \Delta x_3) - \tilde{\Phi}(x_1, x_2, x_3 + \Delta x_3)}{\Delta t} \\ &\quad + \frac{\tilde{\Phi}(x_1, x_2, x_3 + \Delta x_3) - \tilde{\Phi}(x_1, x_2, x_3)}{\Delta t}. \end{aligned} \quad (3.4.9)$$

Wir können nun den ersten Term wie folgt umformen:

$$\begin{aligned} &\frac{\tilde{\Phi}(x_1 + \Delta x_1, x_2 + \Delta x_2, x_3 + \Delta x_3) - \tilde{\Phi}(x_1, x_2 + \Delta x_2, x_3 + \Delta x_3)}{\Delta t} \\ &= \frac{\tilde{\Phi}(x_1 + \Delta x_1, x_2 + \Delta x_2, x_3 + \Delta x_3) - \tilde{\Phi}(x_1, x_2 + \Delta x_2, x_3 + \Delta x_3)}{\Delta x_1} \frac{\Delta x_1}{\Delta t}. \end{aligned} \quad (3.4.10)$$

Da $\tilde{\Phi}$ voraussetzungsgemäß in einer Umgebung von $\underline{x}$ stetig differenzierbar und auch die Kurve bei t differenzierbar ist, ergibt sich für diesen Ausdruck im Limes $\Delta t \rightarrow 0$

$$\frac{\tilde{\Phi}(x_1 + \Delta x_1, x_2 + \Delta x_2, x_3 + \Delta x_3) - \tilde{\Phi}(x_1, x_2 + \Delta x_2, x_3 + \Delta x_3)}{\Delta t} \rightarrow \partial_1 \tilde{\Phi}[\underline{x}(t)] \dot{x}_1. \quad (3.4.11)$$

Genauso geht man für die beiden anderen Terme in (3.4.9) vor, woraus schließlich die Behauptung (3.4.6) folgt.

Nun wollen wir zeigen, dass diese totale Ableitung eine von der Wahl des kartesischen Koordinatensystems unabhängige Bedeutung besitzt. Dazu bemerken wir, dass (3.4.6) in der Form

$$\frac{d}{dt} \tilde{\Phi}[\underline{x}(t)] = \dot{\underline{x}}(t) \cdot \underline{\nabla} \tilde{\Phi}[\underline{x}(t)] \quad (3.4.12)$$

geschrieben werden kann, wenn wir

$$\underline{\nabla} \tilde{\Phi} = \begin{pmatrix} \partial_1 \tilde{\Phi} \\ \partial_2 \tilde{\Phi} \\ \partial_3 \tilde{\Phi} \end{pmatrix} \quad (3.4.13)$$

definieren. Dies legt nahe, dass $\underline{\nabla} \tilde{\Phi}$ bzgl. *kartesischer Koordinaten* ein Vektor ist. Dies werden wir auch sogleich beweisen.

Man nennt diesen Vektor den **Gradienten** des Skalarfeldes $\Phi(\vec{x})$:

$$\text{grad } \Phi(\vec{x}) = \vec{\nabla} \Phi(\vec{x}) = \sum_{j=1}^3 \vec{e}_j \partial_j \Phi(\underline{x}). \quad (3.4.14)$$

Wir wollen nun zeigen, dass dieses **Vektorfeld** tatsächlich unabhängig von der Wahl des *kartesischen* Koordinatensystems ist. Betrachten wir dazu wieder die Basistransformation (2.3.17) bzw. (2.3.20) für die Basisvektoren. Verwenden wir dann (3.4.2), um die Summe auf der rechten Seite von (3.4.14) bzgl. der neuen Basis zu berechnen, folgt

$$\sum_{j=1}^3 \vec{e}'_j \partial'_j \tilde{\Phi}'(\underline{x}') = \sum_{j,k=1}^3 \vec{e}'_j \frac{\partial x_k}{\partial x'_j} \partial_k \tilde{\Phi}(\underline{x}). \quad (3.4.15)$$

3.4. Skalare Felder, Gradient und Richtungsableitung

Dabei haben wir die Kettenregel (3.4.6) für mehrere Veränderliche für die partielle Ableitung nach x'_j verwendet. Gemäß (2.3.17) und (2.3.20) gilt (*warum?*)

$$\underline{x}' = \hat{T} \underline{x} \Rightarrow \underline{x} = \hat{U} \underline{x}' \Rightarrow \frac{\partial x_k}{\partial x'_j} = U_{kj}. \quad (3.4.16)$$

Setzen wir dies in (3.4.15) ein, folgt

$$\sum_{j=1}^3 \vec{e}'_j \partial'_j \tilde{\Phi}'(\underline{x}') = \sum_{j,k=1}^3 \vec{e}'_j U_{kj} \partial_k \tilde{\Phi}(\underline{x}) = \sum_{j,k=1}^3 \vec{e}'_j T_{jk} \partial_k \tilde{\Phi}(\underline{x}). \quad (3.4.17)$$

Dabei haben wir im letzten Schritt $\hat{U} = \hat{T}^{-1} = \hat{T}^T$, also die **Orthogonalität** der Transformationsmatrizen, verwendet. Das bedeutet nämlich, dass $U_{kj} = T_{jk}$ ist. Gemäß (2.3.17) folgt

$$\sum_{j=1}^3 \vec{e}'_j \partial'_j \tilde{\Phi}'(\underline{x}') = \sum_{k=1}^3 \vec{e}_k \partial_k \tilde{\Phi}(\underline{x}) \stackrel{(3.4.14)}{=} \vec{\nabla} \Phi(\vec{x}). \quad (3.4.18)$$

Dies zeigt, dass der Gradient tatsächlich unabhängig von der gewählten Orthonormalbasis ist. Dieser besitzt also eine von der kartesischen Basis unabhängige Bedeutung.

Betrachten wir die Parametrisierung der Kurve in (3.4.6) als Funktion der Bogenlänge, folgt

$$\frac{d}{ds} \Phi[\vec{x}(s)] = \frac{d\vec{x}}{ds} \cdot \vec{\nabla} \Phi[\vec{x}(s)] = \vec{T}(s) \cdot \vec{\nabla} \Phi[\vec{x}(s)]. \quad (3.4.19)$$

Dabei ist $\vec{T}(s)$ der Tangenteneinheitsvektor an die Kurve. Die geometrische Bedeutung ist dabei klar: Bewegt man sich eine infinitesimale Strecke ds entlang der Kurve mit Tangenteneinheitsvektor $\vec{T}(s)$ an der betreffenden Stelle, so ändert sich der Wert des Skalarfeldes $ds \vec{T}(s) \cdot \vec{\nabla} \Phi[\vec{x}(s)]$. Diese Änderung ist dabei freilich unabhängig von der Wahl des kartesischen Koordinatensystems, das wir dazu benutzen, um den Gradienten über die partielle Ableitungen der entsprechenden Funktion $\tilde{\Phi}(\underline{x})$ zu berechnen.

Man bezeichnet daher (3.4.19) als die **Richtungsableitung** des Skalarfeldes an dem betreffenden Punkt in der durch $\vec{T}(s)$ gegebenen Richtung.

Aus der obigen Rechnung folgt noch das Verhalten der Darstellung des Gradienten durch die partiellen Ableitungen unter **orthogonalen Transformationen**. Wegen (3.4.17) ist nämlich

$$\partial'_j \tilde{\Phi}'(\underline{x}') = \sum_{k=1}^3 U_{kj} \partial_k \tilde{\Phi}(\underline{x}) = \sum_{k=1}^3 T_{jk} \partial_k \tilde{\Phi}(\underline{x}). \quad (3.4.20)$$

In Matrix-Vektor-Produktschreibweise bedeutet nun diese Gleichung

$$\underline{\nabla}' \tilde{\Phi}'(\underline{x}') = \hat{U}^T \underline{\nabla} \tilde{\Phi}(\underline{x}) = \hat{T} \underline{\nabla} \tilde{\Phi}(\underline{x}). \quad (3.4.21)$$

Das bedeutet, dass sich der Gradient bzgl. *orthogonaler Transformationen* tatsächlich wie ein Vektor verhält³.

³Wir bemerken zum Schluss noch, dass sich der Gradient bzgl. allgemeiner Basen nicht wie ein Vektor sondern wie ein sogenannte Kovektor, d.h. eine lineare Abbildung $E^3 \rightarrow \mathbb{R}$ transformiert.

3.5 Extremwertaufgaben

Wie bei den Funktionen einer Veränderlichen (vgl. Anhang 1.7.4) ist es auch für skalare Funktionen $f : \mathbb{R}^n \rightarrow \mathbb{R}$ oft interessant, **lokale Extremstellen** zu suchen. Wir können dieses Problem auf das entsprechende Problem für Funktionen einer Veränderlichen zurückführen. Dazu definieren wir

$$g(t) = f(\underline{x}_0 + t\underline{n}). \quad (3.5.1)$$

Dabei sind $\underline{x}_0$ und $\underline{n}$ beliebige Vektoren in $\mathbb{R}^n$. Wir nehmen im Folgenden an, die Funktion f sei zweimal stetig partiell differenzierbar. Aus den Betrachtungen in Anhang 1.7.4 wird klar, dass f an der Stelle $\underline{x}_0$ nur dann ein Extremum annehmen kann, wenn für alle $\underline{n} \in \mathbb{R}^n$

$$\dot{g}(0) = \underline{n} \cdot \underline{\nabla} f(\underline{x}_0) = 0 \quad (3.5.2)$$

ist.

Da wir für $\underline{n}$ beliebige Vektoren einsetzen dürfen, folgt als *notwendige Bedingung für das Vorliegen eines Extremums* bei $\underline{x}_0$, dass der Gradient der skalaren Funktion f dort verschwinden muss

$$\underline{\nabla} f(\underline{x}_0) = \underline{0}. \quad (3.5.3)$$

Weiter wissen wir aus Anhang 1.7.4, dass für das Vorliegen eines **lokalen Minimums** die Bedingung

$$\ddot{g}(0) = \sum_{i,j=1}^n n_i n_j \partial_i \partial_j f(\underline{x}_0) > 0 \quad (3.5.4)$$

hinreichend ist. Da f voraussetzungsgemäß zweimal stetig partiell differenzierbar ist, ist die **Hesse-Matrix**

$$H_{ij}(\underline{x}_0) = \partial_i \partial_j f(\underline{x}_0) \quad (3.5.5)$$

symmetrisch. Da (3.5.4) für alle $\underline{n} \in \mathbb{R}^n$ gelten muss, bedeutet dies, dass ein lokales Minimum vorliegt, falls die durch die Hessematrix definierte **quadratische Form** für alle $\underline{n} \neq \underline{0}$

$$Q(\underline{n}) = \sum_{i,j=1}^n H_{ij}(\underline{x}_0) n_i n_j > 0 \Leftrightarrow Q(\underline{n}) = \underline{n}^T \hat{H}(\underline{x}_0) \underline{n} > 0 \quad (3.5.6)$$

ist. Man nennt solche Bilinearformen **positiv definit**.

Wir kommen auf Kriterien für positive Definitheit von Bilinearformen im nächsten Kapitel noch zurück. Betrachten wir als einfachstes Beispiel die Funktion zweier Veränderlicher

$$f(\underline{x}) = 2x_1^2 + 3x_2^2. \quad (3.5.7)$$

Offenbar besitzt f in $\underline{x}_0 = \underline{0}$ ein lokales Minimum. Dies können wir direkt dem eben hergeleiteten hinreichenden Kriterium entnehmen. Der Gradient ist

$$\underline{\nabla} f(\underline{x}) = \begin{pmatrix} \partial_1 f(\underline{x}) \\ \partial_2 f(\underline{x}) \end{pmatrix} = \begin{pmatrix} 4x_1 \\ 6x_2 \end{pmatrix}. \quad (3.5.8)$$

Der Gradient verschwindet für $x_1 = x_{10} = x_2 = x_{20} = 0$. Die Hesse-Matrix ist offenbar $\hat{H}(\underline{x}_0) = \text{diag}(4, 6)$ und damit

$$Q(\underline{n}) = \underline{n}^T \hat{H}(\underline{x}_0) \underline{n} = 4n_1^2 + 6n_2^2 > 0 \quad (3.5.9)$$

3.5. Extremwertaufgaben

für alle $\underline{n} \neq 0$. Die Hesse-Matrix ist also tatsächlich positiv definit, und damit besitzt f bei $\underline{x}_0 = \underline{0}$ ein Minimum.

Nun ergibt sich bei Funktionen mehrerer Veränderlicher oft auch die Frage, wann sie unter bestimmten **Nebenbedingungen** extremal werden.

Ein praktisches Beispiel ist die Aufgabe, eine zylindrische Blechdose mit vorgegebenen Volumens V zu konstruieren, die eine möglichst kleine Oberfläche besitzt, d.h. den geringsten Materialaufwand erfordert. Ist R der Radius und h die Höhe der Blechdose, so sind Oberfläche und Volumen durch (*nachrechnen!*)

$$O(r, h) = 2\pi r h + 2\pi r^2, \quad V(r, h) = \pi r^2 h \quad (3.5.10)$$

gegeben. Wir suchen also diejenigen Werte für $\underline{x} = (r, h)$, dass $O(r, h) = O(\underline{x})$ minimal wird unter der Bedingung, dass $V(r, h) = V = \text{const}$ vorgegeben ist.

Wir lösen diese Aufgaben nun auf zwei Arten. In diesem Fall lässt sich nämlich die eine Variable (wir wählen hier h) aus der Nebenbedingung als Funktion der anderen Variablen (hier also r) darstellen:

$$h(r) = \frac{V}{\pi r^2}. \quad (3.5.11)$$

Demnach müssen wir also nur dasjenige r finden, für welches

$$\tilde{O}(r) = O[r, h(r)] = \frac{2V}{r} + 2\pi r^2 \quad (3.5.12)$$

minimal wird. Damit ist die Aufgabe auf eine Extremwertbestimmung für eine Funktion einer Veränderlichen $\tilde{O}(r)$ (freilich nun ohne Nebenbedingungen) zurückgeführt. Zunächst muss die erste Ableitung verschwinden:

$$\tilde{O}'(r) = -\frac{2V}{r^2} + 4\pi r \stackrel{!}{=} 0. \quad (3.5.13)$$

Dies ist erfüllt für

$$r = r_0 = \left(\frac{V}{2\pi}\right)^{1/3} \Rightarrow h = h(r_0) = \frac{V}{\pi r_0} = 2r_0. \quad (3.5.14)$$

Nun prüfen wir noch nach, dass diese Lösung tatsächlich ein Minimum der Oberfläche unter der Nebenbedingung konstanten Volumens ist. Dazu berechnen wir

$$\tilde{O}''(r_0) = +\frac{4V}{r_0^3} + 4\pi = 12\pi > 0, \quad (3.5.15)$$

d.h. $\tilde{O}$ wird tatsächlich für $r = r_0$ minimal. Demnach besitzt diejenige Büchse vorgegebenen Volumens V die kleinste Oberfläche, für deren Höhe $h = 2r_0 = (4V/\pi)^{1/3}$ ist.

In ähnlichen Extremwertaufgaben mit Nebenbedingungen kann es aber vorkommen, dass man nicht so einfach einen Teil der Variablen mit Hilfe der Nebenbedingungen eliminieren kann. In diesem Fall verwendet man die **Methode der Lagrange-Multiplikatoren**.

Wir führen dies anhand unseres Beispiels vor. Wir schreiben dazu die Nebenbedingung in der Form

$$N(r, h) = 0 \quad \text{mit} \quad N(r, h) = \pi r^2 h - V. \quad (3.5.16)$$

Dann betrachten wir die Extremwertaufgabe der Funktion

$$f(r, h, \lambda) = O(r, h) + \lambda N(r, h) \quad (3.5.17)$$

ohne Nebenbedingungen. Die notwendige Bedingung, dass diese Funktion extremal wird, verlangt, dass $\nabla f = (\partial_r f, \partial_b f, \partial_\lambda f)^T = \underline{0}$ wird. Aufgrund der Einführung des Lagrange-Parameters wird die Nebenbedingung also als eine der notwendigen Extremalbedingungen berücksichtigt. Wir erhalten aus diesen Bedingungen dieselbe Lösung für r und b wie oben in (3.5.14) angegeben. Auf hinreichende Bedingungen für das Vorliegen eines Extremums mit der Methode der Lagrange-Multiplikatoren gehen wir hier nicht ein. Die Frage, ob tatsächlich ein Minimum oder Maximum an der gefundenen Lösung der hinreichenden Bedingung vorliegt, ist also stets noch gesondert zu untersuchen.

3.6 Vektorfelder, Divergenz und Rotation

Es liegt nun nahe, neben skalaren ortsabhängigen Größen (Skalarfeldern) auch vektorielle ortsabhängige Größen, also **Vektorfelder**, zu betrachten. Typische Beispiele sind das Geschwindigkeitsfeld $\vec{v}(\vec{x})$ einer Flüssigkeit oder eines Gases: Es gibt die Geschwindigkeit eines sich gerade am Ort $\vec{x}$ befindlichen kleinen Flüssigkeitselements an. Andere typische Beispiele sind das elektrische oder magnetische Feld⁴.

Mathematisch gesehen haben wir es mit einer Abbildung $\vec{V} : E^3 \rightarrow V$, $\vec{x} \mapsto \vec{V}(\vec{x})$ zu tun. Um mit diesen Feldern bequem rechnen zu können, führen wir wieder eine beliebige kartesische rechtshändige Basis $(\vec{e}_j)$ ein. Dann ist

$$\vec{V}(\vec{x}) = \sum_{j=1}^3 \vec{e}_j V_j(\underline{x}) \quad \text{mit} \quad \underline{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix}. \quad (3.6.1)$$

Setzen wir wieder $\underline{V} = (V_1, V_2, V_3)^T$, impliziert also die Einführung der Orthonormalbasis $(\vec{e}_j)$ eine umkehrbar eindeutige Abbildung des Vektorfeldes $\vec{V} \in E^3$ auf ein Vektorfeld $\underline{V} : \mathbb{R}^3 \rightarrow \mathbb{R}^3$, $\underline{x} \mapsto \underline{V}(\underline{x})$.

Das Transformationsverhalten der Komponenten $\underline{V}(\underline{x})$ unter orthogonalen Basistransformationen ist dann gemäß (3.4.4) bzw. (3.4.5) gegeben:

$$\underline{V}'(\underline{x}') = \hat{T} \underline{V}(\underline{x}) = \hat{T} \underline{V}(\hat{T}^{-1} \underline{x}'). \quad (3.6.2)$$

Im vorigen Abschnitt haben wir gesehen, dass bzgl. orthogonaler Transformationen der Differentialoperator⁵ $\vec{\nabla}$ Vektorcharakter besitzt.

Es liegt dann nahe, dass die **Divergenz eines Vektorfeldes**

$$\text{div } \vec{V}(\vec{x}) = \vec{\nabla} \cdot \vec{V}(\vec{x}) = \underline{\nabla} \cdot \underline{V}(\underline{x}) = \sum_{j=1}^3 \partial_j V_j(\underline{x}) \quad (3.6.3)$$

ein Skalarfeld ist.

Wir können sofort zeigen, dass es sich dabei tatsächlich um ein Skalarfeld handelt, denn es gilt gemäß (3.4.21)

$$\underline{\nabla}' \cdot \underline{V}'(\underline{x}') = (\hat{T} \underline{\nabla}) \cdot \hat{T} \underline{V}(\underline{x}) = \underline{\nabla}^T \hat{T}^T \hat{T} \underline{V}(\underline{x}) = \underline{\nabla}^T \underline{V}(\underline{x}) = \underline{\nabla} \cdot \underline{V}(\underline{x}). \quad (3.6.4)$$

Die anschauliche Bedeutung dieses Differentialoperators werden wir weiter unten in Abschnitt 3.10 diskutieren.

⁴I.a. werden solche Felder auch noch von der Zeit abhängen. Hier betrachten wir aber nur die Ortsabhängigkeit zu einem festen Zeitpunkt und lassen entsprechend das Zeitargument der Einfachheit halber weg.

⁵Aufgrund der Form des Symbols wird dieser Operator als **Nabla-Operator** bezeichnet. Der Name kommt von einem hebräischen harfenähnlichen Saiteninstrument her, das diese typische Form aufweist.

Wir haben nun in (2.6.70) gesehen, dass das Vektorprodukt zweier Vektoren sich unter *orientierungserhaltenden orthogonalen Transformationen* (also *Drehungen*) wie ein Vektor verhält.

Daraus folgt, dass die **Rotation eines Vektorfeldes**

$$\text{rot } \vec{V} = \vec{\nabla} \times \vec{V} = \sum_{j=1}^3 \vec{e}_j [\nabla \times \underline{V}(\underline{x})]_j = \sum_{j,k,l=1}^3 \vec{e}_j \epsilon_{jkl} \partial_k V_l(\underline{x}) \quad (3.6.5)$$

ein Vektorfeld ist^a.

^aIn der englischsprachigen Literatur heißt die Rotation eines Vektorfeldes curl: $\text{curl } \vec{V} = \vec{\nabla} \times \vec{V}$.

Auf die anschauliche Bedeutung der Rotation als **Wirbeldichte** einer Flüssigkeitsströmung werden wir weiter unten noch zurückkommen.

In kartesischen Komponenten ausgeschrieben ergibt sich

$$\underline{\nabla} \times \underline{V} = \begin{pmatrix} \partial_2 V_3 - \partial_3 V_2 \\ \partial_3 V_1 - \partial_1 V_3 \\ \partial_1 V_2 - \partial_2 V_1 \end{pmatrix}. \quad (3.6.6)$$

Für das Rechnen mit diesen Differentialoperatoren haben sich im Wesentlichen zwei Kalküle etabliert: der **Nabla-Kalkül** und der **Ricci-Kalkül** (Gregorio Ricci-Curbastro, 1853-1925). Der letztere drückt alles mit Hilfe der Komponenten von Vektoren bzw. Vektorfeldern und partiellen Ableitungsoperatoren ∂_j aus. Das hat den Vorteil, dass man nur mit reellen skalaren Größen arbeitet und die üblichen Rechenregeln für Ableitungen direkt anwenden kann. Da dabei ständig Summen über Indizes, die paarweise in einem Ausdruck auftreten, vorkommen, lässt man gewöhnlich die entsprechenden Summenzeichen weg. Diese Schreibweise nennt man die **Einsteinsche Summationskonvention**, denn sie wurde im Zusammenhang mit entsprechenden Rechnungen in der Allgemeinen Relativitätstheorie von Albert Einstein eingeführt. Für das Skalarprodukt zweier Vektoren gilt also im Ricci-Kalkül

$$\vec{v} \cdot \vec{w} = \underline{v} \cdot \underline{w} = \sum_{j=1}^3 v_j w_j = v_j w_j = v_1 w_1 + v_2 w_2 + v_3 w_3. \quad (3.6.7)$$

Betrachten wir als Beispiel zum Ricci-Kalkül die Umformung des Ausdrucks $\text{div}(\Phi \vec{V})$ für ein Skalarfeld Φ und ein Vektorfeld $\vec{V}$:

$$\text{div}(\Phi \vec{V}) = \vec{\nabla} \cdot (\Phi \vec{V}) = \partial_j (\Phi V_j) = (\partial_j \Phi) V_j + \Phi \partial_j V_j = (\vec{\nabla} \Phi) \cdot \vec{V} + \Phi \vec{\nabla} \cdot \vec{V} = \text{grad } \Phi \cdot \vec{V} + \Phi \text{div } \vec{V}. \quad (3.6.8)$$

Dabei haben wir einfach die Produktregel für die partiellen Ableitungen der skalaren Funktionen Φ und V_j verwendet.

Im **Nabla-Kalkül** operiert man mit dem $\vec{\nabla}$ -Operator direkt. Bei der Anwendung der Produktregel ist dabei zu beachten, dass man oft nicht so einfach die Reihenfolge bei Vektoroperationen, insbesondere wenn Kreuzprodukte auftreten, verändern kann. Betrachten wir aber zunächst das obige Beispiel im Nabla-Kalkül. Im ersten Schritt schreiben wir die Produktregel in der Form

$$\vec{\nabla} \cdot (\Phi \vec{V}) = \vec{\nabla} \cdot (\Phi \vec{V}_c) + \vec{\nabla} \cdot (\Phi_c \vec{V}). \quad (3.6.9)$$

Dabei bedeutet der Index c an einem Feld, dass man es bei der Ableitung als konstant ansehen soll. Dann entspricht die obige Formel offenbar der Produktregel, ohne dass man die Reihenfolge der Produkte ändern

3. Vektoranalysis

muss. Jetzt kann man aber die üblichen Rechenregeln anwenden, um die jeweils konstant zu haltenden Teile außerhalb des Nabla-Symbols anzuordnen. Dann kann man den zwischenzeitlich eingeführten Index c wieder weglassen. In unserem Beispiel ist das sehr einfach, weil nur ein Skalar- und ein Vektorfeld involviert sind:

$$\vec{\nabla} \cdot (\Phi \vec{V}) = \vec{V} \cdot \vec{\nabla} \Phi + \Phi \vec{\nabla} \cdot \vec{V}, \quad (3.6.10)$$

was mit (3.6.8) übereinstimmt.

Betrachten wir noch ein etwas komplizierteres Beispiel, nämlich $\text{rot}(\vec{V} \times \vec{W})$ mit zwei Vektorfeldern $\vec{V}$ und $\vec{W}$. Verwenden wir diesmal zuerst den Nabla-Kalkül:

$$\text{rot}(\vec{V} \times \vec{W}) = \vec{\nabla} \times (\vec{V} \times \vec{W}) = \vec{\nabla} \times (\vec{V} \times \vec{W}_c) + \vec{\nabla} \times (\vec{V}_c \times \vec{W}). \quad (3.6.11)$$

Hier drängt sich die Anwendung der „bac-cab-Formel“ für das doppelte Kreuzprodukt auf, wobei man beachten muss, dass für ein Feld, auf das $\vec{\nabla}$ wirken soll (also ohne Index c nach Anwendung der entsprechenden Schreibweise für die Produktregel) das $\vec{\nabla}$ -Symbol links stehen muss, während es für Felder mit einem Index c rechts von diesem Feld stehen muss. Rechnen wir entsprechend unser Beispiel weiter, erhalten wir

$$\text{rot}(\vec{V} \times \vec{W}) = (\vec{W} \cdot \vec{\nabla}) \vec{V} - \vec{W}(\vec{\nabla} \cdot \vec{V}) + \vec{V} \vec{\nabla} \cdot \vec{W} - (\vec{V} \cdot \vec{\nabla}) \vec{W}. \quad (3.6.12)$$

Dieselbe Rechnung im Ricci-Kalkül erfordert die Verwendung des Levi-Civita-Symbols, um die Kreuzprodukte auszudrücken:

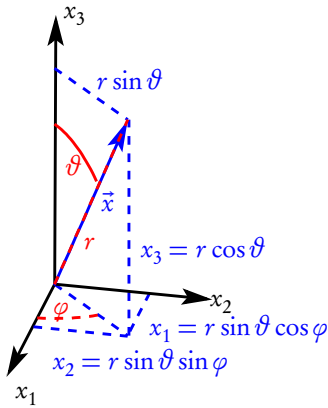
$$\begin{aligned} [\text{rot}(\vec{V} \times \vec{W})]_j &= \epsilon_{jab} \partial_a (\vec{V} \times \vec{W})_b \\ &= \epsilon_{jab} \partial_a (\epsilon_{bcd} V_c W_d) \\ &= \epsilon_{jab} \epsilon_{bcd} \partial_a (V_c W_d) \\ &= \epsilon_{jab} \epsilon_{cdb} (W_d \partial_a V_c + V_c \partial_a W_d) \\ &\stackrel{(2.5.14)}{=} (\delta_{jc} \delta_{ad} - \delta_{jd} \delta_{ac}) (W_d \partial_a V_c + V_c \partial_a W_d) \\ &= W_a \partial_a V_j - W_j \partial_a V_a + V_j \partial_a W_a - V_a \partial_a W_j \\ &= [(\vec{W} \cdot \vec{\nabla}) \vec{V}]_j - W_j \vec{\nabla} \cdot \vec{V} + V_j \vec{\nabla} \cdot \vec{W} - [(\vec{V} \cdot \vec{\nabla}) \vec{W}]_j, \end{aligned} \quad (3.6.13)$$

was offenbar das gleiche bedeutet wie (3.6.12).

3.7 Krummlinige Orthogonalkoordinaten

In den Anwendungen der Vektoranalysis auf physikalische Probleme⁶ ist es aus Symmetriegründen oft vorteilhaft, für den Ortsvektor andere als kartesische Koordinaten, sog. **generalisierte Koordinaten** (q_1, q_2, q_3) zu verwenden. In diesem Abschnitt betrachten wir **Kugel- und Zylinderkoordinaten** als die beiden wichtigsten Standardfälle.

3.7.1 Kugelkoordinaten



In der nebenstehenden Abbildung sind die Kugelkoordinaten durch (r, ϑ, φ) definiert. Man liest den Zusammenhang zwischen den Kugelkoordinaten und den kartesischen Koordinaten

$$\underline{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} r \sin \vartheta \cos \varphi \\ r \sin \vartheta \sin \varphi \\ r \cos \vartheta \end{pmatrix} \quad (3.7.1)$$

unmittelbar aus der Skizze ab. Offenbar liefern die Kugelkoordinaten eine umkehrbar eindeutige Parametrisierung für die Definitionsbereiche

$$r > 0, \quad \vartheta \in (0, \pi), \quad \varphi \in [0, 2\pi), \quad (3.7.2)$$

d.h. sie umfassen alle $\underline{x} \in \mathbb{R}^3$ außer die x_3 -Achse, entlang derer die Kugelkoordinaten singular sind.

In jedem der regulären Punkte können wir nun die drei Tangentenvektoren der Koordinatenlinien definieren. Dabei ist eine Koordinatenlinie dadurch definiert, dass man zwei der Kugelkoordinaten konstant lässt und die dritte variabel. Die drei Tangentenvektoren ergeben sich zu

$$\underline{T}_r = \partial_r \underline{x} = \begin{pmatrix} \sin \vartheta \cos \varphi \\ \sin \vartheta \sin \varphi \\ \cos \vartheta \end{pmatrix}, \quad \underline{T}_\vartheta = \partial_\vartheta \underline{x} = \begin{pmatrix} r \cos \vartheta \cos \varphi \\ r \cos \vartheta \sin \varphi \\ -r \sin \vartheta \end{pmatrix}, \quad \underline{T}_\varphi = \partial_\varphi \underline{x} = \begin{pmatrix} -r \sin \vartheta \sin \varphi \\ r \sin \vartheta \cos \varphi \\ 0 \end{pmatrix}. \quad (3.7.3)$$

Direktes Nachrechnen zeigt, dass diese Tangentenvektoren jeweils paarweise aufeinander senkrecht stehen, d.h. es ist $\underline{T}_j \cdot \underline{T}_k = 0$ für $j \neq k \in \{r, \vartheta, \varphi\}$. Die Kugelkoordinaten sind demnach sog. **krummlinige Orthogonalkoordinaten**. Wie sich gleich zeigen wird, ist es nun vorteilhaft, an jedem Punkt im Raum, die kartesische Basis zu verwenden, die durch Normierung der Tangentenvektoren entstehen.

In unserem Fall sind die entsprechenden Normierungsfaktoren durch

$$|\underline{T}_r| = g_r = 1, \quad |\underline{T}_\vartheta| = g_\vartheta = r, \quad |\underline{T}_\varphi| = g_\varphi = r \sin \vartheta \quad (3.7.4)$$

und die entsprechenden kartesischen Basisvektoren durch

$$\underline{e}'_r = \frac{1}{g_r} \underline{T}_r = \begin{pmatrix} \sin \vartheta \cos \varphi \\ \sin \vartheta \sin \varphi \\ \cos \vartheta \end{pmatrix}, \quad \underline{e}'_\vartheta = \frac{1}{g_\vartheta} \underline{T}_\vartheta = \begin{pmatrix} \cos \vartheta \cos \varphi \\ \cos \vartheta \sin \varphi \\ -\sin \vartheta \end{pmatrix}, \quad \underline{e}'_\varphi = \frac{1}{g_\varphi} \underline{T}_\varphi = \begin{pmatrix} -\sin \varphi \\ \cos \varphi \\ 0 \end{pmatrix} \quad (3.7.5)$$

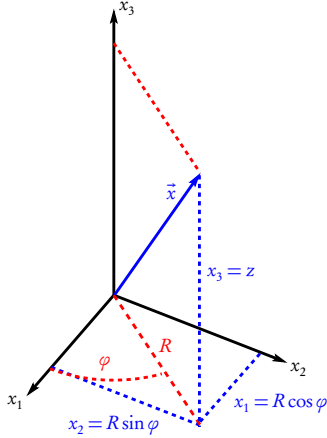
gegeben.

⁶In den Theorievorlesungen für L3 vor allem in der Elektrodynamik, die in Theoretische Physik 2 behandelt wird.

Wir bemerken, dass diese kartesischen Basen in der aufgezählten Reihenfolge überall rechtshändig sind, denn es ist

$$\text{vol}(\underline{e}'_r, \underline{e}'_\vartheta, \underline{e}'_\varphi) = \underline{e}'_r \cdot (\underline{e}'_\vartheta \times \underline{e}'_\varphi) = +1. \quad (3.7.6)$$

3.7.2 Zylinderkoordinaten



In der nebenstehenden Abbildungen sind die **Zylinderkoordinaten** (R, φ, z) definiert. Man liest den Zusammenhang zwischen den Zylinderkoordinaten und den kartesischen Koordinaten

$$\underline{x} = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} R \cos \varphi \\ R \sin \varphi \\ z \end{pmatrix} \quad (3.7.7)$$

unmittelbar aus der Skizze ab. Offenbar liefern die Zylinderkoordinaten eine umkehrbar eindeutige Parametrisierung für die Definitionsbereiche

$$R > 0, \quad \varphi \in [0, 2\pi), \quad z \in \mathbb{R}, \quad (3.7.8)$$

d.h. sie umfassen alle $\underline{x} \in \mathbb{R}^3$ außer die x_3 -Achse, entlang derer die Zylinderkoordinaten singulär sind.

Die drei Tangentenvektoren der Koordinatenlinien ergeben sich zu

$$\underline{T}_R = \partial_R \underline{x} = \begin{pmatrix} \cos \varphi \\ \sin \varphi \\ 0 \end{pmatrix}, \quad \underline{T}_\varphi = \begin{pmatrix} -R \sin \varphi \\ R \cos \varphi \\ 0 \end{pmatrix}, \quad \underline{T}_z = \partial_z \underline{x} = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}. \quad (3.7.9)$$

Direktes Nachrechnen zeigt, dass diese Tangentenvektoren jeweils paarweise aufeinander senkrecht stehen, d.h. es ist $\underline{T}_j \cdot \underline{T}_k = 0$ für $j \neq k \in \{R, \varphi, z\}$.

Auch die Zylinderkoordinaten sind demnach **krummlinige Orthogonalkoordinaten**. Die Normen der Tangentenvektoren ergeben sich zu

$$|\underline{T}_R| = g_R = 1, \quad |\underline{T}_\varphi| = g_\varphi = R, \quad |\underline{T}_z| = g_z = 1, \quad (3.7.10)$$

d.h. die entsprechenden kartesischen Basisvektoren sind

$$\underline{e}'_R = \frac{1}{g_R} \underline{T}_R = \begin{pmatrix} \cos \varphi \\ \sin \varphi \\ 0 \end{pmatrix}, \quad \underline{e}'_\varphi = \frac{1}{g_\varphi} \underline{T}_\varphi = \begin{pmatrix} -\sin \varphi \\ \cos \varphi \\ 0 \end{pmatrix}, \quad \underline{e}'_z = \frac{1}{g_z} \underline{T}_z = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}. \quad (3.7.11)$$

Wir bemerken, dass diese kartesischen Basen in der aufgezählten Reihenfolge überall rechtshändig sind, denn es ist

$$\text{vol}(\underline{e}'_R, \underline{e}'_\varphi, \underline{e}'_z) = \underline{e}'_R \cdot (\underline{e}'_\varphi \times \underline{e}'_z) = +1. \quad (3.7.12)$$

3.7.3 Allgemeine krummlinige Orthogonalkoordinaten

Wir betrachten nun allgemein irgendwelche generalisierten Koordinaten $q = (q_1, q_2, q_3) \in G \subseteq \mathbb{R}^3$, die durch eine Abbildung $\vec{x} : G \rightarrow E^3$ einen bestimmten Raumbereich (ggf. auch den ganzen Raum) umkehrbar eindeutig parametrisieren. Hält man zwei dieser drei Koordinaten fest, erhält man die **Koordinatenlinien**, die

in jedem Punkt von G drei Tangentenvektoren

$$\vec{T}_j(q) = \frac{\partial \vec{x}(q)}{\partial q_j} \quad (3.7.13)$$

definieren.

Diese generalisierten Koordinaten heißen **krummlinige Orthogonalkoordinaten**, falls in jedem Punkt diese Tangentenvektoren aufeinander senkrecht stehen, d.h. es gibt Funktionen $g_j : G \rightarrow \mathbb{R}$ mit $g_j(q) > 0$, so dass

$$\vec{T}_j(q) \cdot \vec{T}_k(q) = g_j^2(q) \delta_{jk} \quad (3.7.14)$$

gilt. Wir definieren weiter die **Einheitsvektoren**

$$\vec{e}'_j(q) = \frac{1}{g_j(q)} \vec{T}_j(q). \quad (3.7.15)$$

Sie bilden in jedem Punkt offenbar **Orthonormalsysteme**. Im Gegensatz zu den kartesischen Koordinatensystemen⁷ hängen sie aber i.a. vom Ort ab.

Im Folgenden nehmen wir an, dass diese Einheitsvektoren in jedem Punkt ein **rechtshändiges Dreibein** sind, d.h. dass überall

$$\vec{e}'_1 \times \vec{e}'_2 = \vec{e}'_3, \quad \vec{e}'_2 \times \vec{e}'_3 = \vec{e}'_1, \quad \vec{e}'_3 \times \vec{e}'_1 = \vec{e}'_2 \quad (3.7.16)$$

gilt. Dies können wir mit Hilfe des Levi-Civita-Symbols auch in der Form

$$\frac{1}{2} \sum_{jk} \epsilon_{jkl} \vec{e}'_j \times \vec{e}'_k = \vec{e}'_l \Leftrightarrow \vec{e}'_j \times \vec{e}'_k = \sum_l \epsilon_{jkl} \vec{e}'_l \quad (3.7.17)$$

schreiben.

3.7.4 Die Differentialoperatoren in krummlinigen Orthogonalkoordinaten

Um nun die Differentialoperatoren grad, div und rot mittels Ableitungen nach den generalisierten Koordinaten q_j auszudrücken, betrachten wir zuerst den Gradienten eines Skalarfeldes. Sei dazu $\vec{e}_j$ eine kartesische rechtshändige Basis. Dann gilt für ein beliebiges Skalarfeld Φ und mit $x_j = \vec{e}_j \cdot \vec{x}$

$$\vec{\nabla} \Phi = \sum_j \vec{e}_j \frac{\partial \Phi}{\partial x_j}. \quad (3.7.18)$$

Die Komponenten dieses Vektorfeldes bzgl. des Basissystems $\vec{e}'_k$ im betreffenden Punkt sind dann

$$\vec{e}'_k \cdot \vec{\nabla} \Phi = \sum_j \vec{e}'_k \cdot \vec{e}_j \frac{\partial \Phi}{\partial x_j} \stackrel{(3.7.15)}{=} \sum_j \frac{1}{g_k} \frac{\partial \vec{x}}{\partial q_k} \cdot \vec{e}_j \frac{\partial \Phi}{\partial x_j}. \quad (3.7.19)$$

Da nun die kartesischen Basisvektoren $\vec{e}_j$ nicht vom Ort abhängen, gilt

$$\frac{\partial \vec{x}}{\partial q_k} \cdot \vec{e}_j = \frac{\partial (\vec{e}_j \cdot \vec{x})}{\partial q_k} = \frac{\partial x_j}{\partial q_k}. \quad (3.7.20)$$

⁷Es ist klar, dass die kartesischen Komponenten des Ortsvektors bzgl. einer beliebigen kartesischen Basis auch Spezialfälle krummliniger Orthogonalkoordinaten sind.

3. Vektoranalysis

Setzen wir dies in (3.7.19) ein und verwenden die Kettenregel, folgt

$$\vec{e}'_k \cdot \vec{\nabla} \Phi = \frac{1}{g_k} \frac{\partial \Phi}{\partial q_k}, \quad (3.7.21)$$

d.h. der Gradient eines Skalarfeldes lässt sich durch Ableitungen bzgl. der orthogonalen krummlinigen Koordinaten q_k in der Form

$$\vec{\nabla} \Phi = \sum_k \vec{e}'_k \frac{1}{g_k} \frac{\partial \Phi}{\partial q_k} \quad (3.7.22)$$

berechnen.

Setzen wir nun $\Phi = q_j$, erhalten wir daraus

$$\vec{\nabla} q_j = \sum_k \vec{e}'_k \frac{1}{g_k} \frac{\partial q_j}{\partial q_k} = \sum_k \vec{e}'_k \frac{1}{g_k} \delta_{jk} = \frac{1}{g_j} \vec{e}'_j. \quad (3.7.23)$$

Sei nun

$$\vec{A} = \sum_k \vec{e}'_k A'_k \quad (3.7.24)$$

ein beliebiges Vektorfeld. Die $A'_k = \vec{e}'_k \cdot \vec{A}$ sind die Komponenten von $\vec{A}$ bzgl. des in jedem Punkt kartesischen Basissystems $\vec{e}'_k$. Berechnen wir nun die Rotation des Vektorfeldes, ergibt sich

$$\begin{aligned} \vec{\nabla} \times \vec{A} &= \sum_k \vec{\nabla} \times (\vec{e}'_k A'_k) \\ &\stackrel{(3.7.23)}{=} \sum_k \vec{\nabla} \times (g_k A'_k \vec{\nabla} q_k) \\ &= \sum_k [\vec{\nabla}(g_k A'_k) \times \vec{\nabla} q_k + g_k A'_k \vec{\nabla} \times \vec{\nabla} q_k] \\ &= \sum_k \vec{\nabla}(g_k A'_k) \times \vec{\nabla} q_k \\ &\stackrel{(3.7.23)}{=} \sum_k \vec{\nabla}(g_k A'_k) \times \frac{1}{g_k} \vec{e}'_k \\ &\stackrel{(3.7.22)}{=} \sum_{j,k} \vec{e}'_j \frac{1}{g_j} \frac{\partial (g_k A'_k)}{\partial q_j} \times \frac{1}{g_k} \vec{e}'_k \\ &= \sum_{j,k} \frac{1}{g_j g_k} \frac{\partial (g_k A'_k)}{\partial q_j} \vec{e}'_j \times \vec{e}'_k \\ &\stackrel{(3.7.17)}{=} \sum_{j,k,l} \frac{1}{g_j g_k} \frac{\partial (g_k A'_k)}{\partial q_j} \epsilon_{jkl} \vec{e}'_l. \end{aligned} \quad (3.7.25)$$

Betrachten wir nun den Beitrag der Summe mit $l = 1$, ergeben sich Beiträge für $j = 2, k = 3$ und $j = 3, k = 2$ und analog für $l = 2$ und $l = 3$. Werten wir dann die entsprechenden Levi-Civita-Symbole aus, ergibt sich

$$\begin{aligned}\vec{\nabla} \times \vec{A} = & \frac{\vec{e}'_1}{g_2 g_3} \left[\frac{\partial(g_3 A'_2)}{\partial q_2} - \frac{\partial(g_2 A'_1)}{\partial q_3} \right] \\ & + \frac{\vec{e}'_2}{g_3 g_1} \left[\frac{\partial(g_1 A'_3)}{\partial q_3} - \frac{\partial(g_3 A'_1)}{\partial q_1} \right] \\ & + \frac{\vec{e}'_3}{g_1 g_2} \left[\frac{\partial(g_2 A'_2)}{\partial q_1} - \frac{\partial(g_1 A'_2)}{\partial q_2} \right].\end{aligned}\quad (3.7.26)$$

Damit haben wir die Rotation mittels der krummlinigen Orthogonalkoordinaten ausgedrückt.

Zur Berechnung der Divergenz gehen wir analog vor, wobei wir hier einen Trick anwenden müssen, indem wir ausnützen, dass die $\vec{e}'_j$ ein rechtshändiges kartesisches Koordinatensystem bilden:

$$\begin{aligned}\vec{\nabla} \cdot \vec{A} &= \sum_k \vec{\nabla} \cdot (\vec{e}'_k A'_k) \\ &\stackrel{(3.7.17)}{=} \frac{1}{2} \sum_{k,l,m} \epsilon_{klm} \vec{\nabla} \cdot (\vec{e}'_l \times \vec{e}'_m A'_k) \\ &\stackrel{(3.7.23)}{=} \frac{1}{2} \sum_{k,l,m} \epsilon_{klm} \vec{\nabla} \cdot (g_l g_m A'_k \vec{\nabla} q_l \times \vec{\nabla} q_m).\end{aligned}\quad (3.7.27)$$

Nun ist

$$\vec{\nabla} \cdot (g_l g_m A'_k \vec{\nabla} q_l \times \vec{\nabla} q_m) = (\vec{\nabla} q_l \times \vec{\nabla} q_m) \cdot \vec{\nabla} (g_l g_m A'_k) + g_l g_m A'_k \vec{\nabla} \cdot (\vec{\nabla} q_l \times \vec{\nabla} q_m). \quad (3.7.28)$$

Der letzte Term verschwindet wegen (*Nachrechnen!*)

$$\vec{\nabla} \cdot (\vec{\nabla} q_l \times \vec{\nabla} q_m) = (\vec{\nabla} \times \vec{\nabla} q_l) \cdot \vec{\nabla} q_m - (\vec{\nabla} \times \vec{\nabla} q_m) \cdot \vec{\nabla} q_l = 0. \quad (3.7.29)$$

Setzen wir also (3.7.28) unter Beachtung von (3.7.29) in (3.7.27) ein, erhalten wir

$$\begin{aligned}\vec{\nabla} \cdot \vec{A} &= \frac{1}{2} \sum_{k,l,m} \epsilon_{klm} (\vec{\nabla} q_l \times \vec{\nabla} q_m) \cdot \vec{\nabla} (g_l g_m A'_k) \\ &\stackrel{(3.7.23)}{=} \frac{1}{2} \sum_{k,l,m} \epsilon_{klm} \frac{1}{g_l g_m} (\vec{e}'_l \times \vec{e}'_m) \cdot \vec{\nabla} (g_l g_m A'_k) \\ &\stackrel{(3.7.22)}{=} \frac{1}{2} \sum_{k,l,m,n} \epsilon_{klm} \frac{1}{g_l g_m} (\vec{e}'_l \times \vec{e}'_m) \cdot \vec{e}'_n \frac{1}{g_n} \frac{\partial(g_l g_m A'_k)}{\partial q_n}.\end{aligned}\quad (3.7.30)$$

Da die $\vec{e}'_k$ stets ein rechtshändiges Orthonormalsystem bilden, gilt

$$(\vec{e}'_l \times \vec{e}'_m) \cdot \vec{e}'_n = \epsilon_{lmn} \quad (3.7.31)$$

und damit

$$\begin{aligned}\vec{\nabla} \cdot \vec{A} &= \frac{1}{2} \sum_{k,l,m,n} \epsilon_{klm} \epsilon_{lmn} \frac{1}{g_l g_m g_n} \frac{\partial(g_l g_m A'_k)}{\partial q_k} \\ &= \sum_{k,n} \delta_{kn} \frac{1}{g_l g_m g_n} \frac{\partial(g_l g_m A'_k)}{\partial q_k} \\ &= \sum_k \frac{1}{g_k g_l g_m} \frac{\partial(g_l g_m A'_k)}{\partial q_k}.\end{aligned}\quad (3.7.32)$$

3. Vektoranalysis

Dabei haben wir in dem Schritt von der 1. zur 2. Zeile über l und m summiert. Wegen der beiden Levi-Civita-Symbole tragen nur Summanden mit $l \neq m$ bei. Da diese Indizes alle in der Menge $\{1, 2, 3\}$ sind kann es nur einen Beitrag geben, wenn $k = n$ den dritten möglichen Wert annimmt. Für $k = n = 1$ tragen nur die Fälle $l = 2, m = 3$ und $l = 3, m = 2$ bei. Die Levi-Civita-Symbole zusammen liefern dann einen Faktor 2. Wir dürfen also einfach die Levi-Civita-Symbole durch einen Faktor $2\delta_{kn}$ ersetzen und die Summen über l und m weglassen, wie in der 2. Zeile geschrieben, und man muss (k, l, m) jeweils als die geraden Permutationen von $(1, 2, 3)$ lesen, d.h. die Summation ist über $(k, l, m) = (1, 2, 3), (2, 3, 1), (3, 1, 2)$ auszuführen. Die Summe ausgeschrieben liefert also schließlich

$$\vec{\nabla} \cdot \vec{A} = \frac{1}{g_1 g_2 g_3} \left[\frac{\partial(g_2 g_3 A'_1)}{\partial q_1} + \frac{\partial(g_3 g_1 A'_2)}{\partial q_2} + \frac{\partial(g_1 g_2 A'_3)}{\partial q_3} \right]. \quad (3.7.33)$$

Wir stellen schließlich zur Übersicht die Differentialoperatoren in **Kugel- und Zylinderkoordinaten** zusammen.

Kugelkoordinaten (r, ϑ, φ) :

$$\vec{\nabla} U = \text{grad } U = \vec{e}_r' \frac{\partial U}{\partial r} + \vec{e}_\vartheta' \frac{1}{r} \frac{\partial U}{\partial \vartheta} + \vec{e}_\varphi' \frac{1}{r \sin \vartheta} \frac{\partial U}{\partial \varphi}, \quad (3.7.34)$$

$$\vec{\nabla} \cdot \vec{A} = \text{div } \vec{A} = \frac{1}{r^2} \frac{\partial(r^2 A'_r)}{\partial r} + \frac{1}{r \sin \vartheta} \frac{\partial(\sin \vartheta A'_\vartheta)}{\partial \vartheta} + \frac{1}{r \sin \vartheta} \frac{\partial A'_\varphi}{\partial \varphi}, \quad (3.7.35)$$

$$\begin{aligned} \vec{\nabla} \times \vec{A} = \text{rot } \vec{A} = & \vec{e}_r' \frac{1}{r \sin \vartheta} \left[\frac{\partial(\sin \vartheta A'_\varphi)}{\partial \vartheta} - \frac{\partial A'_\vartheta}{\partial \varphi} \right] + \vec{e}_\vartheta' \left[\frac{1}{r \sin \vartheta} \frac{\partial A'_r}{\partial \varphi} - \frac{1}{r} \frac{\partial(r A'_\varphi)}{\partial r} \right] \\ & + \vec{e}_\varphi' \frac{1}{r} \left[\frac{\partial(r A'_\vartheta)}{\partial r} - \frac{\partial A'_r}{\partial \vartheta} \right]. \end{aligned} \quad (3.7.36)$$

Wenden wir schließlich (3.7.34) und (3.7.35) hintereinander an, erhalten wir für den **Laplace-Operator**

$$\begin{aligned} \Delta U = \vec{\nabla} \cdot \vec{\nabla} U = \text{div grad } U = & \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\partial U}{\partial r} \right) + \frac{1}{r^2 \sin \vartheta} \frac{\partial}{\partial \vartheta} \left(\sin \vartheta \frac{\partial U}{\partial \vartheta} \right) + \frac{1}{r^2 \sin^2 \vartheta} \frac{\partial^2 U}{\partial \varphi^2} \\ = & \frac{1}{r} \frac{\partial^2}{\partial r^2} (r U) + \frac{1}{r^2 \sin \vartheta} \frac{\partial}{\partial \vartheta} \left(\sin \vartheta \frac{\partial U}{\partial \vartheta} \right) + \frac{1}{r^2 \sin^2 \vartheta} \frac{\partial^2 U}{\partial \varphi^2}. \end{aligned} \quad (3.7.37)$$

Zylinderkoordinaten (R, φ, z) :

$$\vec{\nabla} U = \text{grad } U = \vec{e}_R' \frac{\partial U}{\partial R} + \vec{e}_\varphi' \frac{1}{R} \frac{\partial U}{\partial \varphi} + \frac{\partial U}{\partial z} \vec{e}_z' \quad (3.7.38)$$

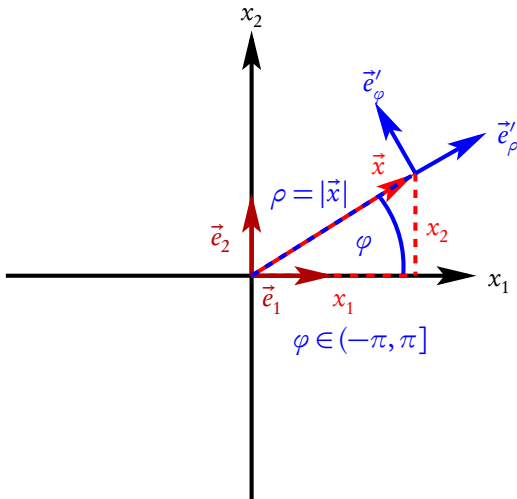
$$\vec{\nabla} \cdot \vec{A} = \text{div } \vec{A} = \frac{1}{R} \frac{\partial (R A_R')}{\partial R} + \frac{1}{R} \frac{\partial A_\varphi'}{\partial \varphi} + \frac{\partial A_z'}{\partial z}, \quad (3.7.39)$$

$$\vec{\nabla} \times \vec{A} = \text{rot } \vec{A} = \vec{e}_R' \left[\frac{1}{R} \frac{\partial A_z'}{\partial \varphi} - \frac{\partial A_\varphi'}{\partial z} \right] + \vec{e}_\varphi' \left[\frac{\partial A_R'}{\partial z} - \frac{\partial A_z'}{\partial R} \right] + \vec{e}_z' \frac{1}{R} \left[\frac{\partial (R A_\varphi')}{\partial R} - \frac{\partial A_R'}{\partial \varphi} \right], \quad (3.7.40)$$

$$\Delta U = \frac{1}{R} \frac{\partial}{\partial R} \left(R \frac{\partial U}{\partial R} \right) + \frac{1}{R^2} \frac{\partial^2 U}{\partial \varphi^2} + \frac{\partial^2 U}{\partial z^2}. \quad (3.7.41)$$

3.7.5 Polarkoordinaten in der Ebene

Natürlich kann man krummlinige Orthogonalkoordinaten auch in der Ebene einführen. Hier sollen die Einheitsvektoren $\vec{e}_j'$ ($j \in \{1, 2\}$) so orientiert sein, dass man durch Drehung von $\vec{e}_1'$ **im Gegenuhrzeigersinn** um den Drehwinkel $\pi/2$ auf $\vec{e}_2'$ kommt. Alternativ kann man die Polarkoordinaten auch in der 12-Ebene eines rechtshändigen kartesischen Koordinatensystems im Raum auffassen, dann gilt für die entsprechenden Einheitsvektoren der positiv orientierte krummlinige Koordinaten in der Ebene stets $\vec{e}_1' \times \vec{e}_2' = \vec{e}_3 = \vec{e}_3'$. Man nennt solche krummlinigen Koordinaten im Raum auch „verallgemeinerte Zylinderkoordinaten“, da $\vec{e}_1'$ und $\vec{e}_2'$ für alle Komponenten x_3 des Ortsvektors $\vec{x}$ gleich sind und $\vec{e}_3'$ ist ein konstanter (also ortsunabhängiger) Einheitsvektor.



Das wichtigste Beispiel für ebene krummlinige Orthogonalkoordinaten sind die **Polarkoordinaten**. Man verwendet zur Charakterisierung des Ortsvektors $\vec{x}$ zunächst ein im obigen Sinne positiv orientiertes Orthonormalsystem $(\vec{e}_1, \vec{e}_2)$ (s. nebenstehende Skizze) und definiert dann als krummlinige Koordinaten $\rho = |\vec{x}| = \sqrt{x_1^2 + x_2^2}$ und den Winkel $\varphi \in (-\pi, \pi]$ zwischen $\vec{e}_1$ und $\vec{x}$, wobei man Winkel in der unteren Halbebene negativ zählt. Aus der Skizze lesen wir dann unmittelbar ab, dass

$$\vec{x} = \rho(\cos \varphi \vec{e}_1 + \sin \varphi \vec{e}_2) \quad (3.7.42)$$

ist. Die Tangentenvektoren an die Koordinatenlinien sind (nachrechnen!)

$$\begin{aligned} \vec{T}_\rho &= \frac{\partial \vec{x}}{\partial \rho} = \cos \varphi \vec{e}_1 + \sin \varphi \vec{e}_2, & g_\rho &= |\vec{T}_\rho| = 1, \\ \vec{T}_\varphi &= \frac{\partial \vec{x}}{\partial \varphi} = \rho(-\sin \varphi \vec{e}_1 + \cos \varphi \vec{e}_2), & g_\varphi &= |\vec{T}_\varphi| = \rho. \end{aligned} \quad (3.7.43)$$

Wir sehen, dass die Polarkoordinaten im Ursprung singulär sind, weil dort $T_\varphi = 0$ ist. Die entsprechenden orthogonalen Einheitsvektoren sind demnach

$$\begin{aligned}\vec{e}'_\rho &= \frac{1}{g_\rho} \vec{T}_\rho = \cos \varphi \vec{e}_1 + \sin \varphi \vec{e}_2, \\ \vec{e}'_\varphi &= \frac{1}{g_\varphi} \vec{T}_\varphi = -\sin \varphi \vec{e}_1 + \cos \varphi \vec{e}_2.\end{aligned}\tag{3.7.44}$$

Etwas schwieriger ist die umgekehrte Transformation von kartesischen Komponenten $\underline{x} = (x_1, x_2)^T$ zu den Polarkoordinaten (ρ, φ) . Für ρ gilt, wie oben angegeben

$$\rho = |\vec{x}| = \sqrt{x_1^2 + x_2^2}.\tag{3.7.45}$$

Für den Winkel φ , der nur für $\vec{x} \neq \vec{0}$ wohldefiniert ist, müssen wir simultan

$$\cos \varphi = \frac{x_1}{\rho}, \quad \sin \varphi = \frac{x_2}{\rho}\tag{3.7.46}$$

erfüllen. Anschaulich ist klar, dass die simultane Lösung dieser Gleichungen eindeutig ist, wenn man zusätzlich als Wertebereich für den Polarwinkel $\varphi \in (-\pi, \pi]$ wählt. Die erste Gleichung ist offenbar nicht eindeutig zu lösen, denn aus ihr folgt, dass es zwei Lösungen

$$\varphi_\pm = \pm \arccos\left(\frac{x_1}{\rho}\right)\tag{3.7.47}$$

gibt. Dabei verstehen wir unter dem arccos den üblicherweise auch auf Taschenrechnern oder in den gängigen Programmiersprachen realisierten **Hauptwert**. Für alle $\xi \in [-1, 1]$ ist $\alpha = \arccos \xi \in [0, \pi]$. Für diesen Wertebereich gilt nun aber $\sin \alpha \geq 0$. Wegen $\sin^2 \alpha + \cos^2 \alpha = 1$ folgt daraus

$$\sin \alpha = +\sqrt{1 - \cos^2 \alpha} = \sqrt{1 - \xi^2}.\tag{3.7.48}$$

Setzen wir also (3.7.47) ein, erhalten wir für beide Vorzeichen

$$\sin(\varphi_\pm) = \pm \sin\left[\arccos\left(\frac{x_1}{\rho}\right)\right] = \pm \sqrt{1 - \frac{x_1^2}{\rho^2}} = \pm \sqrt{\frac{\rho^2 - x_1^2}{\rho^2}} = \pm \sqrt{\frac{x_2^2}{\rho^2}} = \pm \frac{|x_2|}{\rho}.\tag{3.7.49}$$

Damit also auch die zweite Gleichung in (3.7.46) erfüllt ist, müssen wir für $x_2 > 0$ in (3.7.47) das obere und für $x_2 < 0$ das untere Vorzeichen wählen. Für $x_2 = 0$ wählen wir definitionsgemäß stets das obere Vorzeichen, denn dann ist $x_1/\rho = x_1/|x_1|$, und man erhält dann für $x_1 > 0$ stets $\arccos(x_1/\rho) = \arccos 1 = 0$ und für $x_1 < 0$ stets $\arccos(x_1/\rho) = \arccos(-1) = \pi$. Endgültig gilt also

$$\varphi = \begin{cases} \arccos(x_1/\rho) & \text{für } x_2 \geq 0, \\ -\arccos(x_1/\rho) & \text{für } x_2 < 0. \end{cases}\tag{3.7.50}$$

In der Literatur findet man oft auch Formeln, die φ mit Hilfe des arctan bestimmen. Das hat den Vorteil, dass man zur Bestimmung von φ nicht erst ρ berechnen muss. Allerdings benötigt man mehr Fallunterscheidungen bzgl. der Lage von $\vec{x}$ in den unterschiedlichen Quadranten des kartesischen Koordinatensystems.

Dazu bemerken wir, dass für $\xi \in \mathbb{R}$ definitionsgemäß $\arctan \xi \in (-\pi/2, \pi/2)$ gilt. Aus der obigen Skizze lesen wir ab, dass für Ortsvektoren in der rechten offenen Halbebene, also $x_1 > 0$ stets

$$\varphi = \arctan\left(\frac{x_2}{x_1}\right) \quad \text{für } x_1 > 0\tag{3.7.51}$$

3.8. Potentialfelder

gilt. Falls $x_1 < 0$ und $x_2 > 0$, ist $\arctan(x_2/x_1) \in (-\pi/2, 0)$ und definitionsgemäß $\varphi \in (\pi/2, \pi)$. In diesem Fall müssen wir also

$$\varphi = \pi + \arctan\left(\frac{x_2}{x_1}\right) \quad \text{für } x_1 < 0, \quad x_2 > 0 \quad (3.7.52)$$

setzen. Für $x_1 < 0$ und $x_2 < 0$ ist $\arctan(x_2/x_1) \in (0, \pi/2)$ und $\varphi \in (-\pi, -\pi/2)$. Demnach ist dann

$$\varphi = -\pi + \arctan\left(\frac{x_2}{x_1}\right) \quad \text{für } x_1 < 0, \quad x_2 < 0. \quad (3.7.53)$$

Bleibt schließlich der Fall $x_1 = 0$. Dann ist $\varphi = \pi/2$ für $x_2 > 0$ und $\varphi = -\pi/2$ für $x_2 < 0$ zu setzen. Die Polarkoordinaten sind für $x_1 = x_2 = 0$ unbestimmt. Es gilt zwar sicher $\rho = 0$, aber der Winkel φ ist beliebig. Obwohl also hier geometrisch betrachtet keinerlei Besonderheit vorliegt, sind die Polarkoordinaten im Ursprung singulär. Man spricht hier daher von einer **Koordinatensingularität**.

3.8 Potentialfelder

Für die Physik ist es eine sehr wichtige Fragestellung, in welchen Fällen wir ein vorgegebenes Vektorfeld $\vec{V}(\vec{x})$ als **Gradient eines Skalarfeldes** schreiben können. In der Mechanik ist dies insbesondere für **Kraftfelder** interessant. In (3.3.33) haben wir z.B. das Kraftfeld für ein Partikelchen (Planet), das sich im Schwerfeld eines anderen im Ursprung des Koordinatensystems fixierten Punkt (Sonne) bewegt, angegeben. Die Newtonsche Gravitationskraft lautet demnach

$$\vec{F}_G(\vec{x}) = -\frac{\gamma m M}{r^2} \frac{\vec{x}}{r} \quad \text{mit } r = |\vec{x}|. \quad (3.8.1)$$

Wir können nun fragen, ob diese Kraft der Gradient eines Skalarfeldes ist. Aus Gründen, die gegen Ende dieses Abschnitts klar werden, definiert man das Potential so, dass die Kraft durch den negativen Gradienten gegeben ist:

$$\vec{F}_G(\vec{x}) = -\vec{\nabla} \Phi_G(\vec{x}). \quad (3.8.2)$$

Aufgrund der Symmetrie der physikalischen Situation bzgl. beliebiger Drehungen um den Ursprung, liegt es nahe anzunehmen, dass Φ_G nur von r abhängt, denn dies ist die einzige skalare Größe, die wir aus dem Ortsvektor $\vec{x}$ bilden können. Wir können nun den Gradienten leicht berechnen. Verwenden wir Kugelkoordinaten folgt sofort

$$\vec{\nabla} \Phi_G(r) = \vec{e}_r \Phi'_G(r) = \frac{\vec{x}}{r} \Phi'_G(r) \stackrel{!}{=} -\vec{F}_G(\vec{x}) = \frac{\vec{x}}{r} \frac{\gamma m M}{r^2}. \quad (3.8.3)$$

Diese Gleichung wird erfüllt, wenn

$$\Phi'_G(r) = \frac{\gamma m M}{r^2} \Rightarrow \Phi_G(r) = -\frac{\gamma m M}{r} \quad (3.8.4)$$

ist. Dabei haben wir eine eventuelle additive Integrationskonstante weggelassen, weil diese für die Kraft $\vec{F}_G = -\vec{\nabla} \Phi_G$ keine Rolle spielt.

Wir bemerken noch, warum die Existenz eines Potentials für die Kraft wichtig ist. Nehmen wir also an, die Kraft auf ein Partikelchen sei ein Vektorfeld, das ein Potential besitzt. Die Newtonsche Bewegungsgleichung lautet dann

$$m \dot{\vec{v}} = \vec{F}(\vec{x}) = -\vec{\nabla} \Phi(\vec{x}), \quad \vec{v} = \dot{\vec{x}}. \quad (3.8.5)$$

Multiplizieren wir diese Gleichung skalar mit $\vec{v}$, folgt daraus

$$m \vec{v} \cdot \dot{\vec{v}} = -\vec{v} \cdot \vec{\nabla} \Phi(\vec{x}). \quad (3.8.6)$$

3. Vektoranalysis

Nun sind beide Seiten der Gleichung totale Zeitableitungen, wie man aus Produkt- und Kettenregel sofort sieht (*Nachrechnen!*):

$$\frac{d}{dt} \left(\frac{m}{2} \vec{v}^2 \right) = - \frac{d}{dt} \Phi[\vec{x}(t)]. \quad (3.8.7)$$

Bringen wir die rechte auf die linke Seite der Gleichung folgt

$$\frac{d}{dt} \left(\frac{m}{2} \vec{v}^2 + \Phi(\vec{x}) \right) = 0 \Rightarrow E = \frac{m}{2} \vec{v}^2 + \Phi(\vec{x}) = \text{const.} \quad (3.8.8)$$

Die Größe E ist die **Energie** des Teilchens. Sie ist die Summe aus der **kinetischen Energie** $E_{\text{kin}} = T = m\vec{v}^2/2$ und der **potentiellen Energie** $E_{\text{pot}} = \Phi(\vec{x})$, und (3.8.8) besagt, dass entlang der tatsächlichen Trajektorie des Teilchens, also einer beliebigen Lösung der Bewegungsgleichung (3.8.5), die Energie stets erhalten bleibt. Wenn also das Kraftfeld ein Potential besitzt, gilt der **Energieerhaltungssatz**. Man nennt daher Vektorfelder, die ein Potential besitzen, **Potentialfelder** oder auch **konservative Vektorfelder** (von lat. conservare = erhalten).

Das obige Beispiel der Gravitationskraft zeigt, dass es physikalisch wichtige Kraftfelder gibt, die tatsächlich ein Potential besitzen. Freilich ist die eben vorgeführte Methode in Fällen, die nicht eine so hohe Symmetrie aufweisen, zu kompliziert, um zu untersuchen, ob ein Kraftfeld ein **Potentialfeld** ist, d.h. ob es als (negativer) Gradient eines Skalarfeldes darstellbar ist. Im Rest dieses Abschnitts befassen wir uns mit der Frage nach Kriterien dafür, dass ein vorgegebenes Vektorfeld $\vec{V}(\vec{x})$ ein solches **Potentialfeld** ist, ohne dass wir notwendig explizit das Potential ausrechnen müssen. Weiter wollen wir eine kompakte Formel für die Berechnung des Potentials finden, falls ein solches existiert.

Um eine **notwendige Bedingung** für die Existenz eines Potentials zu finden, müssen wir uns zunächst mit den **zweiten partiellen Ableitungen** eines vorgegebenen Skalarfeldes befassen. Es ist klar, dass diese wie bei Funktionen einer unabhängiger Veränderlicher einfach als Hintereinanderausführung zweier Ableitungen definieren lassen, also

$$\partial_j \partial_k \tilde{\Phi}(\underline{x}) = \frac{\partial}{\partial x_j} \left[\frac{\partial}{\partial x_k} \tilde{\Phi}(\underline{x}) \right]. \quad (3.8.9)$$

Hier haben wir *zuerst* nach x_k und *dann* nach x_j differenziert. Die Frage ist nun, ob diese doppelte partielle Ableitung von dieser Reihenfolge abhängt.

Die Antwort ist, dass partielle Ableitungen **kommutieren**, wenn die Funktion **zweimal stetig differenzierbar** ist.

Der Beweis erfordert ein wenig mehr Grundlagen aus der Analysis. Der Einfachheit halber betrachten wir eine Funktion $f: \mathbb{R}^2 \rightarrow \mathbb{R}$. Wir nehmen an, sie sei zweimal partiell stetig differenzierbar. Wir wollen zeigen, dass dann

$$\partial_1 \partial_2 f(x_1, x_2) = \partial_2 \partial_1 f(x_1, x_2) \quad (3.8.10)$$

gilt. Betrachten wir nun für beliebige Δx_1 und Δx_2 die Funktion

$$F(x_1, x_2) = f(x_1 + \Delta x_1, x_2 + \Delta x_2) - f(x_1, x_2). \quad (3.8.11)$$

Weil diese Funktion voraussetzungsgemäß stetig differenzierbar ist, gibt es nach dem Mittelwertsatz der Differentialrechnung (vgl. Abschnitt 1.7.4) eine Zahl ξ_1 zwischen x_1 und $x_1 + \Delta x_1$, so dass

$$F(x_1, x_2) = \Delta x_1 \partial_1 f(\xi_1, x_2) = [\partial_1 f(\xi_1, x_2 + \Delta x_2) - \partial_1 f(\xi_1, x_2)] \Delta x_2. \quad (3.8.12)$$

Da nach Voraussetzung auch die ersten Ableitungen stetig partiell differenzierbar sind, können wir nochmals den Mittelwertsatz der Differentialrechnung anwenden, d.h. es gibt eine Zahl η_1 zwischen x_2 und $x_2 + \Delta x_2$, so dass

$$F(x_1, x_2) = \partial_2 \partial_1 f(\xi_1, \eta_1) \Delta x_1 \Delta x_2. \quad (3.8.13)$$

Dieses Argument können wir aber auch genauso in umgekehrter Reihenfolge anwenden, d.h. den Mittelwertsatz zuerst auf die Variable x_2 anwenden. Demnach existiert eine Zahl η_2 zwischen x_2 und $x_2 + \Delta x_2$, so dass

$$F(x_1, x_2) = \partial_2 F(x_1, \eta_2) \Delta x_2 = [\partial_2 f(x_1 + \Delta x_1, \eta_2) - \partial_2 f(x_1, \eta_2)] \Delta x_2. \quad (3.8.14)$$

Der Mittelwertsatz auf $\partial_2 f$ angewandt liefert dann die Existenz einer Zahl ξ_2 zwischen x_1 und $x_1 + \Delta x_1$, so dass

$$F(x_1, x_2) = \partial_1 \partial_2 f(\xi_2, \eta_2) \Delta x_1 \Delta x_2. \quad (3.8.15)$$

Da dies für alle Δx_1 und Δx_2 gilt, folgt, dass

$$\partial_1 \partial_2 f(\xi_1, \eta_1) = \partial_2 \partial_1 f(\xi_2, \eta_2) \quad (3.8.16)$$

gilt. Lassen wir nun $(\Delta x_1, \Delta x_2) \rightarrow (0, 0)$ streben, gilt $(\xi_1, \eta_1) \rightarrow (x_1, x_2)$ und $(\xi_2, \eta_2) \rightarrow (x_1, x_2)$. Da voraussetzungsgemäß die zweiten partiellen Ableitungen stetig sind, folgt damit aus (3.8.16), dass tatsächlich (3.8.10) gilt.

Jetzt können wir ein **notwendiges Kriterium** für die Existenz eines Potentials für ein vorgegebenes Vektorfeld angeben. Sei also $\vec{V}$ ein Vektorfeld, dessen Komponenten $\underline{V}$ bzgl. einer kartesischen Basis stetig partiell differenzierbar sind. Angenommen, es sei ein Potentialfeld. Dann gilt

$$\underline{V}(\underline{x}) = -\underline{\nabla} \tilde{\Phi}(\underline{x}) \quad \text{bzw.} \quad V_j(\underline{x}) = -\partial_j \tilde{\Phi}(\underline{x}). \quad (3.8.17)$$

Nach Voraussetzung sind die Komponenten V_j stetig partiell differenzierbar, und damit $\tilde{\Phi}$ zweimal stetig partiell differenzierbar. Nach dem eben bewiesenen Satz folgt

$$\partial_k \partial_j \tilde{\Phi} = \partial_j \partial_k \tilde{\Phi} \Rightarrow \partial_k V_j = \partial_j V_k. \quad (3.8.18)$$

Dies gilt für alle $j, k \in \{1, 2, 3\}$. Für $k = j$ ist die Gleichung identisch erfüllt und trägt folglich keine Einschränkung an die Komponenten des Vektorfeldes für die Existenz eines Potentials bei.

Für $j \neq k$ impliziert aber (3.8.18)

$$\text{rot } \underline{V} = \underline{\nabla} \times \underline{V} = \begin{pmatrix} \partial_2 V_3 - \partial_3 V_2 \\ \partial_3 V_1 - \partial_1 V_3 \\ \partial_1 V_2 - \partial_2 V_1 \end{pmatrix} = \underline{0}. \quad (3.8.19)$$

Damit ein Vektorfeld ein Potential besitzt, muss also die Rotation in jedem regulären Punkt, also dort, wo es stetig partiell differenzierbar ist, verschwinden.

Man rechnet leicht nach (*Übung!*), dass für das Gravitationskraftfeld (3.8.1) in der Tat außer im Ursprung des Koordinatensystems, wo das Feld eine Singularität besitzt, $\text{rot } \vec{F}_G = \vec{0}$ gilt.

3.9 Wegintegrale und Potentialfelder

Als nächstes definieren wir **Wegintegrale** von Vektorfeldern. Es sei $\vec{V}$ ein stetiges Vektorfeld und $\mathcal{C} : \vec{x} : [t_1, t_2] \rightarrow E^3$ eine Kurve, die ganz im Inneren eines Bereiches verläuft, in dem $\vec{V}$ definiert und stetig ist.

Dann definieren wir das **Wegintegral entlang des Weges** $\mathcal{C}$ durch

$$\int_{\mathcal{C}} d\vec{x} \cdot \vec{V}(\vec{x}) = \int_{t_1}^{t_2} dt \, \dot{\vec{x}}(t) \cdot \vec{V}[\vec{x}(t)]. \quad (3.9.1)$$

Es ist klar, dass diese Definition nur vom Weg selbst abhängt und nicht von der Parametrisierung, denn angenommen, es ist $\vec{x}' : [\lambda_1, \lambda_2] \rightarrow E^3$ eine andere Parametrisierung desselben Weges, so gilt

$$\vec{x}'(\lambda) = \vec{x}[t(\lambda)] \rightarrow \frac{d}{d\lambda} \vec{x}'(\lambda) = \frac{d}{d\lambda} t(\lambda) \dot{\vec{x}}[t(\lambda)]. \quad (3.9.2)$$

Nach der Substitutionsformel für Integrale folgt daraus in der Tat, dass

$$\int_{\lambda_1}^{\lambda_2} d\lambda \frac{d}{d\lambda} \vec{x}'(\lambda) \cdot \vec{V}[\vec{x}'(\lambda)] = \int_{t_1}^{t_2} dt \frac{d}{dt} \lambda(t) \frac{d}{d\lambda} t(\lambda) \dot{\vec{x}}(t) \cdot \vec{V}[\vec{x}(t)] = \int_{t_1}^{t_2} dt \, \dot{\vec{x}}(t) \cdot \vec{V}[\vec{x}(t)]. \quad (3.9.3)$$

Dabei haben wir im letzten Schritt den Satz von der Ableitung der Umkehrfunktion verwendet, demzufolge

$$\frac{d}{dt} \lambda(t) = \frac{1}{\frac{d}{d\lambda} t(\lambda)} \quad (3.9.4)$$

gilt. Außerdem ist offensichtlich das Wegintegral (3.9.1) ein **Skalar**, d.h. es kann mit Hilfe der Komponenten von $\vec{V}$ bzgl. einer beliebigen kartesischen Basis berechnet werden:

$$\int_{\mathcal{C}} d\vec{x} \cdot \vec{V}(\vec{x}) = \int_{t_1}^{t_2} dt \, \dot{\underline{x}} \cdot \underline{V}[\underline{x}(t)] = \int_{t_1}^{t_2} dt \sum_{j=1}^3 \dot{x}_j V_j[\underline{x}(t)]. \quad (3.9.5)$$

Nehmen wir nun an, $\vec{V}$ sei ein Potentialfeld, d.h. es existiert ein Skalarfeld Φ , so dass $\vec{V} = -\vec{\nabla}\Phi$ ist. Berechnen wir nun das Wegintegral entlang eines beliebigen Weges, so folgt

$$\int_{\mathcal{C}} d\vec{x} \cdot \vec{V}(\vec{x}) = - \int_{\mathcal{C}} d\vec{x} \cdot \vec{\nabla}\Phi(\vec{x}) = - \int_{t_1}^{t_2} dt \, \dot{\vec{x}}(t) \cdot \vec{\nabla}\Phi[\vec{x}(t)]. \quad (3.9.6)$$

Nach der Kettenregel (3.4.12) folgt nun

$$\dot{\vec{x}}(t) \cdot \vec{\nabla}\Phi[\vec{x}(t)] = \dot{\underline{x}}(t) \cdot \vec{\nabla}\Phi[\underline{x}(t)] = \frac{d}{dt} \Phi[\underline{x}(t)] = \frac{d}{dt} \Phi[\vec{x}(t)]. \quad (3.9.7)$$

Da $\vec{V}$ voraussetzungsgemäß stetig ist, sind die partiellen Ableitungen von Φ stetig und damit auch (3.9.7) als Funktion des Kurvenparameters t . Demnach können wir auf (3.9.6) den **Hauptsatz der Differential- und Integralrechnung** anwenden. Dies ergibt

$$\int_{\mathcal{C}} d\vec{x} \cdot \vec{V}(\vec{x}) = - \int_{t_1}^{t_2} dt \frac{d}{dt} \Phi[\vec{x}(t)] = - [\Phi[\vec{x}(t_2)] - \Phi[\vec{x}(t_1)]]. \quad (3.9.8)$$

Dies zeigt, dass für ein Potentialfeld das Wegintegral nur von **Anfangs- und Endpunkt** des Integrationsweges abhängt, nicht von dem Integrationsweg selbst. Dabei ist freilich darauf zu achten, dass dies nur für Wege gilt, die im Definitionsbereich des Vektorfeldes liegen und wo dieses stetig ist.

Aus (3.9.8) folgt auch, wie wir ggf. das Potential berechnen können. Dazu sei $G \subseteq E^3$ ein beliebiges offenes Gebiet⁸.

Existiert dann ein Punkt $\vec{x}_0 \in G$, so dass zu jedem $\vec{x} \in G$ ein differenzierbarer Weg $\mathcal{C}(\vec{x})$ mit Anfangspunkt $\vec{x}_0$ und Endpunkt $\vec{x}$ existiert, so ist gemäß (3.9.8)

$$\Phi(\vec{x}) = \tilde{\Phi}(\underline{x}) = - \int_{\mathcal{C}(\vec{x})} d\vec{x}' \cdot \vec{V}(\vec{x}') \quad (3.9.9)$$

ein Potential von $\vec{V}$.

Um dies zu zeigen, machen wir von der angenommenen **Wegunabhängigkeit** der Wegintegrale (3.9.9) Gebrauch. Dann ist nämlich die Funktion (3.9.9) unabhängig von der konkreten Wahl der Wege $\mathcal{C}(\vec{x})$. Da G offen ist, existiert um $\vec{x} \in G$ eine Kugel $K_\epsilon(\vec{x})$, so dass $K_\epsilon(\vec{x}) \subseteq G$. Berechnen wir für das Potential in der Darstellung des Skalarfeldes $\tilde{\Phi}$ bzgl. einer beliebigen Orthonormalbasis (3.9.9) die partielle Ableitung $\partial_1 \tilde{\Phi}$, indem wir die Definition als Limes anwenden. Dazu sei $|\Delta x_1| < \epsilon$. Dann liegt $\vec{x} + \Delta x_1 \vec{e}_1 \in G$, da konstruktionsgemäß die Kugel $K_\epsilon(\vec{x})$ ganz in G liegt. Damit folgt aus (3.9.9)

$$\tilde{\Phi}(x_1 + \Delta x_1, x_2, x_3) - \tilde{\Phi}(x_1, x_2, x_3) = - \left[\int_{\mathcal{C}(\vec{x} + \Delta x_1 \vec{e}_1)} d\vec{x}' \cdot \vec{V}(\vec{x}') - \int_{\mathcal{C}(\vec{x})} d\vec{x}' \cdot \vec{V}(\vec{x}') \right]. \quad (3.9.10)$$

Offenbar ist nun der Ausdruck in der eckigen Klammer das Wegintegral entlang eines in G gelegenen Weges, der von $\vec{x}$ zu $\vec{x} + \Delta x_1 \vec{e}_1$ verläuft. Da voraussetzungsgemäß das Wegintegral unabhängig von der konkreten Form des Weges ist, können wir dieses Integral durch das Integral entlang der geraden Strecke (parallel zum Basisvektor $\vec{e}_1$ ersetzen. Dieser Weg kann aber durch

$$s: \vec{x}'(t) = \vec{x} + t \Delta x_1 \vec{e}_1, \quad t \in [0, 1] \quad (3.9.11)$$

parametrisiert werden. Damit folgt

$$\begin{aligned} \tilde{\Phi}(x_1 + \Delta x_1, x_2, x_3) - \tilde{\Phi}(x_1, x_2, x_3) &= - \int_0^1 dt \Delta x_1 \vec{e}_1 \cdot \vec{V}(\vec{x} + t \Delta x_1 \vec{e}_1) \\ &= - \Delta x_1 \int_0^1 dt V_1(x_1 + t \Delta x_1, x_2, x_3). \end{aligned} \quad (3.9.12)$$

Dividieren wir diese Gleichung durch Δx_1 , so folgt aus der Stetigkeit von V_1 für den Limes $\Delta x_1 \rightarrow 0$

$$\partial_1 \tilde{\Phi}(\underline{x}) = - \int_0^1 dt V_1(x_1, x_2, x_3) = -V_1(\underline{x}). \quad (3.9.13)$$

Genauso können wir auch für die beiden anderen partiellen Ableitungen zeigen, dass $\partial_2 \tilde{\Phi}(\underline{x}) = -V_2(\underline{x})$ und $\partial_3 \tilde{\Phi}(\underline{x}) = -V_3(\underline{x})$ gilt. Es ist also tatsächlich

$$\underline{V}(\underline{x}) = -\underline{\nabla} \tilde{\Phi}(\underline{x}) \Rightarrow \vec{V}(\vec{x}) = -\vec{\nabla} \Phi(\vec{x}). \quad (3.9.14)$$

Wir kommen auf die Frage, in welchen Fällen das Verschwinden der Rotation eines Vektorfeldes auch hinreichend für die Existenz eines Potentials für dieses Vektorfeldes in Abschnitt 3.11 noch zurück. Dazu benötigen wir nämlich Flächenintegrale und den Stokesschen Integralsatz.

⁸Ein offenes Gebiet $G \subseteq E^3$ ist dabei eine Teilmenge von E^3 , für die zu jedem Punkt $\vec{x} \in E^3$ eine Kugel $K_\epsilon(\vec{x})$ mit Mittelpunkt bei $\vec{x}$ und Radius $\epsilon > 0$ existiert, so dass $K_\epsilon(\vec{x}) \subseteq G$ gilt.

3.10 Flächenintegrale und der Stokessche Integralsatz

In diesem Abschnitt führen wir **Flächenintegrale** über Vektorfeldern ein und besprechen den **Stokesschen Integralsatz**.

3.10.1 Orientierte Flächen im Raum

Analog zu den in Abschnitt 3.1 eingeführten Kurven im Raum führen wir nun **zweidimensionale Flächen** mit Hilfe ihrer **Parameterdarstellung** ein.

Es sei dazu $G \subseteq \mathbb{R}^2$ ein Bereich im $\mathbb{R}^2$. Dann beschreibt eine Abbildung $\vec{x} : G \rightarrow E^3, (q_1, q_2) \mapsto \vec{x}(q_1, q_2)$ eine **Fläche** S im Raum. Wir gehen im Folgenden davon aus, dass diese Abbildung stetig partiell differenzierbar ist. Man bezeichnet dann die Fläche als **glatt**.

Ein wichtiges *Beispiel* ist eine Kugelschale mit dem Radius a um den Ursprung eines rechtshändigen kartesischen Koordinatensystems, die wir durch **sphärische Koordinaten** $\vartheta \in (0, \pi)$ und $\varphi \in [0, 2\pi)$ parametrisieren. Diese Parametrisierung ergibt sich, indem man bei den Kugelkoordinaten $r = R = \text{const}$ setzt, d.h. wegen (3.7.1)

$$\underline{x}(\vartheta, \varphi) = R \begin{pmatrix} \cos \varphi \sin \vartheta \\ \sin \varphi \sin \vartheta \\ \cos \vartheta \end{pmatrix}. \quad (3.10.1)$$

Mit der Parametrisierung der Fläche mittels generalisierter Koordinaten (q_1, q_2) sind durch die partiellen Ableitungen nun auch in jedem Punkt der zwei **Tangentenvektoren** an die Fläche

$$\vec{T}_j(q_1, q_2) = \frac{\partial}{\partial q_j} \vec{x}(q_1, q_2), \quad j \in \{1, 2\} \quad (3.10.2)$$

definiert. Die Parametrisierung der Fläche durch diese generalisierten Koordinaten heißt dann **regulär**, wenn diese beiden Tangentenvektoren **linear unabhängig** sind.

Für unsere oben definierte Kugelschale gilt für die Komponenten der Tangentenvektoren (s. die obige Skizze)

$$\underline{T}_\vartheta(\varphi, \vartheta) = \frac{\partial}{\partial \vartheta} \underline{x}(\vartheta, \varphi) = a \begin{pmatrix} \cos \varphi \cos \vartheta \\ \sin \varphi \cos \vartheta \\ -\sin \vartheta \end{pmatrix}, \quad \underline{T}_\varphi(\vartheta, \varphi) = a \frac{\partial}{\partial \varphi} \underline{x}(\vartheta, \varphi) = R \begin{pmatrix} -\sin \varphi \sin \vartheta \\ \cos \varphi \sin \vartheta \\ 0 \end{pmatrix}. \quad (3.10.3)$$

Diese Tangentenvektoren sind für die oben angegebenen Bereiche für ϑ und φ linear unabhängig. Da für $\vartheta = 0$ oder $\vartheta = \pi$ der Tangentenvektor $\underline{T}_\vartheta = \underline{0}$ wird, sind diese Punkte der Parametrisierung singulär. Dies reflektiert wieder die oben besprochene Koordinatensingularität der sphärischen Koordinaten.

In jedem regulären Punkt der Parametrisierung der Fläche können wir mit Hilfe des Kreuzproduktes den **Flächennormalenvektor** auf die Fläche an diesem Punkt definieren:

$$\vec{N}(q_1, q_2) = \vec{T}_1(q_1, q_2) \times \vec{T}_2(q_1, q_2). \quad (3.10.4)$$

Die Richtung des Normalenvektors hängt dabei von der Reihenfolge der generalisierten Koordinaten ab. Vertauschen wir q_1 und q_2 , kehrt sich wegen der Antisymmetrie des Vektorproduktes der Normalenvektor um. Durch die Wahl der Reihenfolge der beiden generalisierten Koordinaten prägen wir der Fläche also eine

Orientierung auf, indem wir eine der beiden möglichen Orientierungen der Flächennormalenvektoren festlegen. Diese Wahl hängt von der Anwendung ab und ist mathematisch willkürlich. Wir kommen darauf im Folgenden noch zurück.

Für die Kugelschale gilt

$$\underline{N} = \underline{T}_\vartheta \times \underline{T}_\varphi = a^2 \sin \vartheta \begin{pmatrix} \cos \varphi \sin \vartheta \\ \sin \varphi \sin \vartheta \\ \cos \vartheta \end{pmatrix} \quad (3.10.5)$$

Durch die Wahl der Reihenfolge $q_1 = \vartheta$ und $q_2 = \varphi$ ist die Orientierung der Kugelfläche also so, dass die Normalenvektoren radial nach außen weisen⁹.

3.10.2 Definition des Flächenintegrals

Wichtig für die Definition der Flächenintegrale ist die geometrische Bedeutung dieses Normalenvektors: Die beiden Tangentenvektoren (3.10.2) spannen die **Tangentialebene an die Fläche** an dem jeweiligen Punkt auf, und die „infinitesimalen Vektoren“ $dq_1 \vec{T}_1$ und $dq_2 \vec{T}_2$ definieren ein infinitesimales Parallelogramm.

Demnach ist der **Flächenelementvektor**

$$d^2 \vec{f} = dq_1 dq_2 \vec{T}_1 \times \vec{T}_2 = dq_1 dq_2 \vec{N} = d^2 q \vec{N} \quad (3.10.6)$$

ein auf der Tangentialebene senkrecht stehender Vektor, und seine Länge entspricht dem Flächeninhalt des infinitesimalen Parallelogramms.

Das **Flächenintegral** eines Vektorfeldes $\vec{V}$, das in einem offenen Gebiet definiert und (stückweise) stetig ist, ist durch

$$\int_S d^2 \vec{f} \cdot \vec{V} = \int_G d^2 q \vec{N}(q_1, q_2) \cdot \vec{V}[\vec{x}(q_1, q_2)] \quad (3.10.7)$$

definiert. Dabei ist $G \subseteq \mathbb{R}^2$ der Definitionsbereich der generalisierten Koordinaten, die die gesamte Fläche (ggf. bis auf einzelne Punkte) parametrisieren. Das Vorzeichen des Integrals hängt von der oben besprochen Wahl der *Orientierung der Fläche* ab. Das Symbol S steht also stets für eine orientierte Fläche.

Natürlich können wir auch den **Flächeninhalt** der Fläche ausrechnen:

$$A = \int_S |d^2 \vec{f}| = \int_G d^2 q |\vec{N}(q_1, q_2)|. \quad (3.10.8)$$

Wir können z.B. leicht die Oberfläche einer Kugel ausrechnen. Dazu benötigen wir nur den oben berechneten Normalenvektor (3.10.5). Offenbar ist nämlich $|\vec{N}| = a^2 \sin \vartheta$ (man beachte, dass wegen $\vartheta \in (0, \pi)$ stets $\sin \vartheta \geq 0$ ist). Damit folgt das aus der Elementargeometrie bekannte Resultat für die Kugeloberfläche:

$$\begin{aligned} A &= \int_0^\pi d\vartheta \int_0^{2\pi} d\varphi a^2 \sin \vartheta = 2\pi a^2 \int_0^\pi d\vartheta \sin \vartheta \\ &= 2\pi a^2 [-\cos \vartheta]_0^\pi = 4\pi a^2. \end{aligned} \quad (3.10.9)$$

⁹Dies entspricht einer Standardorientierung für geschlossene Flächen, die später im Zusammenhang mit den Volumenintegralen und dem Gaußschen Integralsatz noch wichtig werden wird: Betrachtet man ein Volumen V mit der dann notwendig geschlossenen Oberfläche $S = \partial V$ als Rand, orientiert man diese Oberfläche so, dass die Normalenvektoren nach außen, d.h. von dem Volumen weg, zeigen.

3. Vektoranalysis

Ein wichtiges *Beispiel* für die Integration eines Vektorfeldes ist die Integration des Gravitationsfeldes einer Punktmasse über eine Kugel. Der Einfachheit halber lassen wir die Vorfaktoren weg und integrieren

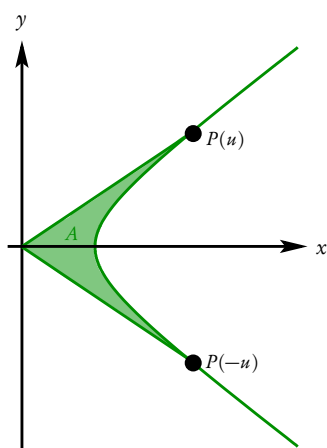
$$\vec{V}(\vec{x}) = \frac{\vec{x}}{r^3}, \quad r = |\vec{x}|. \quad (3.10.10)$$

Für die Werte entlang der Kugelfläche erhalten wir

$$\underline{V}[\underline{x}(\vartheta, \varphi)] = \frac{1}{a^2} \begin{pmatrix} \cos \varphi \sin \vartheta \\ \sin \varphi \sin \vartheta \\ \cos \vartheta \end{pmatrix} \quad (3.10.11)$$

und damit (*nachrechnen!*)

$$\int_S d^2 \vec{f} \cdot \vec{V} = \int_0^\pi d\vartheta \int_0^{2\pi} d\varphi \vec{N}(\vartheta, \varphi) \cdot \vec{V}[\vec{x}(\vartheta, \varphi)] = \int_0^\pi d\vartheta \int_0^{2\pi} d\varphi \sin \vartheta = 4\pi. \quad (3.10.12)$$

Beispiel: Fläche bei der Einheitshyperbel

Als einfaches Beispiel für eine Flächenberechnung betrachten wir die durch

$$\underline{x}(u) = \begin{pmatrix} \cosh u \\ \sinh u \\ 0 \end{pmatrix} \quad (3.10.13)$$

parametrisierte Einheitshyperbel in der x_1 - x_2 -Ebene. Die geometrische Bedeutung von u finden wir durch Berechnung der nebenstehend gezeichneten Fläche A . Wir können die Fläche offenbar mit den Parametern $\lambda \in [0, 1]$ und $u' \in [-u, u]$ durch

$$\vec{x}(\lambda, u') = \begin{pmatrix} \lambda \cosh u' \\ \lambda \sinh u' \\ 0 \end{pmatrix} \quad (3.10.14)$$

parametrisieren. Die Flächennormalenvektoren sind dann (*Nachrechnen!*)

$$\begin{aligned} d^2 \vec{f} &= \partial_\lambda \vec{x} \times \partial_{u'} \vec{x} \\ &= d\lambda du' \begin{pmatrix} \cosh u' \\ \sinh u' \\ 0 \end{pmatrix} \times \begin{pmatrix} \sinh u' \\ \cosh u' \\ 0 \end{pmatrix} \\ &= d\lambda du' \begin{pmatrix} 0 \\ 0 \\ \cosh^2 u' - \sinh^2 u' \end{pmatrix} \\ &= d\lambda du' \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}. \end{aligned} \quad (3.10.15)$$

Damit wird die besagte Fläche

$$A = \int_0^1 d\lambda \int_{-u}^u du' |d^2 \vec{f}| = \int_0^1 d\lambda \int_{-u}^u du' \lambda = u. \quad (3.10.16)$$

d.h. u ergibt die Fläche des entsprechenden Segments einer Einheitshyperbel. Man kann sich klar machen, dass man dies auch ähnlich am Einheitskreis $\underline{x}(\varphi) = (\cos \varphi, \sin \varphi, 0)$, denn dort ist die Fläche des durch die Punkte $\underline{x}(\pm\varphi)$ bestimmten Kreissegments durch $A/A_{\text{Kreis}} = 2\varphi/(2\pi) = \varphi/\pi$ gegeben. Nun ist aber $A_{\text{Kreis}} = \pi$ und damit $A = \varphi$.

3.10.3 Unabhängigkeit des Flächenintegrals von der Parametrisierung

In (3.10.6) haben wir das Flächenintegral auf die einfache Integration über die beiden generalisierten Koordinaten der Fläche zurückgeführt und die Vektoren durch ihre Komponenten bzgl. einer beliebigen rechtsorientierten kartesischen Basis ausgedrückt. Dass das Integral von der Wahl des rechtshändigen Orthonormalsystems unabhängig ist, rechnet man in analoger Weise wie oben bei den Wegintegralen nach und sei dem Leser zur *Übung* überlassen. Man muss nur beachten, dass die auftretenden Operationen wie Skalar- und Kreuzprodukte allesamt unabhängig von der Wahl der rechtshändigen Orthonormalbasis sind.

Nicht offensichtlich ist, dass das Flächenintegral auch unabhängig von der konkreten Wahl der Parametrisierung ist. Es sei also $\vec{x}' : G' \rightarrow E^3, (q'_1, q'_2) \mapsto \vec{x}'(q'_1, q'_2)$ eine beliebige andere Parametrisierung derselben Fläche.

3. Vektoranalysis

Dies impliziert, dass wir die generalisierten Koordinaten (q_1, q_2) der ursprünglichen Parametrisierung als umkehrbar eindeutige Funktionen der neuen generalisierten Koordinaten (q'_1, q'_2) betrachten können.

Um zu zeigen, dass das Flächenintegral unabhängig von der Parametrisierung ist, untersuchen wir zuerst die Transformation der Flächennormalenvektoren zwischen den beiden Parametrisierungen und zeigen dann, dass die Definition des Flächenintegrals mit beiden Parametrisierungen tatsächlich dasselbe Resultat liefert. Für die Tangentenvektoren gilt

$$\vec{T}'_j = \frac{\partial \vec{x}'}{\partial q'_j} = \sum_{k=1}^2 \frac{\partial \vec{x}}{\partial q_k} \frac{\partial q_k}{\partial q'_j} = \sum_{k=1}^2 \vec{T}_k \frac{\partial q_k}{\partial q'_j}, \quad (3.10.17)$$

und daraus folgt für die Normalenvektoren

$$\vec{N}' = \vec{T}'_1 \times \vec{T}'_2 = \sum_{k,l=1}^2 \vec{T}_k \times \vec{T}_l \frac{\partial q_k}{\partial q'_1} \frac{\partial q_l}{\partial q'_2}. \quad (3.10.18)$$

In dieser Summe tragen nun nur die Terme bei, für die entweder $k = 1, l = 2$ oder $k = 2, l = 1$ ist. Schreiben wir also die Summe ausführlich hin, folgt

$$\vec{N}' = \vec{T}_1 \times \vec{T}_2 \frac{\partial q_1}{\partial q'_1} \frac{\partial q_2}{\partial q'_2} + \vec{T}_2 \times \vec{T}_1 \frac{\partial q_2}{\partial q'_1} \frac{\partial q_1}{\partial q'_2}. \quad (3.10.19)$$

Nun ist aber $\vec{N} = \vec{T}_1 \times \vec{T}_2 = -\vec{T}_2 \times \vec{T}_1$. Damit folgt

$$\vec{N}' = \left(\frac{\partial q_1}{\partial q'_1} \frac{\partial q_2}{\partial q'_2} - \frac{\partial q_2}{\partial q'_1} \frac{\partial q_1}{\partial q'_2} \right) \vec{N}. \quad (3.10.20)$$

Den Vorfaktor können wir nun wie folgt übersichtlicher schreiben.

Dazu führen wir die **Jacobi-Matrix** der umkehrbar eindeutigen Transformation zwischen den beiden Sätzen von generalisierten Koordinaten $(q_1, q_2) \leftrightarrow (q'_1, q'_2)$ ein

$$\frac{\partial(q_1, q_2)}{\partial(q'_1, q'_2)} = \begin{pmatrix} \partial q_1 / \partial q'_1 & \partial q_2 / \partial q'_1 \\ \partial q_1 / \partial q'_2 & \partial q_2 / \partial q'_2 \end{pmatrix} \quad (3.10.21)$$

Die Klammer in (3.10.20) ist offenbar die Determinante dieser Matrix, die **Jacobi-Determinante**:

$$\vec{N}' = \det \frac{\partial(q_1, q_2)}{\partial(q'_1, q'_2)} \vec{N}. \quad (3.10.22)$$

Es gilt also

$$\int_{G'} d^2 q' \vec{N}' \cdot \vec{V}[\vec{x}'(q'_1, q'_2)] = \int_{G'} d^2 q' \det \frac{\partial(q_1, q_2)}{\partial(q'_1, q'_2)} \vec{N} \cdot \vec{V}[\vec{x}'(q'_1, q'_2)]. \quad (3.10.23)$$

Da für jedes $(q'_1, q'_2) \in G'$ stets $\vec{x}'(q'_1, q'_2) = \vec{x}(q_1, q_2)$ für die entsprechenden generalisierten Koordinaten $(q_1, q_2) \in G$ gilt, müssen wir zeigen, dass für beliebige Funktionen $f : G \rightarrow \mathbb{R}$

$$\int_{G'} d^2 q' \det \frac{\partial(q_1, q_2)}{\partial(q'_1, q'_2)} f[q(q')] = \int_G d^2 q f(q) \quad (3.10.24)$$

gilt. Dies ist die Verallgemeinerung der Substitutionsformel einfacher Integrale auf Doppelintegrale. Diese Formel ist aber einfach zu verstehen. Dazu müssen wir nur bedenken, dass das Gebiet $G' \subseteq \mathbb{R}^2$ umkehrbar eindeutig auf das Gebiet $G \subseteq \mathbb{R}^2$ abgebildet wird, und zwar durch unsere Koordinatentransformation

$q : G' \rightarrow G$, $q' \mapsto q(q')$. Durch die Koordinatenlinien $q'_1 = \text{const}$ bzw. $q'_2 = \text{const}$ wird das Gebiet G in infinitesimale Parallelogramme zerlegt, und diese besitzen die Flächen

$$dA = d^2 q' \left| \det \frac{\partial(q_1, q_2)}{\partial(q'_1, q'_2)} \right|. \quad (3.10.25)$$

Durch die Koordinatenlinien q_1 und q_2 wird das Gebiet G in infinitesimale Rechtecke zerlegt (wenn man (q_1, q_2) als kartesische Koordinaten in der Ebene auffasst). Daraus folgt aber sofort, dass für stetige Funktionen in der Tat (3.10.24) gilt.

Vorausgesetzt, dass die Jacobi-Determinante der Transformation $(q_1, q_2) \leftrightarrow (q'_1, q'_2)$ *positiv* ist, ist demnach durch die Unabhängigkeit des Flächenintegrals von der Wahl der Parametrisierung bewiesen, denn dann folgt aus (3.10.24)

$$\int_{G'} d^2 q' \vec{N}' \cdot \vec{V}[\vec{x}'(q'_1, q'_2)] = \int_G d^2 q \vec{N} \cdot \vec{V}[\vec{x}(q_1, q_2)] = \int_S d\vec{f} \cdot \vec{V}. \quad (3.10.26)$$

Wir müssen also für unsere Transformation voraussetzen, dass die Jacobi-Determinante positiv ist. Man bezeichnet solche Transformationen zwischen generalisierten Koordinaten einer Fläche als **orientierungserhaltende Transformationen**. Falls irgendeine gewählte Transformation eine negative Jacobi-Determinante ergibt, müssen wir nur die beiden neuen generalisierten Koordinaten umordnen.

3.10.4 Koordinatenunabhängige Definition der Rotation

In Abschnitt 3.6 hatten wir die Rotation eines Vektorfeldes über seine kartesischen Komponenten definiert. Wir können nun die Rotation aber auch koordinatenunabhängig als Grenzwert eines Wegintegrals einführen. Dazu sei $\vec{V} : G \rightarrow E^3$ ein auf einem (offenen) Gebiet $G \subseteq E^3$ definiertes Vektorfeld mit partiell stetig differenzierbaren kartesischen Komponenten. Zu einem Punkt $\vec{x}$ existiert dann eine Umgebung U (z.B. eine Kugel oder ein Quader), die vollständig in G liegt.

Wir betrachten nun eine Schar orientierter Flächen $\Delta S \subseteq U$ mit dem dazu konsistent nach der Rechte-Hand-Regel orientierten Rand $\partial \Delta S$ im Limes $\Delta S \rightarrow 0$, was bedeuten soll, dass die Fläche auf den Punkt $\vec{x}$ zusammengezogen wird.

Der Flächennormalenvektor sei im Limes $\vec{N} \rightarrow \Delta A \vec{n}$, wobei ΔA der Flächeninhalt des infinitesimalen Flächenstücks sein soll und damit $\vec{n}$ der **Flächennormaleneinheitsvektor**. Dann ist

$$\vec{n} \cdot \text{rot } \vec{V}(\vec{x}) = \lim_{\Delta S \rightarrow 0} \frac{1}{\Delta A} \int_{\partial \Delta S} d\vec{r} \cdot \vec{V}(\vec{r}). \quad (3.10.27)$$

Führen wir diese Prozedur nun mit drei Flächen mit Einheitsflächenvektoren $\vec{n}_1$, $\vec{n}_2$ und $\vec{n}_3$, die ein rechtshändiges kartesischen Koordinatensystem (auch lokal im Punkt $\vec{x}$, was weiter unten im Zusammenhang mit **krummlinigen Orthogonalkoordinaten** noch wichtig wird) aufspannen, haben wir so die Rotation als Vektorfeld definiert.

Um zu zeigen, dass dies mit unserer obigen Definition in Abschnitt 3.6 übereinstimmt, wählen wir als Flächenelemente kleine Rechtecke parallel zu den Koordinatenachsen eines rechtshändigen kartesischen Koordinatensystems und berechnen das entsprechende Wegintegral gemäß (3.10.27). Betrachten wir ein Quadrat parallel zur 12-Ebene, das wir im Gegenuhrzeigersinn durchlaufen. Eine bequeme Parametrisierung für die vier Seiten ist dann

$$\underline{r}_1(t) = \begin{pmatrix} x_1 + t \\ x_2 - \epsilon/2 \\ x_3 \end{pmatrix}, \quad \underline{r}_2(t) = \begin{pmatrix} x_1 + \epsilon/2 \\ x_2 + t \\ x_3 \end{pmatrix}, \quad \underline{r}_3(t) = \begin{pmatrix} x_1 - t \\ x_2 + \epsilon/2 \\ x_3 \end{pmatrix}, \quad \underline{r}_4(t) = \begin{pmatrix} x_1 - \epsilon/2 \\ x_2 - t \\ x_3 \end{pmatrix}, \quad (3.10.28)$$

3. Vektoranalysis

wobei jeweils $t \in (-\epsilon/2, \epsilon/2)$ durchläuft und $\epsilon > 0$ so klein ist, dass das umschlossene Quadrat ganz in G liegt. Dessen Einheitsnormalenvektor ist an jedem Punkt offenbar $\vec{e}_3$. Betrachten wir nun das Wegintegral entlang des ersten Wegstücks:

$$\begin{aligned} \int_{\mathcal{C}_1} d\vec{r} \cdot \vec{V}(\vec{r}) &= \int_{-\epsilon/2}^{\epsilon/2} dt \dot{\vec{r}}_1(t) \cdot \underline{V}[(x_1 + t, x_2 - \epsilon/2, x_3)] \\ &= \int_{-\epsilon/2}^{\epsilon/2} dt V_1(x_1 + t, x_2 - \epsilon/2, x_3) \\ &= V_1(\xi_1, x_2 - \epsilon/2, x_3)\epsilon. \end{aligned} \quad (3.10.29)$$

Dabei haben wir im letzten Schritt den Mittelwertsatz der Integralrechnung angewendet. Dabei ist $\xi_1 \in (x_1 - \epsilon/2, x_1 + \epsilon/2)$. Genauso folgt

$$\int_{\mathcal{C}_3} d\vec{r} \cdot \vec{V}(\vec{r}) = -V_1(\xi'_1, x_2 + \epsilon/2, x_3)\epsilon. \quad (3.10.30)$$

Da voraussetzungsgemäß die Komponenten des Vektorfeldes $\vec{V}$ stetig partiell differenzierbar sind, gibt es aufgrund des Zwischenwertsatzes der Differentialrechnung ein $\xi_2 \in (x_2 - \epsilon/2, x_2 + \epsilon/2)$ mit

$$\begin{aligned} \int_{\mathcal{C}_1} d\vec{r} \cdot \vec{V}(\vec{r}) + \int_{\mathcal{C}_3} d\vec{r} \cdot \vec{V}(\vec{r}) &= -[V_1(\xi'_1, x_2 + \epsilon/2, x_3) - V_1(\xi_1, x_2 - \epsilon/2, x_3)]\epsilon \\ &= -\epsilon^2 \partial_2 V_1(\xi''_1, \xi_2, x_3). \end{aligned} \quad (3.10.31)$$

Dabei ist $\xi''_1 \in (x_1 - \epsilon/2, x_1 + \epsilon/2)$. Lassen wir nun $\Delta S \rightarrow 0$, also $\epsilon \rightarrow 0$, gehen folgt wegen $\Delta A = \epsilon^2$

$$\lim_{\epsilon \rightarrow 0} \frac{1}{\epsilon^2} \left[\int_{\mathcal{C}_1} d\vec{r} \cdot \vec{V}(\vec{r}) + \int_{\mathcal{C}_3} d\vec{r} \cdot \vec{V}(\vec{r}) \right] = -\partial_2 V_1(\underline{x}). \quad (3.10.32)$$

Analog zeigt man, dass

$$\lim_{\epsilon \rightarrow 0} \frac{1}{\epsilon^2} \left[\int_{\mathcal{C}_2} d\vec{r} \cdot \vec{V}(\vec{r}) + \int_{\mathcal{C}_4} d\vec{r} \cdot \vec{V}(\vec{r}) \right] = +\partial_1 V_2(\underline{x}). \quad (3.10.33)$$

Gemäß der Definition (3.10.27) ist demnach

$$\underline{e}_3 \cdot \text{rot } \underline{V}(\underline{x}) = \partial_1 V_2(\underline{x}) - \partial_2 V_1(\underline{x}), \quad (3.10.34)$$

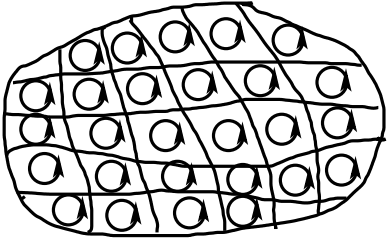
und das stimmt mit der Definition in (3.6.6) überein. Die übrigen Komponenten berechnen sich analog für Quadrate parallel zur 23- und 13-Ebene (*Übung!*).

3.10.5 Der Integralsatz von Stokes

Mit der Definition der Rotation über Wegintegrale im vorigen Abschnitt wird der **Integralsatz von Stokes** fast zu einer Selbstverständlichkeit.

Für ein stetig partiell differenzierbares Vektorfeld gilt demnach für jede orientierte Fläche S mit gemäß der Rechte-Hand-Regel kompatibel mit der Flächenorientierung orientiertem Rand ∂S

$$\int_S d^2 \vec{f} \cdot \text{rot } \vec{V} = \int_{\partial S} d\vec{x} \cdot \vec{V}. \quad (3.10.35)$$



Um dies zu zeigen, muss man nur entsprechend der nebenstehenden Skizze die Fläche in viele kleine Flächenstücke unterteilen und auf jeder dieser Flächenstücke den Mittelwertsatz der Integralrechnung sowie die Definition der Rotation aus dem vorigen Abschnitt anwenden:

$$\int_{\Delta S_j} d^2 \vec{f} \cdot \text{rot } \vec{V} = \Delta A_j \vec{n} \cdot \text{rot } \vec{V}(\vec{\xi}) = \int_{\partial \Delta S_j} d\vec{x} \cdot \vec{V}. \quad (3.10.36)$$

Bei der Summe über alle Flächenstücke heben sich die Wegintegrale über die inneren Linien weg, weil diese jeweils zweimal in entgegengesetztem Sinne durchlaufen werden, und es bleibt nur das Wegintegral über den Rand der Gesamtfläche übrig. Im Limes beliebig feiner Unterteilung der Flächenstücke ergibt sich aus (3.10.46) für die linke Seite wieder das Flächenintegral und die rechte Seite ist stets das Wegintegral über ∂S , womit der **Integralsatz von Stokes** bewiesen ist.

Man kann den Integralsatz von Stokes für den Fall, dass das Gebiet der Parametrisierung ein Rechteck in der (q_1, q_2) -Ebene ist, auch noch auf den Hauptsatz der Differential- und Integralrechnung zurückführen. Sei also $G = [q_{1L}, q_{1R}] \times [q_{2L}, q_{2R}]$. Dann ist

$$\int_S d^2 \vec{f} \cdot \text{rot } \vec{V}(\vec{x}) = \int_{q_{1L}}^{q_{1R}} dq_1 \int_{q_{2L}}^{q_{2R}} dq_2 (\vec{T}_1 \times \vec{T}_2) \cdot (\vec{\nabla} \times \vec{V}), \quad (3.10.37)$$

Dabei sind

$$\vec{T}_1 = \partial_{q_1} \vec{x}, \quad \vec{T}_2 = \partial_{q_2} \vec{x} \quad (3.10.38)$$

die Tangentenvektoren an die Koordinatenlinien der Parametrisierung der Fläche S . Den Integranden können wir nun mit der „bac-cab-Formel“ umformen

$$\begin{aligned} (\vec{T}_1 \times \vec{T}_2) \cdot (\vec{\nabla} \times \vec{V}) &= [(\vec{T}_1 \times \vec{T}_2) \times \vec{\nabla}] \cdot \vec{V} \\ &= \vec{T}_2 \cdot [(\vec{T}_1 \cdot \vec{\nabla}) \vec{V}] - \vec{T}_1 \cdot [(\vec{T}_2 \cdot \vec{\nabla}) \vec{V}]. \end{aligned} \quad (3.10.39)$$

Nun ist aber wegen der Kettenregel

$$(\vec{T}_1 \cdot \vec{\nabla}) \vec{V} = [(\partial_{q_1} \vec{x}) \cdot \vec{\nabla}] \vec{V} = \partial_{q_1} \vec{V}, \quad (3.10.40)$$

$$(\vec{T}_2 \cdot \vec{\nabla}) \vec{V} = [(\partial_{q_2} \vec{x}) \cdot \vec{\nabla}] \vec{V} = \partial_{q_2} \vec{V}. \quad (3.10.41)$$

Damit erhalten wir in (3.10.37)

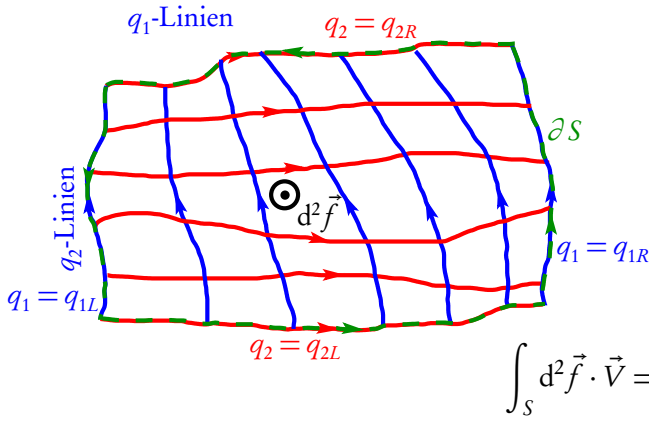
$$\begin{aligned} \int_S d^2 \vec{f} \cdot \text{rot } \vec{V} &= \int_{q_{1L}}^{q_{1R}} dq_1 \int_{q_{2L}}^{q_{2R}} dq_2 [\vec{T}_2 \cdot \partial_{q_1} \vec{V} - \vec{T}_1 \cdot \partial_{q_2} \vec{V}] \\ &= \int_{q_{1L}}^{q_{1R}} dq_1 \int_{q_{2L}}^{q_{2R}} dq_2 [\partial_{q_1} (\vec{T}_2 \cdot \vec{V}) - \vec{V} \cdot \partial_{q_1} \vec{T}_2 - \partial_{q_2} (\vec{T}_1 \cdot \vec{V}) + \vec{V} \cdot \partial_{q_2} \vec{T}_1]. \end{aligned} \quad (3.10.42)$$

Nun ist aber

$$\partial_{q_1} \vec{T}_2 = \partial_{q_1} \partial_{q_2} \vec{x} = \partial_{q_2} \partial_{q_1} \vec{x} = \partial_{q_2} \vec{T}_1. \quad (3.10.43)$$

Damit heben sich im Integranden in (3.10.42) der zweite und der vierte Beitrag gegenseitig weg, und wir finden

$$\begin{aligned} \int_S d^2 \vec{f} \cdot \text{rot } \vec{V} &= \int_{q_{1L}}^{q_{1R}} dq_1 \int_{q_{2L}}^{q_{2R}} dq_2 [\partial_{q_1} (\vec{T}_2 \cdot \vec{V}) - \partial_{q_2} (\vec{T}_1 \cdot \vec{V})] \\ &= \int_{q_{2L}}^{q_{2R}} dq_2 (\vec{T}_2 \cdot \vec{V})_{q_1=q_{1R}} - \int_{q_{1L}}^{q_{1R}} dq_1 (\vec{T}_1 \cdot \vec{V})_{q_2=q_{2R}} \\ &= \int_{q_{2L}}^{q_{2R}} dq_2 [(\vec{T}_2 \cdot \vec{V})_{q_1=q_{1R}} - (\vec{T}_2 \cdot \vec{V})_{q_1=q_{1L}}] \\ &\quad - \int_{q_{1L}}^{q_{1R}} dq_1 [(\vec{T}_1 \cdot \vec{V})_{q_2=q_{2R}} - (\vec{T}_1 \cdot \vec{V})_{q_2=q_{2L}}]. \end{aligned} \quad (3.10.44)$$



Aus der nebenstehenden Skizze entnehmen wir, dass bei Beachtung der durch die Vorzeichen der jeweiligen Integrale bestimmte Orientierung diese vier Integrale zusammengenommen das Wegintegral über die geschlossene Randkurve ∂S der Fläche S ergeben. Dabei ist der Umlaufsinn relativ zur Orientierung der Flächennormalenvektoren $d^2\vec{f}$ gemäß der Rechten-Hand-Regel festgelegt. Damit erhalten wir schließlich wieder den Stokesschen Integralsatz

$$\int_S d^2\vec{f} \cdot \vec{V} = \int_{\partial S} d\vec{x} \cdot \vec{V}. \quad (3.10.45)$$

3.10.6 Der Greensche Satz in der Ebene

Wir können nun einen Integralsatz für **ebene Vektorfelder** als Spezialfall des Stokesschen Satzes herleiten. Wir wählen dazu die Ebene als 12-Ebene eines kartesischen Koordinatensystems im Raum. Dann ist für ein ebenes Vektorfeld

$$\underline{V}(\underline{x}) = \begin{pmatrix} V_1(x_1, x_2) \\ V_2(x_1, x_2) \\ 0 \end{pmatrix}. \quad (3.10.46)$$

Wir betrachten nun ein beliebiges offenes Gebiet S in der 12-Ebene mit im Gegenuhrzeigersinn orientierten Rand ∂S . Dieses Gebiet können wir aber genauso gut als Fläche im dreidimensionalen Raum ansehen. Die Flächennormaleneinheitsvektoren sind dann allesamt $\vec{n} = \vec{e}_3$.

Nun gilt gemäß (3.6.6)

$$\text{rot } \underline{V} = \begin{pmatrix} 0 \\ 0 \\ \partial_1 V_2 - \partial_2 V_1 \end{pmatrix}. \quad (3.10.47)$$

Dann spezialisiert sich der Stokessche Integralsatz auf

$$\int_S d^2\vec{f} \cdot \text{rot } \vec{V} = \int_S df(\partial_1 V_2 - \partial_2 V_1) = \int_{\partial S} d\vec{x} \cdot \vec{V}. \quad (3.10.48)$$

Das ist der **Integralsatz von Green in der Ebene**:

$$\int_S d^2f(\partial_1 V_2 - \partial_2 V_1) = \int_{\partial S} d\vec{x} \cdot \vec{V}. \quad (3.10.49)$$

Dabei muss man die Konvention beachten, dass der Rand des ebenen Gebiets S so zu durchlaufen ist, dass in der Durchlaufrichtung betrachtet dieses Gebiet stets links liegt.

3.11 Das Poincaré-Lemma

Nun können wir die in Abschnitt 3.8 aufgeworfene Frage beantworten, unter welchen Umständen aus $\text{rot } \vec{V} = 0$ folgt, dass $\vec{V}$ ein Potentialfeld ist. In Abschnitt 3.9 haben wir gesehen, dass ein Potential existiert, wenn Wegintegrale in einem Gebiet unabhängig von der konkreten Form des Weges sind und nur von Anfangs- und Endpunkt der Wege abhängen. Sind nun $\mathcal{C}_1$ und $\mathcal{C}_2$, die dieselben Punkte $\vec{x}_1$ und $\vec{x}_2$ zu Anfangs- und Endpunkt haben, so ist der Weg $\mathcal{C} = \mathcal{C}_1 - \mathcal{C}_2$ geschlossen. Dabei bezeichnen wir mit $-\mathcal{C}_2$ den in umgekehrter

Richtung durchlaufenen Weg $\mathcal{C}_2$ und mit $\mathcal{C}_1 - \mathcal{C}_2$ entsprechend den Weg, der zuerst entlang $\mathcal{C}_1$ von $\vec{x}_1$ zu $\vec{x}_2$ und dann entlang $-\mathcal{C}_2$ von $\vec{x}_2$ zurück zu $\vec{x}_1$ führt.

Damit es also für ein stetig partiell differenzierbares Vektorfeld ein Potential gibt, ist es notwendig und hinreichend, dass für alle **geschlossenen Wege**

$$\int_{\mathcal{C}} d\vec{r} \cdot \vec{V} = 0 \quad (3.11.1)$$

gilt.

Ist nun also $\text{rot } \vec{V} = 0$ in einem Gebiet, das so geartet ist, dass man zu jedem geschlossenen Weg $\mathcal{C}$ eine Fläche S finden kann, so dass sein Rand $\partial S = \mathcal{C}$ ist, folgt sofort aus dem Stokesschen Satz

$$0 = \int_S d^2 \vec{f} \cdot \text{rot } \vec{V} = \int_{\partial S} d\vec{x} \cdot \vec{V} = \int_{\mathcal{C}} d\vec{x} \cdot \vec{V}. \quad (3.11.2)$$

Neben der Bedingung, dass die Rotation verschwindet, muss also auch noch das Gebiet G , in dem $\vec{V}$ stetig partiell differenzierbar sein muss die besagte Eigenschaft besitzen, dass jede geschlossene Kurve die Randkurve einer ganz in G gelegenen Fläche ist.

Diese Bedingung können wir auch so formulieren: Es muss möglich sein, eine jede geschlossene Kurve innerhalb von G stetig zu einem Punkt in G zu deformieren, denn bei dieser Deformation überstreicht die so definierte Schar von Kurven eine Fläche mit der ursprünglichen Kurve als Rand. Man nennt solche Gebiete **einfach zusammenhängend**.

I.a. wird ein Vektorfeld in einem **mehrfach zusammenhängenden Gebiet**, für das $\text{rot } \vec{V} = 0$ gilt nur **lokal** ein Potential besitzen. Ist nämlich G offen, kann man um jeden Punkt eine ganz in G gelegene Kugelumgebung finden. Eine Kugel ist nun einfach zusammenhängend, und dort existiert dann auch ein Potential, und dieses ist dort gemäß Abschnitt 3.9 bis auf eine Konstante eindeutig bestimmt. I.a., wird es in mehrfach zusammenhängenden Gebieten nur lokale Potentiale geben, und diese sind i.a. nicht mehr eindeutig bestimmt. Ein schönes physikalisches Beispiel, das dies recht drastisch veranschaulicht, ist der **Potentialwirbel**. Das ist das Geschwindigkeitsfeld einer Strömung, das in kartesischen Koordinaten durch

$$\underline{V}(\vec{x}) = \frac{v_0}{x_1^2 + x_2^2} \begin{pmatrix} -x_2 \\ x_1 \\ 0 \end{pmatrix} \quad (3.11.3)$$

gegeben ist. Es ist offensichtlich in $\underline{x} \in G = \mathbb{R}^3 \setminus \mathbb{R}(0,0,1)^T$ definiert (also überall außer entlang der 3-Achse) und stetig partiell differenzierbar. Außerdem gilt dort auch überall $\text{rot } \vec{V} = 0$, wie man mit Hilfe der Formel (3.6.6) sofort bestätigt (*Übung!*). Dieses Gebiet ist aber offensichtlich *nicht einfach zusammenhängend*, denn man kann jede Kurve, die die 3-Achse umschließt, innerhalb von G nicht stetig zu einem Punkt zusammenziehen. Entsprechend schneidet jede Fläche mit einer solchen Kurve als Rand die 3-Achse und liegt folglich nicht ganz in G .

Betrachten wir nun als Kurve einen Kreis K_R um die 3-Achse in der 12-Ebene des kartesischen Koordinatensystems mit Radius a , den wir durch

$$\underline{x}(\varphi) = a \begin{pmatrix} \cos \varphi \\ \sin \varphi \\ 0 \end{pmatrix}, \quad \varphi \in [-\pi, \pi] \quad (3.11.4)$$

parametrisieren. Es folgt

$$\frac{d}{d\varphi} \underline{x}(\varphi) = a \begin{pmatrix} -\sin \varphi \\ \cos \varphi \\ 0 \end{pmatrix}, \quad \underline{V}[\underline{x}(\varphi)] = \frac{v_0}{a} \begin{pmatrix} -\sin \varphi \\ \cos \varphi \\ 0 \end{pmatrix} \quad (3.11.5)$$

und damit für das Wegintegral

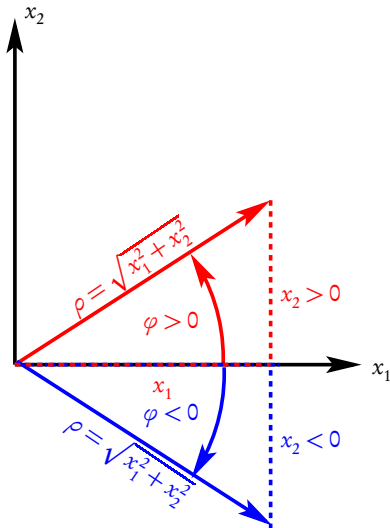
$$\int_{K_R} d\vec{x} \cdot \underline{V}(\vec{x}) = \int_{-\pi}^{\pi} d\varphi \dot{\underline{x}}(\varphi) \cdot \underline{V}[\underline{x}(\varphi)] = v_0 \int_{-\pi}^{\pi} d\varphi = 2\pi v_0. \quad (3.11.6)$$

Obwohl also überall in G stets $\vec{\nabla} = 0$ gilt, verschwindet das Wegintegral entlang der geschlossenen Kreislinie nicht, und damit existiert kein in G eindeutig bestimmtes Potential.

Wir können andererseits aber in jedem **einfach zusammenhängenden Teilgebiet** ein solches Potential finden. Das größtmögliche solche Teilgebiet erhalten wir offenbar, wenn wir eine Halbebene mit der 3-Achse als Rand ausnehmen. Wir wählen willkürlich den Teil der 13-Ebene mit $x_1 \leq 0$, die wir mit $H_{12}^{\leq}$ bezeichnen wollen. Es gibt nun in $\tilde{G} = E^3 \setminus H_{12}^{x_1 \leq 0}$ keine geschlossenen stetigen Kurven, die die 3-Achse umlaufen, da diese unweigerlich unsere ausgeschlossene Halbebene schneiden müssten, und folglich ist dieses Gebiet einfach zusammenhängend und entsprechend verschwinden die Wegintegrale von $\vec{V}$ entlang geschlossener Wege in $\tilde{G}$. Beliebige Wege, die einen festgehaltenen Punkt $\vec{x}_0 \in \tilde{G}$ mit einem beliebigen anderen Punkt $\vec{x} \in \tilde{G}$ verbinden, ergeben dann denselben Wert für das Wegintegral, und dieses Integral definiert dann ein Potential für $\vec{V}$ in dem auf $\tilde{G}$ eingeschränkten Gebiet.

Um es zu berechnen, können wir irgendeinen Punkt $\vec{x}_0$ und irgendeinen Weg, der ihn innerhalb von $\tilde{G}$ mit $\vec{x}$ verbindet, verwenden, um das Potential zu bestimmen. Wir wählen als Anfangspunkt $\underline{x}_0 = (1, 0, 0)^T$. Um einen bequemen Verbindungsweg definieren zu können, verwenden wir **Zylinderkoordinaten** (R, φ, z) gemäß (3.7.7) und wählen als Definitionsbereiche $R > 0$, $\varphi \in (-\pi, \pi)$ und $z \in \mathbb{R}$:

$$\underline{x} = \begin{pmatrix} R \cos \varphi \\ R \sin \varphi \\ z \end{pmatrix} \quad (3.11.7)$$



Um nun $\vec{x}_0$ mit einem Weg mit $\vec{x}(R, \varphi, z)$ zu verbinden, der ein möglichst einfach berechenbares Wegintegral ergibt, beachten wir, dass gerade Linien entlang der 1-Achse und in z -Richtung keinen Beitrag zum Wegintegral liefern, denn beide Wegarten besitzen zu $\vec{V}$ überall senkrechte Tangentenvektoren. Weiter wissen wir von unserer obigen Berechnung des Wegintegrals entlang einer Kreislinie in einer Ebene parallel zur 12-Ebene, dass Integrale entlang von solchen Kreislinienstücken trivial werden. Wir wählen also den folgenden Weg: Von $\underline{x}_0 = (1, 0, 0)$ gehen wir zunächst entlang der 1-Achse zum Punkt $(R, 0, 0)$, dann entlang des Kreislinienstücks

$$\underline{x}(\varphi') = \begin{pmatrix} R \cos \varphi' \\ R \sin \varphi' \\ 0 \end{pmatrix}, \quad \varphi' \in [0, \varphi]. \quad (3.11.8)$$

Schließlich verbinden wir noch $(R \cos \varphi, R \sin \varphi, 0)$ mit $\underline{x} = (R \cos \varphi, R \sin \varphi, z)$ durch die entsprechende zur 3-Achse parallele Strecke. Wie oben gesagt, tragen das erste und das letzte gerade Teilstück dieses Weges $\mathcal{C}$ nichts zum Wegintegral von $\vec{V}$ bei, und wir finden schließlich

$$\Phi(\vec{x}) = - \int_{\mathcal{C}} d\vec{r} \cdot \vec{V}(\vec{r}) = - \int_0^\varphi d\varphi' \dot{\underline{x}}(\varphi') \cdot \underline{V}[\underline{x}(\varphi')] = -v_0 \varphi. \quad (3.11.9)$$

Den Winkel φ erhalten wir aufgrund der nebenstehenden Skizze aus

$$\varphi = \operatorname{sign} x_2 \arccos \left(\frac{x_1}{\sqrt{x_1^2 + x_2^2}} \right). \quad (3.11.10)$$

Man muss dazu bemerken, dass dann $\varphi \in (-\pi, \pi)$ zu liegen kommt, wenn man Punkte auf der negativen 1-Achse ausschließt. Das entspricht genau unserer Wahl des Gebiets $\tilde{G}$. Für Punkte auf der positiven 1-Achse ist eindeutig $\varphi = 0$, so dass die Unbestimmtheit der Signum-Funktion

$$\operatorname{sign} x = \begin{cases} -1 & \text{für } x < 0, \\ +1 & \text{für } x > 0 \end{cases} \quad (3.11.11)$$

hierbei keine Rolle spielt.

Wir können nun leicht verifizieren, dass (3.11.9) mit (3.11.10) für φ tatsächlich ein Potential des Potentialwirbelfeldes in $\tilde{G}$ ist. Es gilt nämlich nach der Kettenregel

$$\partial_1 \Phi = -\operatorname{sign} x_2 \partial_1 \left(\frac{x_1}{\sqrt{x_1^2 + x_2^2}} \right) \arccos' \left(\frac{x_1}{\sqrt{x_1^2 + x_2^2}} \right). \quad (3.11.12)$$

Die Ableitung des arccos ist aber

$$\arccos' x = -\frac{1}{\sqrt{1-x^2}} \quad (3.11.13)$$

und damit

$$\begin{aligned} \partial_1 \Phi &= v_0 \operatorname{sign} x_2 \frac{\sqrt{x_1^2 + x_2^2} - x_1^2 / \sqrt{x_1^2 + x_2^2}}{x_1^2 + x_2^2} \frac{1}{\sqrt{1 - x_1^2 / (x_1^2 + x_2^2)}} = v_0 \frac{x_2}{x_1^2 + x_2^2}, \\ \partial_2 \Phi &= -v_0 \operatorname{sign} x_2 \frac{x_1 x_2}{(x_1^2 + x_2^2)^{3/2}} \frac{1}{\sqrt{1 - x_1^2 / (x_1^2 + x_2^2)}} = -v_0 \frac{x_1}{x_1^2 + x_2^2}, \\ \partial_3 \Phi &= 0. \end{aligned} \quad (3.11.14)$$

Daraus ergibt sich in der Tat

$$\underline{V} = -\underline{\nabla} \Phi = \frac{v_0}{x_1^2 + x_2^2} \begin{pmatrix} -x_2 \\ x_1 \\ 0 \end{pmatrix}. \quad (3.11.15)$$

Wählt man irgendeine andere Halbebene, findet man ein anderes Potential für das entsprechend geänderte einfach zusammenhängende Gebiet $\tilde{G}'$, das sich von dem soeben berechneten in Bereichen in $\tilde{G} \cap \tilde{G}'$ um eine Konstante unterscheidet und entsprechend andere Wertebereiche für φ verwendet, so dass das Potential bei der entsprechenden Halbebene einen Sprung um $\pm 2\pi v_0$ aufweist. Unser Potential ist entlang der negativen x_1 -Achse $-v_0\pi$, wenn man sich von negativen x_2 -Werten her nähert und $+v_0\pi$, wenn man sich von positiven x_2 -Werten her nähert.

3.12 Volumenintegrale, Divergenz und Gaußscher Integralsatz

In diesem Abschnitt definieren wir Volumenintegrale über Skalarfelder, geben die koordinatenunabhängige Definition der Divergenz an und beweisen den Gaußschen Integralsatz.

3.12.1 Definition des Volumenintegrals

Das **Volumenintegral** über ein Skalarfeld ist die Integration über einen dreidimensionalen Bereich $V \subseteq E^3$. Man kann dieses Gebiet durch eine beliebige Parametrisierung mit drei generalisierten Koordinaten (q_1, q_2, q_3) beschreiben.

Das Volumenelement ist dabei unabhängig von der Parametrisierung durch die entsprechende **Jacobi-Determinante** gegeben

$$d^3x = dq_1 dq_2 dq_3 \det \frac{\partial(x_1, x_2, x_3)}{\partial(q_1, q_2, q_3)} = d^3q (\vec{T}_1 \times \vec{T}_2) \cdot \vec{T}_3. \quad (3.12.1)$$

Dabei wählen wir die Reihenfolge der generalisierten Koordinaten so, dass die Jacobi-Determinante positiv ist, d.h. die drei Koordinatenlinien in jedem regulären Punkt der Parametrisierung liefern drei linear unabhängige Tangentenvektoren

$$\vec{T}_j = \frac{\partial \vec{x}}{\partial q_j}, \quad (3.12.2)$$

die relativ zum rechtshändigen kartesischen Koordinatensystem gleichorientiert sind und also auch eine rechtshändige Basis bilden.

Analog wie bei den Flächenintegralen zeigt man (*Übung!*), dass dann das Volumenintegral über das Skalarfeld Φ

$$\int_V d^3x \Phi(\vec{x}) = \int_{\tilde{G}} d^3q \det \frac{\partial(x_1, x_2, x_3)}{\partial(q_1, q_2, q_3)} \tilde{\Phi}[\underline{x}(q)]. \quad (3.12.3)$$

parametrisierungsunabhängig ist. Dabei ist $\tilde{G} \subseteq \mathbb{R}^3$ ein Gebiet, das den Parameterbereich für die generalisierten Koordinaten $q = (q_1, q_2, q_3)$ angibt.

Für krummlinige Orthogonalkoordinaten folgt aus (3.7.13-3.7.16)

$$\det \frac{\partial(x_1, x_2, x_3)}{\partial(q_1, q_2, q_3)} = (\underline{T}_1 \times \underline{T}_2) \cdot \underline{T}_3 = g_1 g_2 g_3. \quad (3.12.4)$$

Für Kugelkoordinaten ergibt sich damit gemäß (3.7.4)

$$\det \frac{\partial(x_1, x_2, x_3)}{\partial(r, \vartheta, \varphi)} = (\underline{T}_r \times \underline{T}_\vartheta) \cdot \underline{T}_\varphi = r^2 \sin \vartheta. \quad (3.12.5)$$

Wie wir sehen, verschwindet die Jacobi-Determinante für $r = 0$ bzw. $\vartheta = 0$ oder $\vartheta = \pi$, also entlang der x_3 -Achse. Dies zeigt, dass die Kugelkoordinaten dort tatsächlich eine Koordinatensingularität besitzen.

Analog ergibt sich für Zylinderkoordinaten mit (3.7.10)

$$\det \frac{\partial(x_1, x_2, x_3)}{\partial(R, \varphi, z)} = (\underline{T}_r \times \underline{T}_\varphi) \cdot \underline{T}_z = R. \quad (3.12.6)$$

Auch hier verschwindet die Jacobi-Determinante entlang der x_3 -Achse, wo auch die Zylinderkoordinaten eine Koordinatensingularität aufweisen.

Als *Beispiele* für Volumenintegrale berechnen wir das Volumen einer Kugel (Radius a und Höhe h) und eines Zylinders (mit Radius a und Höhe h) mit Hilfe der beiden eben definierten Kugel- bzw. Zylinderkoordinaten.

Für die Kugel gilt

$$\begin{aligned}
 V_{\text{Kugel}} &= \int_{K_a} d^3x = \int_0^a dr \int_0^\pi d\vartheta \int_0^{2\pi} \varphi r^2 \sin \vartheta \\
 &= 2\pi \int_0^a dr \int_0^\pi d\vartheta r^2 \sin \vartheta \\
 &= 2\pi \int_0^a dr r^2 [-\sin \vartheta]_0^\pi \\
 &= 4\pi \int_0^a dr r^2 = \frac{4\pi}{3} a^3
 \end{aligned} \tag{3.12.7}$$

und für den Zylinder

$$V_{\text{Zylinder}} = \int_{Z(a,h)} d^3x = \int_0^a dR \int_0^{2\pi} d\varphi \int_{-h/2}^{h/2} dz R = h \int_0^a dR \int_0^{2\pi} d\varphi R = 2\pi h \int_0^a dR R = \pi a^2 h. \tag{3.12.8}$$

3.12.2 Die koordinatenunabhängige Definition der Divergenz

Ähnlich wie wir in Abschnitt 3.10.4 die Rotation durch den Limes eines Wegintegrals definiert haben, geben wir nun eine Definition der Divergenz über ein **Flächenintegral** an. Dazu definieren wir die Orientierung der (notwendig geschlossenen) Randfläche ∂B eines dreidimensionalen Bereiches B so, dass die Flächennormalenvektoren nach außen, also **von dem betrachteten Volumen weg** weisen. Für die Kugel ist die Randfläche in Kugelkoordinaten einfach durch $r = a = \text{const}$ gegeben, und wir gelangen wieder zur Parametrisierung (3.10.1) der entsprechenden Kugelschale. Der Normaleneinheitsvektor ist dabei stets $\vec{e}_r$ und weist entsprechend unserer Definition in die korrekte Richtung weg von der Kugel.

Beim Zylinder zerfällt die Randfläche in drei Teile, nämlich den Zylinder-Mantel $R = a$ (generalisierte Koordinaten (φ, z) (der Flächennormalenvektor $\vec{N} = \vec{T}_\varphi \times \vec{T}_z = R\vec{e}_R$ weist in die korrekte Richtung) sowie die Deckfläche ($z = h/2$ mit generalisierten Koordinaten (R, φ) , $\vec{N} = R\vec{e}_3$) und die Bodenfläche, bei der wir die Orientierung umkehren müssen, damit der Flächennormalenvektor $\vec{N} = -R\vec{e}_3$ ist, also aus dem Zylinder hinausweist.

Betrachten wir nun ein partiell stetig differenzierbares Vektorfeld, können wir die Divergenz durch

$$\text{div } \vec{V}(\vec{x}) = \lim_{\Delta V \rightarrow \{\vec{x}\}} \frac{1}{\text{vol}(\Delta V)} \int_{\partial \Delta V} d^2\vec{f} \cdot \vec{V} \tag{3.12.9}$$

definieren. Dabei ist ΔV ein Volumenbereich mit dem Volumen $\text{vol}(\Delta V)$, der ganz im Definitionsbereich des Vektorfeldes liegt, und der Limes ist so zu verstehen, dass dieses Volumenelement auf den Punkt $\vec{x}$ zusammengezogen wird.

Um zu zeigen, dass in kartesischen Koordinaten diese Definition mit der oben in (3.6.3) gegebenen Definition übereinstimmt, wählen wir ΔV als Würfel der Kantenlänge ϵ mit dem Mittelpunkt in $\vec{x}$ und führen eine ähnliche Rechnung wie in Abschnitt 3.10.4 durch. Betrachten wir als Beispiel den Beitrag der beiden zur 12-Ebene parallelen Würfel Flächen, die bei $x_3 - \epsilon/2$ bzw. $x_3 + \epsilon/2$ liegen. Mit dem Mittelwertsatz ergibt sich für diese beiden Flächenintegrale

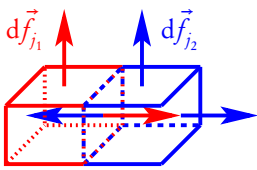
$$\epsilon^2 [V_3(\xi_1, \xi_2, x_3 + \epsilon/2) - V_3(\xi_1, \xi_2, x_3 - \epsilon/2)] = \epsilon^3 \partial_3 V_3(\xi_1'', \xi_2'', \xi_3). \tag{3.12.10}$$

Dies durch das Volumen ϵ^3 des Würfels dividiert und den Limes $\epsilon \rightarrow 0$ gebildet liefert den Beitrag $\partial_3 V_3(\underline{x})$. Entsprechend erhält man für die beiden anderen Seitenflächen des Würfels die beiden anderen Beiträge, so dass (3.12.9) tatsächlich das gleiche Resultat wie (3.6.3) ergibt.

3.12.3 Der Gaußsche Integralsatz

Mit der obigen Definition der Divergenz ist der **Gaußsche Integralsatz** eine selbstverständliche Folgerung. Ist $\vec{V}$ ein stetig partiell differenzierbares Vektorfeld, B ein Volumenbereich und ∂B sein entsprechend der Orientierungsvorschrift, dass die Flächennormalenvektoren aus dem Bereich B herauszeigen, so gilt

$$\int_B d^3x \operatorname{div} \vec{V}(\vec{x}) = \int_{\partial B} d^2\vec{f} \cdot \vec{V}(\vec{x}). \quad (3.12.11)$$



Um dies zu beweisen, müssen wir nur den Bereich B in viele kleine Teilvolumenelemente ΔV_j zerlegen. Nach dem Mittelwertsatz der Integralrechnung ergibt sich dann analog wie beim Satz von Stokes

$$\int_{\Delta B_j} d^3x \operatorname{div} \vec{V} = \operatorname{vol}(\Delta B_j) \operatorname{div} \vec{V}(\vec{\xi}) \simeq \int_{\partial \Delta B_j} d^2\vec{f} \cdot \vec{V}(\vec{x}). \quad (3.12.12)$$

Dabei wird die letztere Beziehung im Limes $\Delta B_j \rightarrow 0$, also bei immer feinerer Unterteilung des Volumens B in Teilvolumina exakt. Summiert man nun die linke Seite dieser Gleichung auf, erhält man das Volumenintegral auf der linken Seite von (3.12.11), und auf der rechten Seite heben sich die Beiträge von den inneren Oberflächenteilen $\partial \Delta B_j$ weg, da diese in der Summe stets zweimal mit unterschiedlicher Orientierung auftauchen (s. Skizze).

3.12.4 Die Greenschen Integralsätze im Raum

Die **Greenschen Integralsätze** im Raum sind Spezialfälle des Gaußschen Integralsatzes. Der **1. Greensche Integralsatz** ergibt sich aus dem Gaußschen Integralsatz, indem wir ihn auf das Vektorfeld

$$\vec{V}(\vec{x}) = \Phi_1(\vec{x}) \vec{\nabla} \Phi_2(\vec{x}) \quad (3.12.13)$$

anwenden. Wir berechnen als erstes die Divergenz. Dazu verwenden wir am einfachsten die Darstellung in kartesischen Komponenten. Zunächst ist

$$V_j(\underline{x}) = \tilde{\Phi}_1(\underline{x}) \partial_j \tilde{\Phi}_2(\underline{x}). \quad (3.12.14)$$

Aus der Produktregel folgt daraus

$$\operatorname{div} \vec{V}(\vec{x}) = \sum_{j=1}^3 \partial_j V_j(\underline{x}) = \sum_{j=1}^3 [\partial_j \tilde{\Phi}_1(\underline{x}) \partial_j \tilde{\Phi}_2(\underline{x}) - \tilde{\Phi}_1(\underline{x}) \partial_j^2 \tilde{\Phi}_2(\underline{x})]. \quad (3.12.15)$$

Dies können wir wieder in koordinatenunabhängiger Form als

$$\operatorname{div} \vec{V}(\vec{x}) = \vec{\nabla} \cdot \vec{V}(\vec{x}) = [\vec{\nabla} \Phi_1(\vec{x})] \cdot [\vec{\nabla} \Phi_2(\vec{x})] + \Phi_1(\vec{x}) \vec{\nabla}^2 \Phi_2(\vec{x}). \quad (3.12.16)$$

Der Differentialoperator $\vec{\nabla}^2$ kommt so häufig in der Feldtheorie vor, dass man dafür ein eigenes Symbol Δ , den **Laplace-Operator**, einführt. Wirkt es auf ein zweimal stetig partiell differenzierbares Skalarfeld, erhält man wieder ein Skalarfeld, und zwar

$$\Delta \Phi(\vec{x}) = \vec{\nabla} \cdot \vec{\nabla} \Phi(\vec{x}) = \operatorname{div} \operatorname{grad} \Phi(\vec{x}) = \vec{\nabla}^2 \Phi(\vec{x}). \quad (3.12.17)$$

In kartesischen Komponenten gilt ausgeschrieben

$$\Delta \tilde{\Phi}(\underline{x}) = \partial_1^2 \Phi(\underline{x}) + \partial_2^2 \Phi(\underline{x}) + \partial_3^2 \Phi(\underline{x}). \quad (3.12.18)$$

Setzen wir also dieses Vektorfeld in den Gaußschen Integralsatz ein, folgt der **1. Greensche Integralsatz**

$$\int_B d^3x \{ [\vec{\nabla} \Phi_1(\vec{x})] \cdot [\vec{\nabla} \Phi_2(\vec{x})] + \Phi_1(\vec{x}) \Delta \Phi_2(\vec{x}) \} = \int_{\partial B} d\vec{f} \cdot \Phi_1(\vec{x}) \vec{\nabla} \Phi_2(\vec{x}). \quad (3.12.19)$$

Der **2. Greensche Integralsatz** folgt daraus, indem wir in dieser Gleichung Φ_1 und Φ_2 vertauschen und das Resultat von der vorigen Gleichung abziehen. Dabei verschwindet im Volumenintegral der erste Term, weil dieser symmetrisch unter dieser Vertauschung ist:

$$\int_B d^3x [\Phi_1(\vec{x}) \Delta \Phi_2(\vec{x}) - \Phi_2(\vec{x}) \Delta \Phi_1(\vec{x})] = \int_{\partial B} d\vec{f} \cdot [\Phi_1(\vec{x}) \vec{\nabla} \Phi_2(\vec{x}) - \Phi_2(\vec{x}) \vec{\nabla} \Phi_1(\vec{x})]. \quad (3.12.20)$$

Dieser Satz wird oft in der **Potentialtheorie** gebraucht. Wir werden ihn in Abschnitt 3.15 verwenden, um die einfachste Aufgabe der Potentialtheorie zu lösen.

3.13 Alternative Herleitung der Differentialoperatoren in krummlinigen Orthogonalkoordinaten

Wir wollen nochmals die bereits in Abschnitt 3.7 hergeleiteten Formeln für die Differentialoperatoren mittels der koordinatenunabhängigen Definitionen für rot und div in den Abschnitten 3.10.4 bzw. 3.12.2 bestätigen. Zur Berechnung der **Divergenz des Vektorfeldes** wenden wir (3.12.9) auf den infinitesimalen Quader ΔQ , der von den Tangentenvektoren der Koordinatenlinien, also

$$\frac{\partial \underline{x}}{\partial q_i} dq_i = \underline{T}_i dq_i = g_i \vec{e}_i' dq_i, \quad (3.13.1)$$

aufgespannt wird, an. Für das Volumenintegral über diesen infinitesimalen Quader ergibt sich

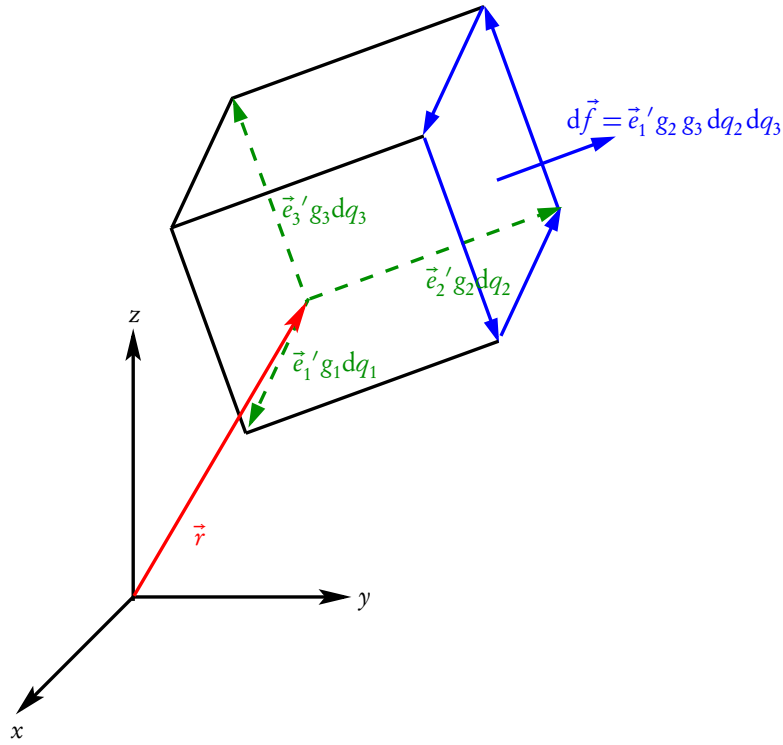
$$\int_{\Delta Q} d^3x [\operatorname{div} \vec{A}(\underline{x})] = d^3x [\operatorname{div} \vec{A}(\underline{x})] = g_1 g_2 g_3 d^3q \operatorname{div} \vec{A}. \quad (3.13.2)$$

Dabei haben wir das Volumenelement d^3x mittels (3.12.1) umgeschrieben.

Das dazugehörige Flächenintegral über den Rand $\partial \Delta Q$ unseres Quaders setzt sich aus den sechs infinitesimalen Seitenflächen des Quaders zusammen. Der Skizze unten entnehmen wir

$$\begin{aligned} \int_{\partial \Delta Q} d^2 \vec{f} \cdot \vec{A}(\underline{x}) &= dq_2 dq_3 ([g_2 g_3 A_1']_{q_1+dq_1, q_2, q_3} - [g_2 g_3 A_1']_{q_1, q_2, q_3}) \\ &\quad + dq_3 dq_1 ([g_3 g_1 A_2']_{q_1, q_2+ dq_2, q_3} - [g_3 g_1 A_2']_{q_1, q_2, q_3}) \\ &\quad + dq_1 dq_2 ([g_1 g_2 A_3']_{q_1, q_2, q_3+ dq_3} - [g_1 g_2 A_3']_{q_1, q_2, q_3}) \\ &= d^3q \left[\frac{\partial (g_2 g_3 A_1')}{\partial q_1} + \frac{\partial (g_3 g_1 A_2')}{\partial q_2} + \frac{\partial (g_1 g_2 A_3')}{\partial q_3} \right]. \end{aligned} \quad (3.13.3)$$

3. Vektoranalysis



Verwenden wir nun (3.13.2) und (3.13.3) in (3.12.9), erhalten wir schließlich

$$\operatorname{div} \vec{A} = \frac{1}{g_1 g_2 g_3} \left[\frac{\partial (g_2 g_3 A'_1)}{\partial q_1} + \frac{\partial (g_3 g_1 A'_2)}{\partial q_2} + \frac{\partial (g_1 g_2 A'_3)}{\partial q_3} \right]. \quad (3.13.4)$$

Analog erhält man die Komponenten für die **Rotation des Vektorfeldes** aus der Anwendung von (3.10.27) auf die drei infinitesimalen Rechteckflächen ΔF_{ij} , die jeweils durch die Tangentenvektoren an die Koordinatenlinien $\underline{T}_i dq_i = g_i \vec{e}_i'$ und $\underline{T}_j dq_j = g_j \vec{e}_j'$ aufgespannt werden, wobei nacheinander $(i, j) \in \{(2, 3); (3, 1); (1, 2)\}$ gesetzt wird. Verwenden wir z.B. das erste Indexpaar, erhalten wir die erste Komponente der Rotation bzgl. der krummlinigen Koordinaten (s. wieder die Skizze oben). Das Flächenintegral ist

$$\int_{\Delta F_{23}} d^2 \vec{f} \cdot \operatorname{rot} \vec{A} = dq_2 dq_3 g_2 g_3 (\vec{e}_2' \times \vec{e}_3') \cdot \operatorname{rot} \vec{A} = dq_2 dq_3 g_2 g_3 \vec{e}_1' \cdot \operatorname{rot} \vec{A} = dq_2 dq_3 g_2 g_3 (\operatorname{rot} \vec{A})'_1, \quad (3.13.5)$$

und das dazugehörige Linienintegral entlang des Randes $\partial \Delta F_{23}$ lautet

$$\begin{aligned} \int_{\partial \Delta F_{23}} d\vec{x} \cdot \vec{A} &= dq_3 \left[(g_3 A'_3)_{q_1, q_2 + dq_2, q_3} - (g_3 A'_3)_{q_1, q_2, q_3} \right] \\ &\quad - dq_2 \left[(g_2 A'_2)_{q_1, q_2, q_3 + dq_3} - (g_2 A'_2)_{q_1, q_2, q_3} \right] \\ &= dq_2 dq_3 \left[\frac{\partial (g_3 A'_3)}{\partial q_2} - \frac{\partial (g_2 A'_2)}{\partial q_3} \right]. \end{aligned} \quad (3.13.6)$$

Wegen des Stokesschen Integralsatzes sind (3.13.5) und (3.13.6) gleich, so dass sich schließlich

$$(\operatorname{rot} \vec{A})'_1 = \frac{1}{g_2 g_3} \left[\frac{\partial(g_3 A'_3)}{\partial q_2} - \frac{\partial(g_2 A'_2)}{\partial q_3} \right] \quad (3.13.7)$$

ergibt. Die beiden übrigen Komponenten finden wir auf analoge Weise durch Verwendung der anderen beiden infinitesimalen Rechteckflächen. Wir erhalten sie jedoch auch einfach durch **zyklische Vertauschung** der Indizes:

$$\begin{aligned} (\operatorname{rot} \vec{A})'_2 &= \frac{1}{g_3 g_1} \left[\frac{\partial(g_1 A'_1)}{\partial q_3} - \frac{\partial(g_3 A'_3)}{\partial q_1} \right], \\ (\operatorname{rot} \vec{A})'_3 &= \frac{1}{g_1 g_2} \left[\frac{\partial(g_2 A'_2)}{\partial q_1} - \frac{\partial(g_1 A'_1)}{\partial q_2} \right]. \end{aligned} \quad (3.13.8)$$

3.14 Solenoidalfelder und Vektorpotentiale

In den Abschnitten 3.9 und 3.11 haben wir die Frage untersucht, unter welchen Umständen ein vorgegebenes Vektorfeld $\vec{V}$ der Gradient eines Skalarfeldes ist und wie ggf. dieses Skalarfeld, das **Potential** des Vektorfeldes, berechnet werden kann. Die Antwort war das **Lemma von Poincaré**, dass falls $\vec{V}$ auf einem einfach zusammenhängenden Gebiet stetig partiell differenzierbar ist und dort $\operatorname{rot} \vec{V} = 0$ gilt, in diesem Gebiet ein bis auf eine Konstante eindeutig bestimmtes Potential Φ existiert, so dass $\vec{V} = -\operatorname{grad} \Phi$ gilt. Das Potential ist dabei durch das Wegintegral entlang eines beliebigen Weges innerhalb dieses Gebiets von einem festen Punkt $\vec{x}_0$ zum Punkt $\vec{x}$ gegeben und bis auf eine additive Konstante eindeutig bestimmt.

Hier fragen wir nun, unter welchen Bedingungen ein stetig partiell differenzierbares Vektorfeld $\vec{B}$ die Rotation eines anderen Vektorfeldes $\vec{A}$ ist, d.h.

$$\vec{B}(\vec{x}) = \vec{\nabla} \times \vec{A}(\vec{x}). \quad (3.14.1)$$

Man nennt dann $\vec{A}$ ein **Vektorpotential** für $\vec{B}$.

Ein wichtiges Beispiel für Solenoidalfelder in der Physik sind **Magnetfelder**.

Zuerst leiten wir wieder eine notwendige Bedingung her. Nehmen wir also an, es gäbe ein Vektorpotential für $\vec{B}$, so dass (3.14.1) erfüllt ist. Dann ist

$$\operatorname{div} \vec{B} = \vec{\nabla} \cdot \vec{B} = \vec{\nabla} \cdot (\vec{\nabla} \times \vec{A}). \quad (3.14.2)$$

Wäre $\vec{\nabla}$ ein gewöhnlicher Vektor, würde die rechte Seite verschwinden. Wir zeigen nun, dass diese Annahme tatsächlich zutrifft, indem wir die Divergenz in kartesischen Koordinaten explizit ausrechnen. Zunächst gilt

$$\begin{aligned} B_1 &= \partial_2 A_3 - \partial_3 A_2 \Rightarrow \partial_1 B_1 = \partial_1 \partial_2 A_3 - \partial_1 \partial_3 A_2, \\ B_2 &= \partial_3 A_1 - \partial_1 A_3 \Rightarrow \partial_2 B_2 = \partial_2 \partial_3 A_1 - \partial_2 \partial_1 A_3, \\ B_3 &= \partial_1 A_2 - \partial_2 A_1 \Rightarrow \partial_3 B_3 = \partial_3 \partial_1 A_2 - \partial_3 \partial_2 A_1. \end{aligned} \quad (3.14.3)$$

Addieren wir nun die Gleichungen nach den Folgepfeilen, ergibt sich die Divergenz

$$\operatorname{div} \vec{B} = \partial_1 B_1 + \partial_2 B_2 + \partial_3 B_3 = \partial_1 \partial_2 A_3 - \partial_1 \partial_3 A_2 + \partial_2 \partial_3 A_1 - \partial_2 \partial_1 A_3 + \partial_3 \partial_1 A_2 - \partial_3 \partial_2 A_1. \quad (3.14.4)$$

3. Vektoranalysis

Da voraussetzungsgemäß die partiellen Ableitungen der Komponenten von $\vec{B}$ stetig sind, gilt dies notwendig auch für die zweiten Ableitungen der Komponenten von $\vec{A}$. Dann vertauschen aber die Ableitungsoperatoren, und daher folgt aus (3.14.4) tatsächlich, dass

$$\operatorname{div} \vec{B} = 0 \quad (3.14.5)$$

ist.

Für die Existenz eines Vektorpotentials ist also *notwendig* (3.14.5) erfüllt.

Wir wollen nun zeigen, dass (3.14.5) auch hinreichend ist für die **lokale Existenz eines Vektorpotentials**. Leider gibt es keine einfache Formel wie beim Skalarpotential, das wir wie oben gezeigt als Wegintegral darstellen können. Wir müssen also versuchen, die Gleichung (3.14.1) nach $\vec{A}$ aufzulösen.

Dazu bemerken wir aber als erstes, dass offensichtlich das Vektorpotential nur **bis auf den Gradienten eines Skalarfeldes** bestimmt ist, denn wie wir oben gesehen haben, gilt für jedes zweimal stetig differenzierbare Skalarfeld χ

$$\vec{\nabla} \times \vec{\nabla} \chi = \operatorname{rot} \operatorname{grad} \chi = 0. \quad (3.14.6)$$

Das bedeutet, dass für jedes Vektorpotential $\vec{A}$ von $\vec{B}$ auch

$$\vec{A}' = \vec{A} + \vec{\nabla} \chi \quad (3.14.7)$$

ein Vektorpotential von $\vec{B}$ ist. Das Vektorpotential eines Vektorfeldes ist also, falls es überhaupt existiert, nur bis auf ein Gradientenfeld bestimmt.

Wir können also zur Vereinfachung der Auflösung von (3.14.1) nach $\vec{A}$ eine **Nebenbedingung** fordern. Wie wir sehen werden, erweist sich die Bedingung

$$A_3 = 0 \quad (3.14.8)$$

als sehr hilfreich. Um zu zeigen, dass man dies durch Wahl eines geeigneten Skalarfeldes χ stets erreichen kann, nehmen wir an $\vec{A}$ sei irgendein Vektorpotential von $\vec{B}$ und suchen ein Skalarfeld χ , so dass

$$A_3' = A_3 + \partial_3 \chi = 0 \quad (3.14.9)$$

ist. Diese Gleichung lässt sich aber sehr leicht lösen. Wir müssen sie nur bzgl. x_3 integrieren:

$$\chi(x_1, x_2, x_3) = - \int_{x_{30}}^{x_3} dx_3' A_3(x_1, x_2, x_3'). \quad (3.14.10)$$

Dabei gehen wir davon aus, dass es um den Punkt $\vec{x}$ eine offene quaderförmige Umgebung parallel zu den Koordinatenachsen gibt, wo $\vec{A}$ stetig differenzierbar ist. Dies ist sicher erfüllt, wenn $\vec{B}$ in einem offenen Gebiet stetig ist. Dann liegt der Integrationsbereich von (3.14.10) ganz in dieser Umgebung, und das Integral ist daher eine Lösung von (3.14.9). Wir dürfen also annehmen, dass die Nebenbedingung (3.14.8) erfüllt ist.

In kartesischen Komponenten ausgeschrieben lautet dann die Gleichung (3.14.1)

$$\underline{B} = \underline{\nabla} \times \begin{pmatrix} A_1 \\ A_2 \\ 0 \end{pmatrix} = \begin{pmatrix} -\partial_3 A_2 \\ \partial_3 A_1 \\ \partial_1 A_2 - \partial_2 A_1 \end{pmatrix}. \quad (3.14.11)$$

Aus der ersten Komponente dieser Gleichung folgt

$$A_2(x_1, x_2, x_3) = - \int_{x_{30}}^{x_3} dx_3' B_1(x_1, x_2, x_3') + A_2'(x_1, x_2) \quad (3.14.12)$$

und aus der zweiten

$$A_1(x_1, x_2, x_3) = \int_{x_{30}}^{x_3} dx'_3 B_1(x_1, x_2, x'_3) + A'_1(x_1, x_2). \quad (3.14.13)$$

Dabei haben wir beim Integrieren berücksichtigt, dass die Vorgabe der partiellen Ableitung nach x_3 eine Funktion nur bis auf Funktionen, die von den beiden anderen unabhängigen Variablen x_1 und x_2 abhängen, bestimmt ist. Diese können wir nun bestimmen, indem wir (3.14.12) und (3.14.13) in die dritte Komponente der Gleichung (3.14.11) einsetzen:

$$B_3(x_1, x_2, x_3) = - \int_{x_{30}}^{x_3} dx'_3 [\partial_1 B_1(x_1, x_2, x'_3) + \partial_2 B_2(x_1, x_2, x'_3)] + \partial_1 A'_2(x_1, x_2) - \partial_2 A'_1(x_1, x_2). \quad (3.14.14)$$

Wegen $\operatorname{div} \vec{B} = \partial_1 B_1 + \partial_2 B_2 + \partial_3 B_3 = 0$ folgt daraus

$$B_3(x_1, x_2, x_3) = \int_{x_{30}}^{x_3} dx'_3 \partial'_3 B_3(x_1, x_2, x'_3) + \partial_1 A'_2(x_1, x_2) - \partial_2 A'_1(x_1, x_2). \quad (3.14.15)$$

Da voraussetzungsgemäß $\vec{B}$ stetig partiell differenzierbar sein soll, können wir das Integral auswerten, woraus sich

$$\begin{aligned} B_3(x_1, x_2, x_3) &= B_3(x_1, x_2, x_3) - B_3(x_1, x_2, x_{30}) + \partial_1 A'_2(x_1, x_2) - \partial_2 A'_1(x_1, x_2) \\ &\Rightarrow \partial_1 A'_2(x_1, x_2) - \partial_2 A'_1(x_1, x_2) = B_3(x_1, x_2, x_{30}) \end{aligned} \quad (3.14.16)$$

ergibt. Dies zeigt, dass dank der Divergenzfreiheit des Vektorfeldes $\vec{B}$ die Lösungen (3.14.14) tatsächlich konsistent mit der Existenz des Vektorpotentials $\vec{A}$ mit $A_3 = 0$ sind.

Die Nebenbedingung legt allerdings offensichtlich das Vektorpotential immer noch nicht eindeutig fest. Das ist unmittelbar einleuchtend, denn die Addition des Gradienten eines nur von x_1 und x_2 abhängigen Skalarfeldes ändert an den Eigenschaften des Vektorpotentials, $\operatorname{rot} \vec{A} = \vec{B}$ und $A_3 = 0$ zu erfüllen nichts. Wir können also in (3.14.13) $A'_1(x_1, x_2) = 0$ setzen. Dann folgt aus (3.14.16)

$$A'_1 = 0 \Rightarrow \partial_1 A'_2(x_1, x_2) = B_3(x_1, x_2, x_{30}). \quad (3.14.17)$$

Damit wird durch die Wahl

$$A'_2(x_1, x_2) = \int_{x_{20}}^{x_2} dx'_2 B_3(x_1, x'_2, x_{30}) \quad (3.14.18)$$

das Vektorpotential vervollständigt. Sammeln wir also die Lösungsschritte (3.14.12), (3.14.17) und (3.14.18) erhalten wir als eine mögliche Lösung für das Vektorpotential

$$\begin{aligned} A_1(\underline{x}) &= \int_{x_{30}}^{x_3} dx'_3 B_2(x_1, x_2, x'_3), \\ A_2(\underline{x}) &= - \int_{x_{20}}^{x_2} dx'_2 B_1(x_1, x_2, x'_2) + \int_{x_{20}}^{x_2} dx'_2 B_3(x_1, x'_2, x_{30}), \\ A_3(\underline{x}) &= 0. \end{aligned} \quad (3.14.19)$$

3.15 Die Poisson-Gleichung und Green-Funktionen

In diesem Abschnitt beschäftigen wir uns mit einer typischen Fragestellung der Feldtheorie. Dazu betrachten wir das Newtonsche Gravitationsgesetz etwas genauer. Wie schon in Abschnitt 3.8 gesagt, ist die Gravitationskraft einer als punktförmig angenommenen im Ursprung des Koordinatensystems sitzenden Masse M auf eine andere bei $\vec{x}$ gelegene Punktmasse m , die wir im folgenden **Problemasse** nennen, durch das Kraftfeld

$$\vec{F}(\vec{x}) = -\gamma m M \frac{\vec{x}}{r^3}, \quad r = |\vec{x}| \quad (3.15.1)$$

gegeben.

Dividieren wir diese Kraft durch die Masse der Probemasse, erhalten wir die **Gravitationsbeschleunigung**, die unabhängig von der Probemasse ist. Wir gelangen so zur Interpretation der Gravitation als **Feldwirkung**.

Demnach erzeugt die Masse M ein Gravitationsfeld

$$\vec{g}(\vec{x}) = -\gamma M \frac{\vec{x}}{r^3}. \quad (3.15.2)$$

Man weist es einfach dadurch nach, dass man an die Stelle $\vec{x}$ eine Probemasse m setzt und die auf sie wirkende Gravitationskraft $\vec{F}(\vec{x}) = m\vec{g}(\vec{x})$ (3.15.1) misst. Dadurch haben wir die Gravitationskraft als „Nahwirkungstheorie“ formuliert. In dieser Interpretation wirkt die Kraft nicht über eine Fernwirkung sondern weil die Präsenz der Masse M im Ursprung das **Gravitationsfeld** (3.15.2) impliziert, und die Kraft wirkt dann lokal an der Stelle $\vec{x}$ der Probemasse aufgrund der Gegenwart dieses Feldes.

In Abschnitt 3.8 haben wir auch gezeigt, dass dieses Feld ein Potentialfeld ist, das wegen (3.8.4) durch

$$\Phi(\vec{x}) = -\frac{\gamma M}{r} \quad (3.15.3)$$

gegeben ist. Wir nennen dieses Feld das **Gravitationspotential**. Dabei haben wir die willkürliche Konstante $C = 0$ gesetzt. Dabei folgen wir der Konvention, dass das Gravitationspotential im Unendlichen verschwinden soll.

Wir fragen nun, wie die Gravitationskraft einer ausgedehnten Massenverteilung auf eine Punktmasse zu berechnen ist. Dazu denken wir uns die Masse kontinuierlich gemäß der **Massendichteverteilung** $\rho(\vec{x})$ über den Körper verteilt. Mit Newton gehen wir nun davon aus, dass jedes Massenelement am Punkt $\vec{x}'$, also $dM = d^3x' \rho(\vec{x}')$ additiv zur Kraft auf eine Probemasse beiträgt. Dabei müssen wir nur beachten, dass der Ursprung dieser infinitesimalen **Quelle des Gravitationsfeldes** bei $\vec{x}'$ und nicht im Koordinatenursprung sitzt. Der entsprechende Beitrag zum Gravitationspotential ist demnach

$$d\Phi = -\frac{\gamma dM}{|\vec{x} - \vec{x}'|}. \quad (3.15.4)$$

Insgesamt ist das Gravitationspotential der Massenverteilung also durch

$$\Phi(\vec{x}) = -\gamma \int_V d^3x' \frac{\rho(\vec{x}')}{|\vec{x} - \vec{x}'|} \quad (3.15.5)$$

gegeben. Dabei nehmen wir an, dass V so gewählt ist, dass der gesamte Körper ganz im Inneren dieses Volumens liegt, also $\rho(\vec{x}') = 0$ für $\vec{x}' \in E^3 \setminus V$ gilt. Der Bequemlichkeit halber sei auch $\rho(\vec{x}')|_{\vec{x}' \in \partial V} = 0$.

Im Prinzip können wir mit (3.15.5) das Gravitationspotential einer beliebigen Massenverteilung ausrechnen. In der Physik haben sich aber **lokale Feldgleichungen** als weitaus nützlicher erwiesen, d.h. wir suchen eine **partielle Differentialgleichung**, die das Gravitationspotential bzw. die Gravitationsbeschleunigung mit der Massenverteilung $\rho(\vec{x})$ verknüpft.

Dazu betrachten wir noch einmal das Gravitationsfeld (3.15.2) einer einzelnen Punktmasse im Ursprung des Koordinatensystems. Eine solche Punktmasse ist in der Feldtheorie ein wenig problematisch, denn offenbar ist für solch eine Massenverteilung die Massendichte extrem singulär. Außer im Ursprung ist die Massendichte exakt 0, und da wir eine endliche Masse im Ursprung vereint haben, ist dort die Massendichte unendlich groß.

Wir können nun aber die im Ursprung konzentrierte Masse M durch ein **Flächenintegral** über das Gravitationsfeld (3.15.2) erhalten. Dazu sei V ein Volumen, bei dem der Koordinatenursprung durch eine kleine Kugel mit Radius ϵ ausgespart ist. Sei V' nun also ein Volumen, das den Ursprung im Inneren enthält mit auf übliche Weise orientierter Randfläche $\partial V'$ (Flächennormalenvektoren weisen aus dem Volumen heraus) und sei die Kugel mit Radius ϵ um den Ursprung $K_\epsilon(0) \subset V'$. Wir wenden nun den Gaußschen Satz auf das Gravitationsfeld (3.15.2) und integrieren über das Volumen $V = V' \setminus K_\epsilon(0)$ an. Dann besteht der auf die übliche Weise orientierte Rand von V zum Einen aus dem Rand von V' und zum anderen aus der Kugeloberfläche $S_\epsilon(0)$, die so orientiert ist, dass die Normalenvektoren radial auf den Ursprung weisen, also ebenfalls weg vom Volumen V .

Durch direktes Nachrechnen in kartesischen Koordinaten weist man nun nach, dass

$$\operatorname{div} \vec{g}(\vec{x}) = \vec{\nabla} \cdot \vec{g}(\vec{x}) = 0 \quad \text{für} \quad \vec{x} \neq \vec{0} \quad (3.15.6)$$

gilt (*Nachrechnen!*). Demnach folgt aus dem Gaußschen Integralsatz

$$0 = \int_V d^3x \operatorname{div} \vec{g}(\vec{x}) = \int_{\partial V} d^2\vec{f} \cdot \vec{g}(\vec{x}) = \int_{\partial V'} d^2\vec{f} \cdot \vec{g}(\vec{x}) - \int_{S_\epsilon(0)} d^2\vec{f} \cdot \vec{g}(\vec{x}), \quad (3.15.7)$$

wobei wir in dem letzteren Flächenintegral die Orientierung wieder wie üblich mit den Normalenvektoren vom Kugelursprung weg gerichtet haben. Da für den Gaußschen Integralsatz allerdings die Orientierung über die Oberfläche der ausgesparten Kugel um den Ursprung umgekehrt zu richten ist, trägt das entsprechende Flächenintegral mit dem negativen Vorzeichen bei. Jedenfalls folgt aus (3.15.7)

$$\int_{\partial V'} d^2\vec{f} \cdot \vec{g}(\vec{x}) = \int_{S_\epsilon(0)} d^2\vec{f} \cdot \vec{g}. \quad (3.15.8)$$

Das rechtsstehende Integral lässt sich nun sehr leicht ausrechnen. Dazu verwenden wir unsere Standardparametrisierung der Kugelschale (3.10.1) und das entsprechende Resultat (3.10.12):

$$\int_{\partial V'} d^2\vec{f} \cdot \vec{g}(\vec{x}) = -4\pi\gamma M. \quad (3.15.9)$$

Für irgendeine den Ursprung umschließende geschlossene Oberfläche $\partial V'$ erhalten wir also aus dem Gravitationspotential bis auf die Vorfaktoren $-4\pi\gamma$ die in dem umschlossenen Volumen enthaltene felderzeugende Masse.

Nun liegt die Vermutung nahe, dass dies auch für die kontinuierliche Massenverteilung gilt. Dies führt auf die Idee, dass für eine kontinuierliche singularitätenfreie Massenverteilung das Gravitationspotential überall eine wohldefinierte Divergenz besitzt, und dann können wir den Gaußschen Integralsatz auf irgendein Volumen V anwenden, ohne Singularitäten durch kleine Kugeln ausschließen zu müssen. Diese Überlegung liefert

$$\int_V d^3x \operatorname{div} \vec{g}(\vec{x}) = \int_{\partial V} d^2\vec{f} \cdot \vec{g}(\vec{x}) \stackrel{!}{=} -4\pi\gamma \int_V d^3x \rho(\vec{x}) = -4\pi\gamma M_V, \quad (3.15.10)$$

wobei M_V die gesamte im Volumen V befindliche Masse ist. Da dies für jedes Volumen V gilt, können wir die Definition der Divergenz über ein Flächenintegral verwenden (s. Abschnitt 3.12.2). Ist nämlich unsere Hypothese (3.15.10) korrekt, so folgt

$$\operatorname{div} \vec{g}(\vec{x}) = \lim_{\Delta V \rightarrow 0} \frac{1}{\operatorname{vol}(\Delta V)} \int_{\partial \Delta V} d^2\vec{f}' \cdot \vec{g}(\vec{x}') = -4\pi\gamma \lim_{\Delta V \rightarrow 0} \frac{M_{\Delta V}}{\Delta V} = -4\pi\gamma \rho(\vec{x}). \quad (3.15.11)$$

Dies führt uns auf die Differentialgleichung

$$\operatorname{div} \vec{g}(\vec{x}) = -4\pi\gamma \rho(\vec{x}). \quad (3.15.12)$$

Nun erinnern wir uns, dass das Gravitationsfeld ein Potential besitzt. Es gilt demnach

$$\operatorname{div} \vec{g}(\vec{x}) = -\operatorname{div} \operatorname{grad} \Phi(\vec{x}) = -\Delta \Phi(\vec{x}) = -4\pi\gamma\rho(\vec{x}). \quad (3.15.13)$$

Dabei haben wir den in (3.12.17) eingeführten Laplace-Operator verwendet.

Um nun zu zeigen, dass die Annahme in (3.15.10) wirklich zutrifft, zeigen wir, dass (3.15.5) tatsächlich die **Poisson-Gleichung** (3.15.13) löst. Falls $\vec{x} \notin V$ ist, können wir einfach den Laplaceoperator auf (3.15.5) anwenden, indem wir ihn in das Integral ziehen, denn für diesen Fall treten keine Singularitäten im Integranden auf. Nun ist aber

$$\Delta \frac{1}{|\vec{x} - \vec{x}'|} = \Delta' \frac{1}{|\vec{x} - \vec{x}'|} = 0, \quad \text{für } \vec{x} \neq \vec{x}', \quad (3.15.14)$$

wie man durch Bildung der partiellen Ableitungen direkt nachrechnet (*Übung!*). Daraus folgt, dass dann

$$\Delta \Phi(\vec{x}) = -\gamma \int_V d^3x' \rho(\vec{x}') \Delta \frac{1}{|\vec{x} - \vec{x}'|} = 0 \quad \text{für } \vec{x} \notin V \quad (3.15.15)$$

gilt. Da voraussetzungsgemäß für $\vec{x} \notin V$ stets $\rho(\vec{x}) = 0$ gelten soll, ist Die Poisson-Gleichung (3.15.13) für diesen Fall erfüllt.

Nun betrachten wir den Fall, dass $\vec{x} \in V$. Dazu wenden wir den 2. Greenschen Satz (3.12.20) an, wobei wir $\Phi_1 = \Phi$ und $\Phi_2(\vec{x}') = 1/|\vec{x} - \vec{x}'|$ setzen. Als Integrationsvolumen wählen wir $\tilde{V} = V \setminus K_\epsilon(\vec{x})$, wobei $K_\epsilon(\vec{x})$ eine Kugel mit Radius ϵ mit Mittelpunkt in $\vec{x}$ bezeichnet. Dann ist der Rand $\partial \tilde{V}$ durch ∂V und die Kugelschale $S_\epsilon(\vec{x})$ gegeben. Entsprechend der Standardorientierung beim Gaußschen Integralsatz muss man für letztere die Normalvektoren in Richtung auf den Mittelpunkt $\vec{x}$ weisend wählen (also vom Integrationsvolumen weg). Der Greensche Satz lautet für unseren Fall also

$$\int_{\tilde{V}} d^3\vec{x}' \left[\Phi(\vec{x}') \Delta \frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x} - \vec{x}'|} \Delta \Phi(\vec{x}') \right] = \int_{\partial \tilde{V}} d\vec{f}' \cdot \left[\Phi(\vec{x}') \vec{\nabla}' \frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x} - \vec{x}'|} \vec{\nabla}' \Phi(\vec{x}') \right]. \quad (3.15.16)$$

Im Volumenintegral auf der linken Seite fällt der erste Term wegen (3.15.14) weg und unter der Annahme, dass Φ der Poisson-Gleichung $\Delta \Phi = 4\pi\gamma\rho$ genügt, folgt

$$\int_{\tilde{V}} d^3\vec{x}' \left[\Phi(\vec{x}') \Delta \frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x} - \vec{x}'|} \Delta \Phi(\vec{x}') \right] = -4\pi\gamma \int_{\tilde{V}} d^3\vec{x}' \frac{\rho(\vec{x}')}{|\vec{x} - \vec{x}'|}. \quad (3.15.17)$$

Das Oberflächenintegral auf der rechten Seite von (3.15.16) werten wir nun für den Fall aus, dass wir $V = E^3$ setzen, wobei wir annehmen, dass der entsprechende Beitrag von ∂V , der jetzt im Unendlichen liegt, zum Oberflächenintegral nichts beiträgt, weil $\Phi(\vec{x})$ für große r mindestens wie $1/r$ abfällt. Dann bleibt von $\partial \tilde{V}$ nur noch die Oberfläche der kleinen Kugel, also $S_\epsilon(\vec{x})$ übrig, und dieses parametrisieren wir in den üblichen Kugelkoordinaten:

$$\underline{x}' = \underline{x} + \epsilon \begin{pmatrix} \cos \varphi \sin \vartheta \\ \sin \varphi \sin \vartheta \\ \cos \vartheta \end{pmatrix}. \quad (3.15.18)$$

Der Flächenelementvektor berechnet sich gemäß (3.10.6) zu (*Übung!*)

$$d\underline{f}' = d\vartheta d\varphi \epsilon^2 \sin \vartheta \begin{pmatrix} \cos \varphi \sin \vartheta \\ \sin \varphi \sin \vartheta \\ \cos \vartheta \end{pmatrix}, \quad (3.15.19)$$

mit der Standardorientierung vom Kugelzentrum weg, was wir wie bereits oben erwähnt im Flächenintegral auf der rechten Seite von (3.15.16) durch ein zusätzliches Vorzeichen berücksichtigen müssen. Für das Flächenintegral benötigen wir noch

$$\underline{\nabla}' \frac{1}{|\vec{x} - \vec{x}'|} = -\frac{\vec{x}' - \vec{x}}{|\vec{x} - \vec{x}'|^3} \Rightarrow d\underline{f}' \cdot \underline{\nabla}' \frac{1}{|\vec{x} - \vec{x}'|} = -d\vartheta d\varphi \sin \vartheta \quad (3.15.20)$$

und

$$d\vec{f}' \cdot \vec{\nabla}' \Phi(\vec{x}') = \epsilon^2 \partial_\epsilon \tilde{\Phi}(\underline{x}). \quad (3.15.21)$$

Dabei haben wir verwendet, dass in Kugelkoordinaten $(\epsilon, \vartheta, \varphi)$

$$\underline{e}_r = \underline{T}_r = \begin{pmatrix} \cos \varphi \sin \vartheta \\ \sin \varphi \sin \vartheta \\ \cos \vartheta \end{pmatrix} \quad (3.15.22)$$

und somit wegen der Kettenregel

$$\vec{e}_r \cdot \vec{\nabla} \Phi(\vec{x}') = \frac{\partial x}{\partial \epsilon} \cdot \underline{\nabla} \Phi(\underline{x}') = \partial_\epsilon \tilde{\Phi}(\underline{x}') \quad (3.15.23)$$

gilt. Setzen wir also (3.15.19-3.15.23) in die rechten Seite von (3.15.16) ein, so folgt

$$\int_{\partial \tilde{V}} d\vec{f}' \cdot \left[\Phi(\vec{x}') \vec{\nabla}' \frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x} - \vec{x}'|} \vec{\nabla}' \Phi(\vec{x}') \right] = - \int_0^\pi d\vartheta \int_0^{2\pi} d\varphi \sin \vartheta \left[-\tilde{\Phi}(\underline{x}') - \epsilon \partial_\epsilon \tilde{\Phi}(\underline{x}') \right]. \quad (3.15.24)$$

Da $\tilde{\Phi}$ zweimal partiell stetig differenzierbar ist, verschwindet das zweite Integral für $\epsilon \rightarrow 0$ und nach dem Zwischenwertsatz ergibt das erste

$$\int_{\partial \tilde{V}} d\vec{f}' \cdot \left[\Phi(\vec{x}') \vec{\nabla}' \frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x} - \vec{x}'|} \vec{\nabla}' \Phi(\vec{x}') \right] \stackrel{\epsilon \rightarrow 0}{=} 4\pi \Phi(\vec{x}). \quad (3.15.25)$$

Gleichsetzen mit der rechten Seite von (3.15.17) liefert schließlich das gewünschte Resultat

$$\Phi(\vec{x}) = -\gamma \lim_{\epsilon \rightarrow 0} \int_{\tilde{V}} d^3 x' \frac{\rho(\vec{x}')}{|\vec{x} - \vec{x}'|} = -\gamma \int_V d^3 x' \frac{\rho(\vec{x}')}{|\vec{x} - \vec{x}'|}, \quad (3.15.26)$$

wobei wir verwendet haben, dass $\rho(\vec{x}')$ nur im Inneren von V von 0 verschieden ist.

Der obige Beweis zeigt aber, dass dieses Resultat auch dann noch gilt, wenn $\rho(\vec{x}')$ im Unendlichen schnell genug abfällt, so dass die oben betrachteten Integrale über $\tilde{V}$ allesamt existieren.

In gewisser Weise haben wir eine Umkehroperation zum Laplaceoperator gefunden, denn wie wir gezeigt haben, folgt aus

$$\Delta \Phi_1(\vec{x}) = \Phi_2(\vec{x}), \quad (3.15.27)$$

dass

$$\Phi_1(\vec{x}) = - \int_{E^3} d^3 x' G(\vec{x}, \vec{x}') \Phi_2(\vec{x}') \quad \text{mit} \quad G(\vec{x}, \vec{x}') = \frac{1}{4\pi |\vec{x} - \vec{x}'|} \quad (3.15.28)$$

gilt, sofern nur $\Phi_2(\vec{x})$ hinreichend schnell verschwindet, so dass das Volumenintegral existiert. Dabei erfüllt die Lösung Φ_1 die **Randbedingung**, dass sie im Unendlichen verschwindet. Unter dieser Bedingung ist (3.15.28) eine eindeutige Lösung der Poisson-Gleichung (3.15.27). Man nennt in diesem Zusammenhang G eine **Green-Funktion des Differentialoperators** $(-\Delta)^a$.

^aGreen-Funktionen von linearen Differentialoperatoren werden Ihnen in allen Theorie-Vorlesungen während des gesamten Physikstudiums begegnen, insbesondere in der Elektrodynamik, Quantenmechanik und schließlich der Quantenfeldtheorie.

3.16 Der Helmholtzsche Zerlegungssatz der Vektoranalysis

In der Physik muss man oft ein Vektorfeld aus der Vorgabe seiner Divergenz („Quellen“) und seiner Rotation („Wirbel“) bestimmen.

Gegeben sei ein skalares Feld $\rho(\vec{x})$ und ein Vektorfeld $\vec{w}(\vec{x})$, und wir fragen nach der Existenz eines Vektorfeldes $\vec{V}$, so dass

$$\operatorname{div} \vec{V} = \rho, \quad \operatorname{rot} \vec{V} = \vec{w} \quad (3.16.1)$$

gilt. Es ist klar, dass für die (zumindest lokale) Existenz eines solchen Vektorfeldes $\vec{V}$ die Konsistenzbedingung

$$\operatorname{div} \vec{w} = 0 \quad (3.16.2)$$

erfüllt sein muss, d.h. $\vec{w}$ ist ein Solenoidalfeld, dessen Vektorpotential $\vec{V}$ sein soll.

Die Betrachtungen in den vorigen Abschnitten legen die Zerlegung des Vektorfeldes in zwei Anteile nahe:

$$\vec{V} = \vec{V}_1 + \vec{V}_2 \quad \text{mit} \quad \operatorname{rot} \vec{V}_1 = 0 \quad \text{und} \quad \operatorname{div} \vec{V}_2 = 0, \quad (3.16.3)$$

d.h. wir spalten $\vec{V}$ in ein Potentialfeld $\vec{V}_1$ und ein Solenoidalfeld $\vec{V}_2$ auf, so dass

$$\operatorname{div} \vec{V} = \operatorname{div} \vec{V}_1 = \rho, \quad \operatorname{rot} \vec{V}_1 = 0, \quad (3.16.4)$$

$$\operatorname{rot} \vec{V} = \operatorname{rot} \vec{V}_2 = \vec{w}, \quad \operatorname{div} \vec{V}_2 = 0 \quad (3.16.5)$$

gilt (vgl. (3.16.1)). Im Folgenden zeigen wir, dass eine solche Zerlegung existiert, vorausgesetzt es sind bestimmte Glattheitsbedingungen und ein hinreichend schnelles Abfallen der Quellen ρ und $\vec{w}$ im Unendlichen erfüllt. Wir kommen auf die genauen Bedingungen weiter unten noch zurück.

3.16.1 Bestimmung des Potentialfeldanteils

Wir beginnen mit der Aufgabe, $\vec{V}_1$ zu bestimmen. Wegen $\operatorname{rot} \vec{V}_1 = 0$ existiert nach dem Poincaréschen Lemma (zumindest lokal) ein skalares Feld Φ , so dass

$$\vec{V}_1 = -\operatorname{grad} \Phi \quad (3.16.6)$$

gilt.

Setzen wir dies in (3.16.4) ein, ergibt sich

$$\operatorname{div} \vec{V}_1 = -\operatorname{div} \operatorname{grad} \Phi = -\Delta \Phi = \rho. \quad (3.16.7)$$

Dies bezeichnet man als die inhomogene Potentialgleichung oder auch als **Poisson-Gleichung**. Diese haben wir bereits im vorigen Abschnitt gelöst. Verwenden wir also (3.15.27) und (3.15.28), ergibt sich bereits die Lösung

$$\Phi(\vec{x}) = \int_{E^3} d^3x' \frac{\rho(\vec{x}')}{4\pi|\vec{x} - \vec{x}'|}. \quad (3.16.8)$$

Der Potentialfeldanteil selbst berechnet sich durch Gradientenbildung gemäß unseres Ansatzes (3.16.6):

$$\vec{V}_1(\vec{x}) = -\operatorname{grad} \Phi(\vec{x}) = \int_{E^3} d^3x' \frac{\rho(\vec{x}')}{4\pi} \frac{\vec{x} - \vec{x}'}{|\vec{x} - \vec{x}'|^3}. \quad (3.16.9)$$

Betrachten wir (3.16.8), muss offenbar für die Existenz des Integrals über den ganzen Raum $\rho(\vec{x}')$ schneller als $1/|\vec{x}'|^2$ für $|\vec{x}'| \rightarrow \infty$ abfallen.

Andererseits ist Φ aber ohnehin nur bis auf eine additive Konstante definiert, und wir können statt (3.16.8)

$$\tilde{\Phi}(\vec{x}) = \int_{E^3} d^3x' \frac{\rho(\vec{x}')}{4\pi} \left(\frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x}_0 - \vec{x}'|} \right) \quad (3.16.10)$$

als Lösung für die Poisson-Gleichung (3.16.7) verwenden. Dabei ist $\vec{x}_0$ irgendein geeigneter fixierter Punkt. Da der Ausdruck in der Klammer für $|\vec{x}'| \rightarrow \infty$ nun wie $1/|\vec{x}'|^2$ verschwindet, ist es für die Konvergenz des Integrals im Unendlichen hinreichend, wenn $\rho(\vec{r}')$ schneller als $1/|\vec{x}'|$ verschwindet.

3.16.2 Bestimmung des Solenoidalfeldanteils

Wir müssen nun noch den Solenoidalfeldanteil des Ausgangsvektorfeldes bestimmen, d.h. wir haben (3.16.5) zu lösen. Da $\text{div } \vec{V}_2 = 0$ ist, existiert aufgrund unserer Betrachtungen in Abschnitt 3.14 ein **Vektorpotential** $\vec{A}$, so dass

$$\vec{V}_2 = \text{rot } \vec{A} \quad (3.16.11)$$

ist. Wir wie ebenfalls in Abschnitt 3.14 erläutert haben, ist $\vec{A}$ nur bis auf ein Potentialvektorfeld bestimmt, so dass wir eine Nebenbedingung fordern können. Diesmal wird sich die folgende als **Coulomb-Bedingung**¹⁰

$$\text{div } \vec{A} = 0 \quad (3.16.12)$$

als besonders bequem erweisen. Setzen wir nämlich (3.16.11) in die erste Gleichung von (3.16.5) ein, finden wir für diese Nebenbedingung (*nachrechnen!*)

$$\text{rot rot } \vec{A} = \vec{\nabla} \times (\vec{\nabla} \times \vec{A}) = \vec{\nabla}(\vec{\nabla} \cdot \vec{A}) - (\vec{\nabla} \cdot \vec{\nabla})\vec{A} = \text{grad div } \vec{A} - \Delta \vec{A} = -\Delta \vec{A} \stackrel{!}{=} \vec{w}. \quad (3.16.13)$$

Es sei auch an dieser Stelle nochmals betont, dass diese Gleichung **nur auf kartesische Komponenten** des Vektorfeldes $\vec{A}$ angewendet werden darf, da nur für kartesische Orthonormalsysteme die Basisvektoren ortsunabhängig sind.

Wir werden also wieder auf **Poisson-Gleichungen** für die **kartesischen** Komponenten von $\vec{A}$ geführt, so dass wir wieder die Green-Funktion des Laplace-Operators anwenden können:

$$\vec{A}(\vec{x}) = \int_{E^3} d^3x' \frac{\vec{w}(\vec{x}')}{4\pi|\vec{x} - \vec{x}'|} \quad (3.16.14)$$

bzw.

$$\vec{A}(\vec{x}) = \int_{E^3} d^3x' \frac{\vec{w}(\vec{x}')}{4\pi} \left(\frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x}_0 - \vec{x}'|} \right). \quad (3.16.15)$$

Für die letztere Form der Lösung genügt es für die Konvergenz des Integrals wieder, wenn $\vec{w}(\vec{x}')$ schneller als $1/|\vec{x}'|$ für $|\vec{x}'| \rightarrow \infty$ verschwindet.

Allerdings müssen wir uns nun noch vergewissern, dass diese Lösung auch tatsächlich die Coulomb-Bedingung (3.16.12) erfüllt, die wir ja verlangt haben, um für $\vec{A}$ eine Poisson-Gleichung zu erhalten. In der Tat gilt

$$\text{div } \vec{A}(\vec{x}) = \int_{E^3} d^3x' \frac{1}{4\pi} \vec{w}(\vec{x}') \vec{\nabla}_x \frac{1}{|\vec{x} - \vec{x}'|} = \int_{E^3} d^3x' \frac{1}{4\pi} \vec{w}(\vec{x}') \left(-\vec{\nabla}_{x'} \frac{1}{|\vec{x} - \vec{x}'|} \right). \quad (3.16.16)$$

¹⁰Diese Bezeichnung kommt aus der Elektrodynamik, wo statische Magnetfelder ein Vektorpotential besitzen.

Wegen der Konsistenzbedingung (3.16.2), also $\operatorname{div} \vec{w} = 0$, können wir dies mit Hilfe des Gaußschen Satzes in ein Oberflächenintegral umwandeln. Da der Rand des gesamten Raumes E^3 aber im „Unendlichen“ liegt, können wir unter der Annahme, dass $\vec{w}$ im Unendlichen hinreichend schnell verschwindet, schließen, dass auch dieses Oberflächenintegral verschwindet: Nach der Produktregel gilt

$$\operatorname{div}_{x'} \left[\frac{\vec{w}(\vec{x}')}{|\vec{x} - \vec{x}'|} \right] = [\operatorname{div}_{x'} \vec{w}(\vec{x}')] \frac{1}{|\vec{x} - \vec{x}'|} + \vec{w}(\vec{x}') \cdot \operatorname{grad}_{x'} \frac{1}{|\vec{x} - \vec{x}'|} \stackrel{(3.16.2)}{=} \vec{w}(\vec{x}') \cdot \operatorname{grad}_{x'} \frac{1}{|\vec{x} - \vec{x}'|}. \quad (3.16.17)$$

Wir haben also nach dem Gaußschen Satz

$$\operatorname{div} \vec{A}(\vec{x}) = -\frac{1}{4\pi} \int_{E^3} d^3 x' \operatorname{div}_{x'} \left[\vec{w}(\vec{x}') \frac{1}{|\vec{x} - \vec{x}'|} \right] = \lim_{R \rightarrow \infty} \frac{1}{4\pi} \int_{\partial K_R} d^2 \vec{f}(\vec{x}') \frac{\vec{w}(\vec{x}')}{|\vec{x} - \vec{x}'|} = 0, \quad (3.16.18)$$

falls $w(\vec{x}')$ im Unendlichen schneller als $\mathcal{O}(|\vec{x}'|^{-1})$ verschwindet. Also ist die Coulomb-Eichbedingung erfüllt und somit (3.16.14) bzw. (3.16.15) tatsächlich eine Lösung für das Vektorpotential.

Für den Solenoidalfeldanteil finden wir schließlich

$$\begin{aligned} \vec{V}_2(\vec{x}) &= \operatorname{rot} \vec{A}(\vec{x}) = \int_{E^3} d^3 x' \vec{\nabla}_x \times \frac{\vec{w}(\vec{x}')}{4\pi|\vec{x} - \vec{x}'|} \\ &= - \int_{E^3} d^3 x' \vec{w}(\vec{x}') \times \vec{\nabla}' \frac{1}{4\pi|\vec{x} - \vec{x}'|} = \int_{E^3} d^3 x' \vec{w}(\vec{x}') \times \frac{\vec{x} - \vec{x}'}{4\pi|\vec{x} - \vec{x}'|^3}. \end{aligned} \quad (3.16.19)$$

Damit haben wir den **Helmholtzschen Zerlegungssatz** für Vektorfelder bewiesen. Wir fassen ihn noch einmal übersichtlich zusammen:

Seien $\rho(\vec{x})$ und $\vec{w}(\vec{x})$ im E^3 definierte Felder, die im Unendlichen schneller als mit der Ordnung $\mathcal{O}(|\vec{x}|^{-1})$ abfallen und erfülle $\vec{w}$ die Bedingung $\operatorname{div} \vec{w} = 0$. Dann lässt sich ein gegebenes Vektorfeld $\vec{V}$ eindeutig in einen **Potentialfeldanteil** und einen **Solenoidalfeldanteil** zerlegen, so dass

$$\vec{V} = \vec{V}_1 + \vec{V}_2, \quad \operatorname{div} \vec{V} = -\rho, \quad \operatorname{rot} \vec{V} = \vec{w}, \quad \text{wobei} \quad (3.16.20)$$

$$\vec{V}_1 = -\operatorname{grad} \Phi \quad \text{mit} \quad \Phi(\vec{x}) = \int_{E^3} d^3 x' \frac{\rho(\vec{x}')}{4\pi} \left(\frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x}_0 - \vec{x}'|} \right), \quad (3.16.21)$$

$$\vec{V}_2 = \operatorname{rot} \vec{A} \quad \text{mit} \quad \vec{A}(\vec{x}) = \int_{E^3} d^3 x' \frac{\vec{w}(\vec{x}')}{4\pi} \left(\frac{1}{|\vec{x} - \vec{x}'|} - \frac{1}{|\vec{x}_0 - \vec{x}'|} \right). \quad (3.16.22)$$

Die Felder selbst berechnen sich aus der **Quelle** ρ und der **Wirbelstärke** $\vec{w}$ des Vektorfeldes zu

$$\vec{V}_1(\vec{x}) = \int_{E^3} d^3 x' \frac{\rho(\vec{x}')}{4\pi} \frac{\vec{x} - \vec{x}'}{|\vec{x} - \vec{x}'|^3}, \quad \vec{V}_2(\vec{x}) = \int_{E^3} d^3 x' \vec{w}(\vec{x}') \times \frac{\vec{x} - \vec{x}'}{4\pi|\vec{x} - \vec{x}'|^3}. \quad (3.16.23)$$

3.17 Verallgemeinerte Funktionen (Distributionen)

In den vorigen Abschnitten haben wir die Green-Funktion für den Operator $(-\Delta)$ als

$$G(\vec{x}) = \frac{1}{4\pi r}, \quad r = |\vec{x}| \quad (3.17.1)$$

definiert. Diese Funktion ist natürlich bei $\vec{x} = \vec{0}$ singular, und es gilt

$$\Delta G(\vec{x}) = 0, \quad \vec{x} \neq 0. \quad (3.17.2)$$

Wir haben aber auch gesehen, dass die **Poisson-Gleichung**

$$\Delta \Phi(\vec{x}) = -\rho(\vec{x}) \quad (3.17.3)$$

durch

$$\Phi(\vec{x}) = \int_{\mathbb{R}^3} d^3 x' G(\vec{x} - \vec{x}') \rho(\vec{x}') = \int_{\mathbb{R}^3} d^3 x' \frac{\rho(\vec{x}')}{4\pi |\vec{x} - \vec{x}'|} \quad (3.17.4)$$

gelöst wird, wobei wir voraussetzen, dass das Integral existiert. Wenden wir hierauf den Laplace-Operator an und vertauschen naiv die Ableitungen nach den Komponenten von $\vec{x}$ mit der Integration bzgl. $\vec{x}'$ erhalten wir

$$\Delta \Phi(\vec{x}) = \int_{\mathbb{R}^3} d^3 x' \Delta G(\vec{x} - \vec{x}') \rho(\vec{x}'). \quad (3.17.5)$$

Dies scheint nun (3.17.3) zu widersprechen, denn gemäß (3.17.2) ist

$$\Delta G(\vec{x} - \vec{x}') = 0 \quad \text{für} \quad \vec{x} \neq \vec{x}'. \quad (3.17.6)$$

Da also außer für den einzigen Punkt $\vec{x}' = \vec{x}$ der Integrand in (3.17.5) verschwindet, erhalten wir mit dieser Rechnung 0 im Widerspruch zu (3.17.3). Dies liegt natürlich daran, dass die Vertauschung von Ableitungen und Integration in diesem Fall mathematisch *nicht* gerechtfertigt werden kann. Oben haben wir uns dieses Problems entledigt, indem wir zunächst eine kleine Kugel vom Radius ϵ um die singuläre Stelle $\vec{x}' = \vec{x}$ herausgeschnitten haben und mit Hilfe des Greenschen Integralsatzes bestätigt haben, dass (3.17.4) in der Tat die Poisson-Gleichung (3.17.3) löst.

Wir können dieses Problem aber auch auf einfachere Weise lösen, indem wir die Singularität der Green-Funktion (*3.17.1) bei $r = 0$ zunächst **regularisieren**. Als für unsere Zwecke bequem erweist sich die Regularisierung

$$G_\epsilon(\vec{x}) = \frac{1}{4\pi\sqrt{r^2 + \epsilon^2}}, \quad \epsilon > 0. \quad (3.17.7)$$

Dann ist die Funktion in für alle $\vec{x} \in \mathbb{R}^3$ definiert und beliebig oft partiell integrierbar. Offenbar gilt (*Nachrechnen!*)

$$\vec{\nabla} G_\epsilon(\vec{x}) = (\vec{\nabla} r) \frac{d}{dr} \left(\frac{1}{4\pi\sqrt{r^2 + \epsilon^2}} \right) = -\frac{1}{4\pi} \frac{\vec{x}}{r} \frac{r}{(r^2 + \epsilon^2)^{3/2}} = -\frac{1}{4\pi(r^2 + \epsilon^2)^{3/2}} \vec{x} \quad (3.17.8)$$

und damit

$$\begin{aligned} -\Delta G_\epsilon(\vec{x}) &= \vec{\nabla} \cdot \vec{\nabla} G_\epsilon(\vec{x}) = \vec{x} \cdot \vec{\nabla} \left(\frac{1}{4\pi(r^2 + \epsilon^2)^{3/2}} \right) + \frac{1}{4\pi(r^2 + \epsilon^2)^{3/2}} \vec{\nabla} \cdot \vec{x} \\ &= \vec{x} \cdot \left(-\frac{3}{2} \frac{2r}{(r^2 + \epsilon^2)^{5/2}} \frac{\vec{x}}{r} \right) + \frac{3}{4\pi(r^2 + \epsilon^2)^{3/2}} \\ &= \frac{3}{4\pi(r^2 + \epsilon^2)^{3/2}} + \frac{3r^2}{4\pi(r^2 + \epsilon^2)^{5/2}} \\ &\Rightarrow -\Delta G_\epsilon(\vec{x}) = \frac{3\epsilon^2}{4\pi(r^2 + \epsilon^2)^{5/2}}. \end{aligned} \quad (3.17.9)$$

Wie zu erwarten, ergibt sich nun

$$-\lim_{\epsilon \rightarrow 0} \Delta G_\epsilon(\vec{x}) = \begin{cases} 0 & \text{falls } \vec{x} \neq \vec{0}, \\ \infty & \text{falls } \vec{x} = \vec{0}. \end{cases} \quad (3.17.10)$$

3. Vektoranalysis

Allerdings benötigen wir die regularisierte Green-Funktion in (3.17.5) d.h. zur Definition dieses Integrals als

$$\Delta\Phi(\vec{x}) = \lim_{\epsilon \rightarrow 0} \int_{\mathbb{R}^3} d^3x' \Delta G_\epsilon(\vec{x} - \vec{x}') \rho(\vec{x}') = - \lim_{\epsilon \rightarrow 0} \int_{\mathbb{R}^3} d^3x' \frac{3\epsilon^2}{4\pi[(\vec{x} - \vec{x}')^2 + \epsilon^2]^{5/2}} \rho(\vec{x}'). \quad (3.17.11)$$

In diesem Integral substituieren wir $\vec{r} = (\vec{x}' - \vec{x})/\epsilon$, $d^3x' = \epsilon^3 d^3r$:

$$\Delta\Phi(\vec{x}) = - \lim_{\epsilon \rightarrow 0} \int_{\mathbb{R}^3} d^3r \frac{3}{4\pi(r^2 + 1)^{5/2}} \rho(\vec{x} + \epsilon\vec{r}) = -\rho(\vec{x}) \int_{\mathbb{R}^3} d^3r \frac{3}{4\pi(r^2 + 1)^{5/2}}. \quad (3.17.12)$$

Für das verbliebene Integral führen wir zunächst Kugelkoordinaten (r, ϑ, φ) ein. Die Winkelintegration ergibt dann einfach einen Faktor 4π und das verbliebene Integral ist

$$\Delta\Phi(\vec{x}) = -\rho(\vec{x}) \int_0^\infty dr \frac{3r^2}{(r^2 + 1)^{5/2}}. \quad (3.17.13)$$

Hierin substituieren wir wiederum $r = \sinh \lambda$, $dr = d\lambda \cosh \lambda$, was schließlich auf das zu erwartende Resultat

$$\Delta\Phi(\vec{x}) = -\rho(\vec{x}) \int_0^\infty d\lambda \frac{3 \sin^2 \lambda}{\cosh^4 \lambda} = -\rho(\vec{x}) \int_0^\infty d\lambda 3 \frac{\tanh^2 \lambda}{\cosh^2 \lambda} = -\rho(\vec{x}) \tanh \lambda \Big|_{\lambda=0}^\infty = -\rho(\vec{x}). \quad (3.17.14)$$

Es ist also in der Tat

$$\Delta\Phi(\vec{x}) = \lim_{\epsilon \rightarrow 0} \int_{\mathbb{R}^3} d^3x' \Delta G(\vec{x} - \vec{x}') \rho(\vec{x}') = -\rho(\vec{x}). \quad (3.17.15)$$

für alle hinreichend glatten und schnell im Unendlichen abfallenden Funktionen $\rho(\vec{x})$. Hier darf man natürlich nicht im üblichen Sinne die Limesbildung mit der Integration vertauschen. Man muss vielmehr zuerst über die glatte „Testfunktion“ $\rho(\vec{x}')$ integrieren und dann den Grenzwert ausführen.

Wir können nun aber auch einen anderen Limes-Begriff einführen, den sog. **schwachen Limes** (engl. weak limit):

$$\text{w-lim}_{\epsilon \rightarrow 0} [-\Delta G_\epsilon(\vec{x} - \vec{x}')] = \delta^{(3)}(\vec{x} - \vec{x}'). \quad (3.17.16)$$

Darunter ist dann die oben beschriebene Prozedur „erst über eine Testfunktion integrieren und dann den Limes bilden“ zu verstehen. Führt man dann formal den Grenzwert unter dem Integral aus, erhält man in diesem Sinne die Vorschrift

$$\begin{aligned} - \lim_{\epsilon \rightarrow 0} \int_{\mathbb{R}^3} d^3x' \Delta G_\epsilon(\vec{x} - \vec{x}') \rho(\vec{x}') &= \int_{\mathbb{R}^3} d^3x' \rho(\vec{x}') \text{w-lim}_{\epsilon \rightarrow 0} [-\Delta G_\epsilon(\vec{x} - \vec{x}')] \\ &= \int_{\mathbb{R}^3} d^3x' \rho(\vec{x}') \delta^{(3)}(\vec{x} - \vec{x}') = \rho(\vec{x}). \end{aligned} \quad (3.17.17)$$

Man nennt das Symbol $\delta^{(3)}(\vec{x} - \vec{x}')$ eine **verallgemeinerte Funktion** oder auch **Distribution**. Sie ist durch die Vorschrift definiert, dass für alle **Testfunktionen** ρ gilt

$$\int_{\mathbb{R}^3} d^3x' \rho(\vec{x}') \delta^{(3)}(\vec{x} - \vec{x}') = \rho(\vec{x}). \quad (3.17.18)$$

Anschaulich ergibt sich aus (3.17.9),

$$\delta^{(3)}(\vec{x}) = \text{w-lim} \frac{3\epsilon^2}{4\pi(r^2 + \epsilon^2)^{5/2}}. \quad (3.17.19)$$

Nimmt man den Limes im „naiven Sinne“ ist

$$\delta^{(3)}(\vec{x}) = \begin{cases} 0 & \text{für } \vec{x} \neq 0, \\ \infty & \text{für } \vec{x} = 0, \end{cases} \quad (3.17.20)$$

so dass

$$\int_V d^3x \delta^{(3)}(\vec{x}) = \begin{cases} 0 & \text{falls } \vec{0} \notin V, \\ 1 & \text{falls } \vec{0} \in V. \end{cases} \quad (3.17.21)$$

Diese Idee wurde insbesondere durch Dirac in seiner Formulierung der Quantenmechanik eingeführt (Paul Adrien Maurice Dirac, 1902-1984) aber nicht mathematisch rigoros begründet. Aufgrund der Nützlichkeit des Konzepts entwickelten daraus verschiedene Mathematiker die moderne **Funktionalanalysis**. Wir werden die Diracsche δ -Distribution im Folgenden im nichtrigorösen Sinne weiterverwenden.

Die δ -Distribution lässt sich natürlich auch auf d -dimensionale Integrale verallgemeinern. Insbesondere gilt für $d = 1$

$$\int_{\mathbb{R}} dx' \delta(x - x') f(x') = f(x) \quad (3.17.22)$$

für alle hinreichend oft stetig differenzierbaren im Unendlichen schnell abfallende Funktionen.

Man kann dann auch weitere Distributionen einführen. Wichtig ist z.B. die **Heavisidesche Einheitssprungfunktion** (Oliver Heaviside, 1850-1925)

$$\Theta(x) = \begin{cases} 1 & \text{für } x > 0, \\ 0 & \text{für } x < 0. \end{cases} \quad (3.17.23)$$

Dann rechnet man formal, als hätte man es mit einer überall differenzierbaren (!) Funktion zu tun, wobei man allerdings im Sinne von Distributionen über eine Testfunktion integriert. Dann erhält man für $\Theta'(x)$

$$\int_{\mathbb{R}} dx \Theta'(x) f(x) = - \int_{\mathbb{R}} \Theta(x) f'(x), \quad (3.17.24)$$

d.h. wir haben einfach partiell integriert und dabei verwendet, dass für die Testfunktion $f'(x) \rightarrow 0$ für $x \rightarrow \pm\infty$ sein soll. Dann ergibt sich aber, da ja $f'(x)$ auch als stetig vorausgesetzt wird

$$\int_{\mathbb{R}} dx \Theta'(x) f(x) = - \int_0^\infty f'(x) = -f(x)|_{x=0}^\infty = f(0). \quad (3.17.25)$$

Daraus folgt die im Sinne von Distributionen gültige Ableitungsregel

$$\Theta'(x) = \delta(x). \quad (3.17.26)$$

Dies lässt sich freilich auch wieder im Sinne eines schwachen Limes' definieren. Dazu nähern wir die Sprungfunktion durch eine geeignete stetig differenzierbare Funktion an

$$\Theta_\epsilon(x) = \frac{1}{2} [1 + \operatorname{artanh}(x/\epsilon)], \quad \Theta(x) = \lim_{\epsilon \rightarrow +0} \Theta_\epsilon(x). \quad (3.17.27)$$

In der Tat gilt ja

$$\lim_{y \rightarrow \pm\infty} \operatorname{artanh}(y) = \pm 1, \quad (3.17.28)$$

d.h. für $x \neq 0$

$$\lim_{\epsilon \rightarrow 0^+} \operatorname{artanh}(x/\epsilon) = \lim_{y \rightarrow \infty \operatorname{sign} x} = \operatorname{sign} x. \quad (3.17.29)$$

Für $x = 0$ folgt mit der Regularisierung (3.17.27) für diesen zunächst unbestimmten Punkt $\Theta(0) = 1/2$. Um nun (3.17.26) zu beweisen, bemerken wir, dass

$$\delta_\epsilon(x) = \Theta'_\epsilon(x) = \frac{1}{2\epsilon(1+x^2/\epsilon^2)} = \frac{\epsilon}{2(x^2 + \epsilon^2)} \quad (3.17.30)$$

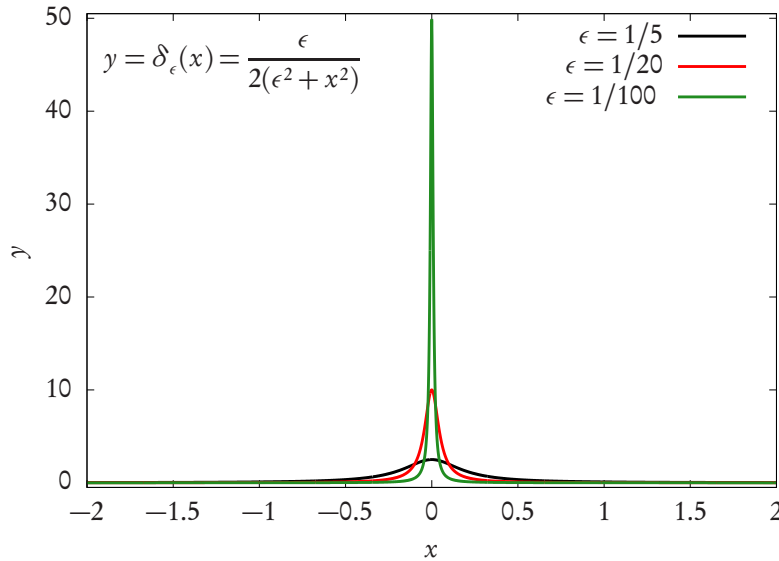
ist. Für eine beliebige Testfunktion ist dann mit der Substitution $x = \epsilon y$, $dx = dy\epsilon$

$$\begin{aligned} \lim_{\epsilon \rightarrow 0^+} \int_{\mathbb{R}} dx f(x) \Theta'_\epsilon(x) &= \lim_{\epsilon \rightarrow 0^+} \int_{\mathbb{R}} dx f(x) \frac{\epsilon}{2(x^2 + \epsilon^2)} = \lim_{\epsilon \rightarrow 0^+} \int_{\mathbb{R}} dy f(\epsilon y) \frac{1}{2(1+y^2)} \\ &= \frac{1}{2} f(0) \int_{\mathbb{R}} dy \frac{1}{1+y^2} \\ &= \frac{1}{2} f(0) \operatorname{artanh} y \Big|_{y \rightarrow -\infty}^{y \rightarrow \infty} = f(0). \end{aligned} \quad (3.17.31)$$

Dies beweist auch formal, dass

$$\Theta'(x) = \text{w-lim}_{\epsilon \rightarrow 0^+} \Theta'_\epsilon(x) = \delta(x). \quad (3.17.32)$$

Plottet man $\delta_\epsilon(x)$ (3.17.30) für verschiedene ϵ , erkennt man, dass sich in der Tat für $\epsilon \rightarrow 0^+$ immer schärfer um $x = 0$ gepeakte Funktionen ergibt, für die aber die Fläche unter dem Graphen stets 1 ist, was genau der anschaulichen Definition der Diracschen δ -Distribution entspricht.



3.18 Transporttheoreme

Manchmal benötigt man die Ableitungen von Weg-, Flächen- und Volumenintegralen, für die die Integrationsbereiche und die zu integrierenden Felder von einem Parameter abhängen, nach diesem Parameter. Dabei ist der Parameter oft die Zeit. Die anschaulichsten Beispiele stammen aus der Hydrodynamik, wo man die Strömung einer Flüssigkeit oder eines Gases beschreibt, indem man die physikalisch interessanten Größen wie Dichte oder Geschwindigkeit an jedem Ort $\vec{x}$ als Funktion der Zeit angibt. Entsprechend lassen sich Dichten von Energie, Impuls und Drehimpuls definieren. Will man daraus die Zeitentwicklung dieser Größen für einen endlich ausgedehnten Teil des Fluids bestimmen, benötigt man die Beschreibung eines mit dem gegebenen Teil des Fluids bewegten Volumens und dessen Randes. Bildet man dann die Zeitableitung dieser

„integralen Größe“ ist die entsprechende zeitliche Änderung des Volumens und dessen Randes zu berücksichtigen. Da hierbei der Transport von Größen wie Energie und Impuls beschrieben wird, nennt man die entsprechenden Formeln für die zeitliche Ableitung solcher über Integrale definierten Größen **Transporttheoreme**. Insbesondere heißt die Version für Volumenintegrale **Reynoldssches Transporttheorem** (Osborne Reynolds 1842-1912).

3.18.1 Transporttheorem für Wegintegrale

Wir betrachten das Wegintegral eines Vektorfeldes $\vec{W}(t, \vec{x})$ entlang eines von der Zeit t abhängigen Weges C . Diesen Weg parametrisieren wir mittels eines Parameters q als Funktion der Zeit

$$\vec{x} = \vec{\xi}(t, q), \quad q \in [a, b], \quad (3.18.1)$$

wobei wir die Parametrisierung so wählen, dass das Intervall $[a, b]$ unabhängig von der Zeit ist. Dann definieren wir das Wegintegral

$$I(t) = \int_C d\vec{x} \cdot \vec{W}(t, \vec{x}) = \int_a^b dq \partial_q \vec{\xi}(t, q) \cdot \vec{W}[t, \vec{\xi}(t, q)]. \quad (3.18.2)$$

Die Zeitableitung können wir nun berechnen, indem wir Integration und Ableitung vertauschen, wobei wir voraussetzen, dass alle Funktion unter dem Integral stetig differenzierbar sind. Mit der Produkt- und Kettenregel folgt

$$\begin{aligned} \frac{d}{dt} \left\{ \partial_q \vec{\xi}(t, q) \cdot \vec{W}[t, \vec{\xi}(t, q)] \right\} &= \partial_t \partial_q \vec{\xi}(t, q) \cdot \vec{W}[t, \vec{\xi}(t, q)] \\ &+ \partial_q \vec{\xi}(t, q) \cdot [\partial_t \vec{W}(t, \vec{x}) + (\partial_t \vec{\xi}(t, q) \cdot \vec{\nabla}) \vec{W}(t, \vec{x})]_{\vec{x}=\vec{\xi}(t, q)}. \end{aligned} \quad (3.18.3)$$

Wir können nun davon ausgehen, dass die Abbildung (3.18.1) $q \mapsto \vec{x}$ für jedes t eindeutig ist und wir daher die Funktion zu jedem Zeitpunkt t

$$\vec{v}(t, q) = \partial_t \vec{\xi}(t, q) \quad (3.18.4)$$

auch als Funktion von $\vec{x}$ entlang der Kurve auffassen können. Dann gilt

$$\vec{v}(t, q) \equiv \vec{v}(t, \vec{x})|_{\vec{x}=\vec{\xi}(t, q)}. \quad (3.18.5)$$

Damit folgt

$$\partial_t \partial_q \vec{\xi}(t, q) = \partial_q \partial_t \vec{\xi}(t, q) = \partial_q \vec{v}(t, q) = \partial_q \vec{\xi} \cdot \vec{\nabla} \vec{v}(t, \vec{x})|_{\vec{x}=\vec{\xi}(t, q)}. \quad (3.18.6)$$

Setzen wir all dies in (3.18.3) ein und integrieren bzgl. q , folgt

$$\begin{aligned} \frac{d}{dt} I(t) &= \int_a^b dq \partial_q \vec{\xi}(t, q) \cdot [\partial_t \vec{W}(t, \vec{x}) + (\vec{v}(t, \vec{x}) \cdot \vec{\nabla}) \vec{W}(t, \vec{x})]_{\vec{x}=\vec{\xi}(t, q)} \\ &+ \int_a^b dq \vec{W}(t, \vec{x}) \cdot (\partial_q \vec{\xi}(t, q) \cdot \vec{\nabla}) \vec{v}(t, \vec{x})|_{\vec{x}=\vec{\xi}(t, q)} \end{aligned} \quad (3.18.7)$$

Dies können wir schließlich wieder als Wegintegral schreiben:

$$\frac{d}{dt} \int_C d\vec{x} \cdot \vec{W} = \int_C d\vec{x} \cdot [\partial_t \vec{W} + (\vec{v} \cdot \vec{\nabla}) \vec{W}] + \int_C [(d\vec{x} \cdot \vec{\nabla}) \vec{v}] \cdot \vec{W}. \quad (3.18.8)$$

3.18.2 Transporttheorem für Flächenintegrale

Wir betrachten nun ein Flächenintegral

$$I(t) = \int_A d^2 \vec{f} \cdot \vec{W}(t, \vec{x}) \quad (3.18.9)$$

für eine durch

$$\vec{x} = \vec{\xi}(t, q_1, q_2) \quad (3.18.10)$$

parametrisierte Fläche. Dabei sei $\Omega \subseteq \mathbb{R}^2$ der von t unabhängige Bereich der Parameter (q_1, q_2) . Es gilt

$$I(t) = \int_{\Omega} d^2 q (\partial_{q_1} \vec{\xi} \times \partial_{q_2} \vec{\xi}) \cdot \vec{W}(t, \vec{\xi}). \quad (3.18.11)$$

Die Zeitableitung erfolgt wieder durch Vertauschen mit dem Integral. Dabei benötigen wir

$$\frac{d}{dt} (\partial_{q_1} \vec{\xi} \times \partial_{q_2} \vec{\xi}) = \partial_{q_1} \vec{v} \times \partial_{q_2} \vec{\xi} + \partial_{q_1} \vec{\xi} \times \partial_{q_2} \vec{v}, \quad (3.18.12)$$

wobei $\vec{v} = \partial_t \vec{\xi}$. Fassen wir nun wieder, wie oben beim Wegintegral, $\vec{v}$ als Funktion von $\vec{x}$ auf, folgt

$$\partial_{q_j} \vec{v} = (\partial_{q_j} \vec{\xi} \cdot \vec{\nabla}) \vec{v}. \quad (3.18.13)$$

Dies in (3.18.12) eingesetzt liefert (mit $\partial_k = \partial / \partial x_k$)

$$\frac{d}{dt} (\partial_{q_1} \vec{\xi} \times \partial_{q_2} \vec{\xi})_i = \sum_{a,b,k=1}^3 \epsilon_{iab} (\partial_k v_a \partial_{q_1} \xi_k \partial_{q_2} \xi_b + \partial_k v_b \partial_{q_1} \xi_a \partial_{q_2} \xi_k). \quad (3.18.14)$$

Im zweiten Summanden vertauschen wir nun die Summationsindizes a und b und verwenden $\epsilon_{iab} = -\epsilon_{iba}$:

$$\frac{d}{dt} (\partial_{q_1} \vec{\xi} \times \partial_{q_2} \vec{\xi})_i = \sum_{a,b,k=1}^3 \epsilon_{iab} \partial_k v_a (\partial_{q_1} \xi_k \partial_{q_2} \xi_b - \partial_{q_1} \xi_b \partial_{q_2} \xi_k). \quad (3.18.15)$$

Nun ist

$$d^2 q (\partial_{q_1} \xi_k \partial_{q_2} \xi_b - \partial_{q_1} \xi_b \partial_{q_2} \xi_k) = \sum_{c=1}^3 \epsilon_{ckb} d^2 f_c \quad (3.18.16)$$

und damit

$$\begin{aligned} d^2 q \frac{d}{dt} (\partial_{q_1} \vec{\xi} \times \partial_{q_2} \vec{\xi})_i &= \sum_{a,b,c,k=1}^3 \partial_k v_a \epsilon_{iab} \epsilon_{ckb} d^2 f_c \\ &= \sum_{a,b,c,k=1}^3 \partial_k v_a (\delta_{ic} \delta_{ak} - \delta_{ik} \delta_{ac}) d^2 f_c \\ &= d^2 f_i \vec{\nabla} \cdot \vec{v} - d^2 f_c \partial_i v_c. \end{aligned} \quad (3.18.17)$$

Damit erhalten wir schließlich

$$\frac{d}{dt} I(t) = \int_A d^2 f_c \cdot [\partial_t \vec{W} + (\vec{v} \cdot \vec{\nabla}) \vec{W} + \vec{W} (\vec{\nabla} \cdot \vec{v}) - (\vec{W} \cdot \vec{\nabla}) \vec{v}]. \quad (3.18.18)$$

Wir können dies noch etwas vereinfachen, indem wir beachten, dass (*nachrechnen!*)

$$\vec{\nabla} \times (\vec{W} \times \vec{v}) = (\vec{v} \cdot \vec{\nabla}) \vec{W} + \vec{W} (\vec{\nabla} \cdot \vec{v}) - (\vec{W} \cdot \vec{\nabla}) \vec{v} - \vec{v} (\vec{\nabla} \cdot \vec{W}) \quad (3.18.19)$$

ist. Setzen wir dies in (3.18.18) ein, folgt nach einiger Rechnung (*Übung!*)

$$\frac{d}{dt}I(t) = \int_A d^2\vec{f} \cdot [\partial_t \vec{W} + \vec{v}(\vec{\nabla} \cdot \vec{W}) + \vec{\nabla} \times (\vec{W} \times \vec{v})]. \quad (3.18.20)$$

Wenden wir im letzten Term noch den Stokesschen Integralsatz an, folgt schließlich

$$\frac{d}{dt} \int_A d^2\vec{f} \cdot \vec{W} = \int_A d^2\vec{f} \cdot [\partial_t \vec{W} + \vec{v}(\vec{\nabla} \cdot \vec{W})] + \int_{\partial A} d\vec{x} \cdot (\vec{W} \times \vec{v}). \quad (3.18.21)$$

3.18.3 Reynoldssches Transporttheorem für Volumenintegrale

Es sei

$$I(t) = \int_V d^3x U(t, \vec{x}) \quad (3.18.22)$$

mit dem durch

$$\vec{x} = \vec{\xi}(t, q_1, q_2, q_3), \quad (q_1, q_2, q_3) \in \Omega \quad (3.18.23)$$

definierten zeitabhängigen Volumen V . Dabei sei $\Omega \subseteq \mathbb{R}^3$ wieder ein zeitunabhängiger Bereich für die Parameter $q = (q_1, q_2, q_3)$. Dann gilt

$$I(t) = \int_{\Omega} d^3q J U(t, \vec{\xi}). \quad (3.18.24)$$

Für die Zeitableitung unter dem Integral benötigen wir zunächst die Ableitung der Jacobi-Determinante

$$J = \det \left(\frac{\partial(\xi_1, \xi_2, \xi_3)}{\partial(q_1, q_2, q_3)} \right). \quad (3.18.25)$$

Dazu schreiben wir die Determinante zuerst mit Hilfe des Levi-Civita-Symbols aus, wobei wir die **Einstein-sche Summenkonvention** verwenden, d.h. über zwei doppelt in einem Ausdruck auftretende Indizes wird summiert

$$J = \epsilon_{j_1 j_2 j_3} J_{j_1 1} J_{j_2 2} J_{j_3 3} = \frac{1}{3!} \epsilon_{j_1 j_2 j_3} \epsilon_{k_1 k_2 k_3} J_{j_1 k_1} J_{j_2 k_2} J_{j_3 k_3} \quad \text{mit} \quad J_{ab} = \frac{\partial \xi_a}{\partial q_b}. \quad (3.18.26)$$

Dabei haben wir im letzten Schritt verwendet, dass durch Umordnen von Zeilen oder Spalten bis auf das Vorzeichen stets J herauskommt. Die Vorzeichen werden aber durch das zweite ϵ -Symbol wieder kompensiert. Außerdem wird über alle k_1, k_2 und k_3 gemäß der Summenkonvention summiert, d.h. man erhält $3!J$, was durch $1/3!$ kompensiert wird. Daraus folgt

$$\frac{d}{dt}J = \frac{1}{3!} \epsilon_{j_1 j_2 j_3} \epsilon_{k_1 k_2 k_3} [(\partial_{q_{j_1}} v_{j_1}) J_{j_2 k_2} J_{j_3 k_3} + J_{j_1 k_1} (\partial_{q_{j_2}} v_{j_2}) J_{j_3 k_3} + J_{j_1 k_1} J_{j_2 k_2} (\partial_{q_{j_3}} v_{j_3})]. \quad (3.18.27)$$

Die drei Summanden in der eckigen Klammer liefern beim Summieren offenbar jeweils den gleichen Beitrag, d.h.

$$\frac{d}{dt}J = \frac{1}{2!} \epsilon_{j_1 j_2 j_3} \epsilon_{k_1 k_2 k_3} (\partial_{q_{j_1}} v_{j_1}) J_{j_2 k_2} J_{j_3 k_3} \quad (3.18.28)$$

Nun ist

$$J'_{j_1 k_1} = \frac{1}{2} \epsilon_{j_1 j_2 j_3} \epsilon_{k_1 k_2 k_3} J_{j_2 k_2} J_{j_3 k_3} = J[(\hat{J}^{-1})^T]_{j_1 k_1}, \quad (3.18.29)$$

denn es ist wegen (3.18.26)

$$J_{j_1 k_1} J'_{j_2 k_2} = \frac{1}{2} \epsilon_{k_1 k_2 k_3} \epsilon_{l_1 l_2 l_3} J_{j_1 l_1} J_{j_2 l_2} J_{j_3 l_3} = \frac{1}{2} \epsilon_{k_1 k_2 k_3} \epsilon_{j_1 j_2 j_3} J = J \delta_{j_1 k_1}. \quad (3.18.30)$$

Fassen wir nun wieder die Geschwindigkeiten $\vec{v}$ als Funktionen von $\vec{x}$ auf, folgt aus (3.18.27)

$$\frac{d}{dt}J = J_{lk}J'_{ik}\partial_l v_i = J\partial_i v_i = J\operatorname{div}\vec{v}. \quad (3.18.31)$$

Leiten wir also (3.18.24) nach der Zeit ab, folgt

$$\begin{aligned} \frac{d}{dt}I &= \int_{\Omega} d^3q J[\partial_t U + (\vec{v} \cdot \vec{\nabla})U + U\vec{\nabla} \cdot \vec{v}] \\ &= \int_V d^3x [\partial_t U + (\vec{v} \cdot \vec{\nabla})U + U(\vec{\nabla} \cdot \vec{v})]. \end{aligned} \quad (3.18.32)$$

Nun ist aber (*nachrechnen!*)

$$\vec{v} \cdot \vec{\nabla} U + U(\vec{\nabla} \cdot \vec{v}) = \vec{\nabla} \cdot (\vec{v} U). \quad (3.18.33)$$

Setzen wir dies in (3.18.32) ein und verwenden den Gaußschen Integralsatz, folgt schließlich das **Reynoldssche Transporttheorem** für Volumenintegrale

$$\frac{d}{dt} \int_V d^3x U = \int_V d^3x \partial_t U + \int_{\partial V} d^2\vec{f} \cdot (\vec{v} U). \quad (3.18.34)$$

3.19 Allgemein kovariante Formulierung der Vektoranalysis

3.19.1 Transformationsverhalten von Vektoren und Tensoren

Im Folgenden wollen wir noch kurz auf die **allgemein kovariante Formulierung** der Vektoranalysis im E^3 eingehen. Damit ist gemeint, dass wir die Differentialoperatoren grad, div und rot als auch Weg-, Flächen- und Volumenintegrale durch Operationen mit beliebigen **generalisierten Koordinaten** formuliert werden können. Hierzu bedienen wir uns eines präzisierten **Ricci-Kalküls**, bei dem **obere und untere** Indizes für die Basisvektoren und die darauf bezogenen Komponenten von Vektoren oder Tensoren eingeführt werden. Dies erleichtert nämlich die Anwendung der Transformationen von einem Satz generalisierter Koordinaten zu einem anderen erheblich. Die generalisierten Koordinaten (q^α) erhalten daher nunmehr definitionsgemäß einen mit griechischen Buchstaben geschriebenen **oberen Index**. Dabei ist der Definitionsbereich der (q^α) irgendein offenes Gebiet $G \subseteq \mathbb{R}^3$. Die Abbildung $(q^\alpha) \in G \mapsto \vec{x}(q) \in E^3$ wird dabei i.a., nur einen offenen Teilbereich des E^3 umkehrbar eindeutig abbilden.

Als zu dieser Parametrisierung gehörige Basisvektoren verwenden wir die **Tangentenvektoren der Koordinatenlinien**, die durch einen unteren Index gekennzeichnet werden:

$$\vec{T}_\alpha = \frac{\partial \vec{x}}{\partial q^\alpha} = \partial_\alpha \vec{x}. \quad (3.19.1)$$

Die Transformation zu einem neuen Satz generalisierter Koordinaten ergibt dann

$$\vec{T}'_\alpha = (\partial'_\alpha q^\beta) \partial_\beta \vec{x} = U^\beta_\alpha \partial_\beta \vec{x} = U^\beta_\alpha \vec{T}_\beta \quad \text{mit} \quad U^\beta_\alpha = \partial'_\alpha q^\beta. \quad (3.19.2)$$

Die Komponenten eines beliebigen Vektors $\vec{V}$ tragen dann definitionsgemäß obere Indizes:

$$\vec{V} = V^\beta \vec{T}_\beta = V'^\alpha \vec{T}'_\alpha = U^\beta_\alpha V'^\alpha \vec{T}_\beta \Rightarrow V^\beta = U^\beta_\alpha V'^\alpha. \quad (3.19.3)$$

Umgekehrt folgt mit der inversen Matrix $\hat{T} = \hat{U}^{-1}$

$$T^\alpha_\beta = \partial_\beta q'^\alpha \quad (3.19.4)$$

aus (3.19.3)

$$V'^\alpha = T^\alpha_\beta V^\beta. \quad (3.19.5)$$

Wir bezeichnen irgendwelche Symbole mit einem unteren Index, die sich analog zum Transformationsverhalten wie die Basisvektoren gemäß (3.19.2) transformieren, als **kovariante Objekte** und solche mit einem oberen Index, die sich gemäß (3.19.5) transformieren, als **kontravariante Objekte**. Entsprechend heißen die Komponenten V^α des Vektors $\vec{V}$ genauer die **kontravarianten Komponenten** dieses Vektors bzgl. der Basis $\vec{T}_\alpha$.

Wir betrachten nun das Skalarprodukt zweier Vektoren. Offenbar gilt

$$\vec{V} \cdot \vec{W} = (V^\alpha \vec{T}_\alpha) \cdot (W^\beta \vec{T}_\beta) = V^\alpha W^\beta \vec{T}_\alpha \cdot \vec{T}_\beta = g_{\alpha\beta} V^\alpha W^\beta. \quad (3.19.6)$$

Dazu haben wir die **kovarianten Komponenten** der das Skalarprodukt definierenden symmetrischen und positiv definiten Bilinearform durch (bzw. der **Metrik**)

$$g_{\alpha\beta} = \vec{T}_\alpha \cdot \vec{T}_\beta = g_{\beta\alpha} \quad (3.19.7)$$

definiert. Es ist klar, dass sich diese Komponenten in der Tat kovariant transformieren, und zwar bzgl. jedes der beiden Indizes, denn wegen (3.19.2) ist

$$g'_{\alpha\beta} = \vec{T}'_\alpha \cdot \vec{T}'_\beta = U^\gamma_\alpha U^\delta_\beta \vec{T}_\gamma \cdot \vec{T}_\delta = U^\gamma_\alpha U^\delta_\beta g_{\gamma\delta}. \quad (3.19.8)$$

Entsprechend ist das Skalarprodukt zweier Vektoren unabhängig von der Wahl der Basis durch die jeweiligen Metrikkoeffizienten. In der Tat folgt unter Verwendung von (3.19.3) und (3.19.8)

$$g'_{\alpha\beta} V'^\alpha W'^\beta = U^\gamma_\alpha U^\delta_\beta g_{\gamma\delta} V'^\alpha V'^\beta = g_{\gamma\delta} V^\gamma W^\delta = \vec{V} \cdot \vec{W}, \quad (3.19.9)$$

wie es sein muss.

Da $\overleftrightarrow{g} = (g_{\alpha\beta})$ eine symmetrische positiv definite Matrix ist, kann man sie diagonalisieren, und alle drei Eigenvektoren sind positiv. Es ist also $g = \det \overleftrightarrow{g} > 0$ und somit $\overleftrightarrow{g}$ invertierbar. Die Komponenten der

3. Vektoranalysis

inversen Matrix bezeichnen wir mit $\overleftrightarrow{g} = (g^{\alpha\beta})$. In Matrixschreibweise lautet (3.19.8)

$$\overleftrightarrow{g}' = \hat{U}^T \overleftrightarrow{g} \hat{U}. \quad (3.19.10)$$

Damit ist

$$\overleftrightarrow{g}' = \overleftrightarrow{g}'^{-1} = \hat{U}^{-1} \overleftrightarrow{g}^{-1} (\hat{U}^{-1})^T = \hat{T} \overleftrightarrow{g} \hat{T}^T. \quad (3.19.11)$$

Im Ricci-Kalkül folgt

$$g'^{\alpha\beta} = T^\alpha_\gamma T^\beta_\delta g^{\gamma\delta}, \quad (3.19.12)$$

d.h. $g^{\alpha\beta}$ transformiert sich, bezogen auf jeden der beiden Indizes, tatsächlich wie ein kontravariantes Objekt. Allgemein erhält man ein solches Transformationsverhalten für die ko- bzw. kontravarianten Komponenten $T_{\alpha_1 \dots \alpha_n} = T(\vec{T}_{\alpha_1}, \dots, \vec{T}_{\alpha_n})$ von Multilinearformen $T : (E^3)^n \rightarrow \mathbb{R}$, d.h. Abbildungen, die n Vektoren aus E^3 in die reellen Zahlen abbilden und linear in allen Argumenten sind. Es ist unmittelbar klar, dass $T_{\alpha_1 \dots \alpha_n}$ sich bzgl. aller n Indizes wie kovariante Vektorkomponenten transformiert (*warum?*), d.h.

$$T'_{\alpha_1 \dots \alpha_n} = U^{\beta_1}_{\alpha_1} \dots U^{\beta_n}_{\alpha_n} T_{\beta_1 \dots \beta_n}. \quad (3.19.13)$$

Schließlich müssen wir uns noch mit der Berechnung des Vektorprodukts mittels der kontravarianten oder kovarianten Komponenten beliebiger Vektoren beschäftigen.

Wir können mit Hilfe der Metrikkomponenten $g_{\mu\nu}$ bzw. $g^{\mu\nu}$ kontravariante Indizes „nach unten“ bzw. kovariante Indizes „nach oben ziehen“. Die entsprechenden Komponenten transformieren sich dann auch gemäß der Indexstellung ko- bzw. kontravariant.

Insbesondere können wir auch die zu $\vec{T}_\alpha$ gehörige **duale Basis**, die sich kontravariant transformiert via

$$\vec{T}^\alpha = g^{\alpha\beta} \vec{T}_\beta \quad (3.19.14)$$

definieren.

Dann kann man einen Vektor sowohl über seine kontravarianten als auch seine kovarianten Indizes ausdrücken:

$$\vec{V} = V^\alpha \vec{T}_\alpha = V_\beta \vec{T}^\beta = g_{\beta\alpha} V^\alpha \vec{T}^\beta = g_{\alpha\beta} V^\alpha \vec{T}^\beta. \quad (3.19.15)$$

Nun gilt

$$\vec{T}_\alpha \cdot \vec{T}^\beta = g^{\beta\gamma} \vec{T}_\alpha \cdot \vec{T}_\gamma = g^{\beta\gamma} g_{\alpha\gamma} = \delta_\alpha^\beta. \quad (3.19.16)$$

Dabei definieren wir die Kronecker-Symbole unabhängig von der Stellung der Indizes durch

$$\delta_{\alpha\beta} = \delta^{\alpha\beta} = \delta_\alpha^\beta = \begin{cases} 1 & \text{für } \alpha = \beta, \\ 0 & \text{für } \alpha \neq \beta. \end{cases} \quad (3.19.17)$$

Im Zusammenhang mit dem allgemein kovarianten Tensorkalkül definieren wir das bekannte Levi-Civita-Symbol als $\epsilon_{\alpha\beta\gamma} = \epsilon^{\alpha\beta\gamma}$. Dabei ist $\epsilon_{123} = \epsilon^{123} = 1$. Weiter ist es antisymmetrisch unter Vertauschen beliebiger Indexpaare. Insbesondere wird es 0, wenn wenigstens zwei Indizes gleich sind.

Wir können nun die duale Basis auch anders definieren. Dazu bemerken wir, dass für eine gegebene kovariante Basis $\vec{T}_\alpha$ die dazugehörigen dualen Basisvektoren $\vec{T}^\beta$ offenbar eindeutig durch (3.19.16) charakterisiert sind. Wir gehen im Folgenden davon aus, dass die Basisvektoren **rechtshändig** sind, d.h.

$$\vec{T}_1 \cdot (\vec{T}_2 \times \vec{T}_3) > 0. \quad (3.19.18)$$

Weiter bezeichnen wir mit lateinischen Indizes Vektor- und Tensorkomponenten bzgl. einer beliebigen kartesischen Basis $\vec{e}'_j$ und solche bzgl. der Tangentenvektoren $\vec{T}_\alpha$ der Koordinatenlinien weiterhin mit griechischen Vektoren. Nun ist

$$\vec{T}_1 \cdot (\vec{T}_2 \times \vec{T}_3) = (\partial_1 x^a)(\partial_2 x^b)(\partial_3 x^c) e_{abc} = J, \quad (3.19.19)$$

da wir das Vektorprodukt mittels kartesischer Vektorkomponenten mit dem Levi-Civita-Symbol bilden dürfen. Dabei ist offenbar

$$J = \det \frac{\partial(x^1, x^2, x^3)}{\partial(q^1, q^2, q^3)} > 0 \quad (3.19.20)$$

Die Jacobi-Determinante der Transformation von kartesischen Koordinaten (x^1, x^2, x^3) zu den beliebigen generalisierten Koordinaten. Es ist also

$$\vec{T}_\alpha \cdot (\vec{T}_\beta \times \vec{T}_\gamma) = J e_{\alpha\beta\gamma}. \quad (3.19.21)$$

Mit Hilfe des Transformationsverhaltens der $\vec{T}_\alpha$ zu neuen generalisierten Koordinaten q'_α (3.19.5) sieht man nach einer kurzen Rechnung (*Nachrechnen!*), dass sich

$$\epsilon_{\alpha\beta\gamma} = J e_{\alpha\beta\gamma} \quad (3.19.22)$$

wie die kovarianten Komponenten eines Tensors transformiert. Wir nennen den dadurch definierten Tensor den **Levi-Civita-Tensor**.

Wir können nun die Jacobi-Determinante J mit Hilfe der Metrikkomponenten $g_{\alpha\beta}$ ausdrücken, denn es ist

$$\begin{aligned} g &:= \det(g_{\alpha\beta}) = \det(\vec{T}_\alpha \cdot \vec{T}_\beta) = \det(\partial_\alpha x^j \partial_\beta x^k \delta_{jk}) = \det(T^j_\alpha T^j_\beta) \\ &= \det(\hat{T}^T \hat{T}) = (\det \hat{T})^2 = J^2 \Rightarrow J = \sqrt{g}. \end{aligned} \quad (3.19.23)$$

Wir können also (3.19.22) auch in der Form

$$\epsilon_{\alpha\beta\gamma} = \sqrt{g} e_{\alpha\beta\gamma} \quad (3.19.24)$$

schreiben.

Für die kontravarianten Komponenten des Levi-Civita-Tensors ergibt sich daraus

$$\epsilon^{\delta\epsilon\eta} = g^{\alpha\delta} g^{\beta\epsilon} g^{\gamma\eta} \epsilon_{\alpha\beta\gamma} = \sqrt{g} e^{\delta\epsilon\eta} \det(g^{\mu\nu}) = \frac{1}{\sqrt{g}} e^{\delta\epsilon\eta}, \quad (3.19.25)$$

Dabei haben wir verwendet, dass die Matrix $(g^{\mu\nu})$ definitionsgemäß die Inverse der Matrix $(g_{\mu\nu})$ und somit $\det(g^{\mu\nu}) = 1/\det(g_{\mu\nu}) = 1/g$ ist.

Wir können nun offenbar die duale Basis auch als **reziproke Basis** mittels

$$\vec{T}^\alpha = \frac{1}{2J} \epsilon^{\alpha\beta\gamma} (\vec{T}_\beta \times \vec{T}_\gamma) = \frac{1}{2\sqrt{g}} \epsilon^{\alpha\beta\gamma} (\vec{T}_\beta \times \vec{T}_\gamma) = \frac{1}{2} \epsilon^{\alpha\beta\gamma} \vec{T}_\beta \times \vec{T}_\gamma \quad (3.19.26)$$

berechnen, denn dann folgt

$$\begin{aligned} \vec{T}_\delta \cdot \vec{T}^\alpha &= \frac{1}{2J} \epsilon^{\alpha\beta\gamma} \vec{T}_\delta \cdot (\vec{T}_\beta \times \vec{T}_\gamma) = \frac{1}{2J} \epsilon^{\alpha\beta\gamma} \epsilon_{\delta\beta\gamma} \det \hat{T} \\ &= \frac{1}{2} (\delta_\delta^\alpha \delta_\beta^\beta - \delta_\delta^\beta \delta_\beta^\alpha) = \frac{1}{2} (3 - 1) \delta_\delta^\alpha = \delta_\delta^\alpha, \end{aligned} \quad (3.19.27)$$

d.h. (3.19.16) ist erfüllt, und damit bilden die $\vec{T}^\alpha$ tatsächlich die duale Basis zu den $\vec{T}_\alpha$.

Wir können nun auch das Kreuzprodukt zweier Vektoren durch die kontravarianten Komponenten bzgl. der Basis $\vec{T}_\alpha$ ausdrücken:

$$(\vec{V} \times \vec{W})^\alpha = \vec{T}^\alpha \cdot V^\beta W^\gamma (\vec{T}_\beta \times \vec{T}_\gamma) \stackrel{(3.19.26)}{=} \frac{1}{2} V^\beta W^\gamma \epsilon^{\alpha\delta\epsilon} (\vec{T}_\delta \times \vec{T}_\epsilon) \cdot (\vec{T}_\beta \times \vec{T}_\gamma). \quad (3.19.28)$$

Nun ist

$$(\vec{T}_\delta \times \vec{T}_\epsilon) \cdot (\vec{T}_\beta \times \vec{T}_\gamma) = \vec{T}_\delta \cdot [\vec{T}_\epsilon \times (\vec{T}_\beta \times \vec{T}_\gamma)] = \vec{T}_\delta \cdot [\vec{T}_\beta (\vec{T}_\epsilon \cdot \vec{T}_\gamma) - \vec{T}_\gamma (\vec{T}_\epsilon \cdot \vec{T}_\beta)] = g_{\delta\beta} g_{\epsilon\gamma} - g_{\delta\gamma} g_{\epsilon\beta}. \quad (3.19.29)$$

Damit folgt durch Einsetzen in (3.19.28)

$$(\vec{V} \times \vec{W})^\alpha = \epsilon^{\alpha\delta\epsilon} V_\delta W_\epsilon = \epsilon^{\alpha\beta\gamma} V_\beta W_\gamma. \quad (3.19.30)$$

Wir erhalten also das Kreuzprodukt durch „Überschieben“ mit den entsprechenden Komponenten des Levi-Civita-Tensors. Dabei ist allerdings auf die Indexstellung der Objekte zu achten.

Man kann (3.19.30) auch für die kovarianten Komponente des Vektorprodukts schreiben. Dabei brauchen wir nur zu beachten, dass $\epsilon^{\alpha\delta\epsilon}$ die in allen drei Indizes kontravarianten Komponenten des Levi-Civita-Tensors sind, d.h. man kann die Indizes mit den $g_{\alpha\beta}$ nach unten ziehen. Daraus ergibt sich

$$(\vec{V} \times \vec{W})_\alpha = \epsilon_{\alpha\beta\gamma} V^\beta W^\gamma. \quad (3.19.31)$$

3.19.2 Die Differentialoperatoren grad, div und rot

Wir können nun auch die Differentialoperatoren direkt mittels der generalisierten Koordinaten und der dazugehörigen Basisvektoren ausdrücken, indem wir die bekannten Ausdrücke in kartesischen Koordinaten auf die entsprechenden Komponenten bzgl. der $\vec{T}_\alpha$ umrechnen. Dabei ist zu beachten, dass sich für kartesische Koordinaten (x^1, x^2, x^3) für die Tangentenvektoren $\vec{T}_j = \vec{e}_j = \text{const}$ mit $g'_{jk} = \vec{T}_j \cdot \vec{T}_k = \vec{e}_j \cdot \vec{e}_k = \delta_{jk}$, d.h.

$$V'_j = g'_{jk} V'^k = V'^j, \quad (3.19.32)$$

ergibt; d.h. in diesem Fall müssen wir nicht zwischen ko- und kontravarianten Komponenten bzw. Basisvektoren unterscheiden, was die Konvention, die wir weiter oben ständig verwendet haben, d.h. dass alle Indizes unten stehen, rechtfertigt.

Für den **Gradienten** eines Skalarfeldes ergibt sich

$$\text{grad} \Phi = \vec{e}^i \partial'_i \Phi = \vec{T}^i \partial'_i \Phi = T^i_{\alpha} \vec{T}^{\alpha} \partial'_i \Phi = \vec{T}^{\alpha} \partial_{\alpha} x^i \partial'_i \Phi = \vec{T}^{\alpha} \partial_{\alpha} \Phi. \quad (3.19.33)$$

Im letzten Schritt haben wir die Kettenregel verwendet. Die kovarianten Komponenten des Gradienten sind demnach einfach die partiellen Ableitungen nach den generalisierten Koordinaten

$$\nabla_{\alpha} \Phi = \partial_{\alpha} \Phi. \quad (3.19.34)$$

Bei der Ableitung von Vektoren müssen wir beachten, dass die Basisvektoren $\vec{T}_{\alpha}$ von den (q^{β}) abhängen. Beginnen wir mit der Berechnung der **Divergenz** eines Vektorfeldes. Wir gehen von dem bekannten Ausdruck für die kartesischen Komponenten aus und verwenden die obigen Transformationsformel für die Vektorkomponenten sowie die Kettenregel für die Ableitungen nach den x'^i :

$$\text{div } \vec{V} = \partial'_i V^i = (\partial'_i q^{\beta}) \partial_{\beta} (\partial_{\alpha} x'^i V^{\alpha}). \quad (3.19.35)$$

Unser Ziel ist es nun, dies allein durch die generalisierten Koordinaten auszudrücken, also ohne auf die kartesischen Komponenten x'^i zurückzugreifen. Dazu verwenden wir, dass $U^{\gamma}_k = \partial'_k q^{\gamma}$ invers zu $T^i_{\beta} = \partial_{\beta} x'^i$ ist. Man kann nun die Inverse einer Matrix über die Unterdeterminanten bestimmen. Wir drücken dadurch $\partial'_k q^{\gamma}$ mit Hilfe von Ableitungen nach den q^{α} aus. Es gilt

$$U^{\gamma}_k = \partial'_k q^{\gamma} = \frac{1}{J} e_{ijk} e^{\alpha\beta\gamma} (\partial_{\alpha} x'^i) (\partial_{\beta} x'^j), \quad (3.19.36)$$

wobei J die Jacobi-Determinante der Transformation zwischen kartesischen und generalisierten Koordinaten (3.19.20) ist. Zunächst beweisen wir (3.19.36). In der Tat ist

$$\frac{1}{J} e_{ijk} e^{\alpha\beta\gamma} (\partial_{\alpha} x'^i) (\partial_{\beta} x'^j) \partial_{\delta} x'^k = \frac{1}{J} e^{\alpha\beta\gamma} J e_{\alpha\beta\delta} = \delta^{\gamma}_{\delta}, \quad (3.19.37)$$

wobei wir im letzten Schritt die „bac-cab-Formel“ für die Levi-Civita-Symbole verwendet haben. Jedenfalls ist damit (3.19.36) bewiesen.

Mit (3.19.23) folgt durch Einsetzen von (3.19.36) in (3.19.35)

$$\begin{aligned} \text{div } \vec{V} &= \partial_{\alpha} V^{\alpha} + \frac{1}{2\sqrt{g}} e_{ijk} e^{\gamma\delta\beta} (\partial_{\gamma} x'^j) (\partial_{\delta} x'^k) \partial_{\beta} [(\partial_{\alpha} x'^i) V^{\alpha}] \\ &= \frac{1}{2\sqrt{g}} e_{ijk} e^{\gamma\delta\beta} \{ \partial_{\beta} [(\partial_{\gamma} x'^j) (\partial_{\alpha} x'^i) (\partial_{\delta} x'^k) V^{\alpha}] - (\partial_{\alpha} x'^i) (\partial_{\delta} x'^k) V^{\alpha} (\partial_{\beta} \partial_{\gamma} x'^j) \\ &\quad - (\partial_{\alpha} x'^i) (\partial_{\gamma} x'^j) V^{\alpha} (\partial_{\beta} \partial_{\delta} x'^k) \}. \end{aligned} \quad (3.19.38)$$

Die beiden letzten Terme in der geschweiften Klammer verschwinden beim Überschieben mit $e^{\gamma\delta\beta}$, weil diese Terme symmetrisch unter Vertauschung der Indexpaare $\beta\gamma$ bzw. $\beta\delta$ sind, während das Levi-Civita-Symbol beim Vertauschen dieser Indizes antisymmetrisch ist. Wendet man die Summation über die Indizes i, j und k an, erhalten wir

$$\text{div } \vec{V} = \frac{1}{2\sqrt{g}} e^{\gamma\delta\beta} e_{\alpha\gamma\delta} \partial_{\beta} (\sqrt{g} V^{\alpha}) = \frac{1}{\sqrt{g}} \partial_{\alpha} (\sqrt{g} V^{\alpha}), \quad (3.19.39)$$

wobei wir im letzten Schritt wieder die „bac-cab-Formel“ für die Levi-Civita-Symbole verwendet haben. Damit ist die Divergenz tatsächlich allein durch die auf die generalisierten Koordinaten bezogenen Größen zurückgeführt.

Für die Rotation eines Vektorfeldes gehen wir wieder von der Formel in kartesischen Koordinaten

$$\begin{aligned}
 (\text{rot } \vec{V})^i &= \epsilon^{ijk} \partial_j' V_k' = \frac{1}{2} \epsilon^{ijk} (\partial_j' V_k' - \partial_k' V_j') = \frac{1}{2} \epsilon^{ijk} \{ \partial_j' [(\partial_k' q^\alpha) V_\alpha] - \partial_k' [(\partial_j' q^\alpha) V_\alpha] \} \\
 &= \frac{1}{2} \epsilon^{ijk} [(\partial_k' q^\alpha) \partial_j' V_\alpha - (\partial_j' q^\alpha) \partial_k' V_\alpha] \\
 &= \epsilon^{ijk} (\partial_k' q^\alpha) \partial_j' V_\alpha = \epsilon^{ijk} (\partial_k' q^\alpha) (\partial_j' q^\beta) \partial_\beta V_\alpha = -\epsilon^{ijk} (\partial_j' q^\alpha) (\partial_k' q^\beta) \partial_\beta V_\alpha
 \end{aligned} \tag{3.19.40}$$

aus und transformieren sie in die allgemeinen Koordinaten:

$$\begin{aligned}
 (\text{rot } \vec{V})^\gamma &= -(\partial_i' q^\gamma) \epsilon^{ijk} (\partial_j' q^\alpha) (\partial_k' q^\beta) \partial_\beta V_\alpha \\
 &= -\frac{1}{\sqrt{g}} \epsilon^{\gamma\alpha\beta} \partial_\beta V_\alpha = +\frac{1}{\sqrt{g}} \epsilon^{\gamma\alpha\beta} \partial_\alpha V_\beta,
 \end{aligned} \tag{3.19.41}$$

wobei wir im letzten Schritt

$$\det[(\partial_i' q^\alpha)] = \det \hat{U} = \det \hat{T}^{-1} = \frac{1}{\det \hat{T}} \stackrel{(3.19.23)}{=} \frac{1}{\sqrt{g}} \tag{3.19.42}$$

verwendet haben. Mit den kontravarianten Komponenten des Levi-Civita-Tensors (3.19.25) erhalten wir schließlich

$$(\text{rot } \vec{V})^\gamma = \epsilon^{\gamma\alpha\beta} \partial_\alpha V_\beta. \tag{3.19.43}$$

3.19.3 Kovariante Ableitungen

Wir können nun auch noch allgemeinere Ableitungsoperatoren für Felder als den Gradienten eines Skalarfeldes und Divergenz und Rotation eines Vektorfeldes einführen, die aus beliebigen Tensorfeldern neue Tensorfelder bilden.

Es ist unmittelbar klar, dass in kartesischen Koordinaten

$$W_i'^j = \partial_i' V'^j \tag{3.19.44}$$

ein Tensor zweiter Stufe ist. Wir wollen diesen Tensor nun wieder in Komponenten bzgl. der allgemeinen generalisierten Koordinaten ausdrücken:

$$W_\alpha^\beta = (\partial_\alpha x'^i) (\partial_j' q^\beta) \partial_{i'} [(\partial_\delta x'^j) V^\delta] = (\partial_\alpha x'^i) (\partial_j' q^\beta) (\partial_i' q^\gamma) \partial_\gamma [(\partial_\delta x'^j) V^\delta]. \tag{3.19.45}$$

Nun ist

$$(\partial_\alpha x'^i) (\partial_i' q^\gamma) = \partial_\alpha q^\gamma = \delta_\alpha^\gamma. \tag{3.19.46}$$

Dies in (3.19.45) eingesetzt ergibt

$$W_\alpha^\beta = (\partial_j' q^\beta) \partial_\gamma [(\partial_\delta x'^j) V^\delta] = \partial_\alpha V^\beta + \Gamma_{\alpha\delta}^\beta V^\delta. \tag{3.19.47}$$

Dabei haben wir das sog. **Christoffel-Symbol**

$$\Gamma_{\gamma\delta}^\beta = (\partial_j' q^\beta) (\partial_\gamma \partial_\delta x'^j) \tag{3.19.48}$$

eingeführt. Es ist wichtig zu bemerken, dass es sich hierbei *nicht um Tensorkomponenten* handelt.

Man kann sich leicht davon überzeugen, dass $\partial_\alpha V^\beta$ ebenfalls keine Tensorkomponenten bilden. Der Zusatzterm mit dem Christoffelsymbol (3.19.47) trägt der Abhängigkeit der Basis $\vec{T}_\alpha$ von den generalisierten Koordinaten Rechnung und sorgt dafür, dass (3.19.47) tatsächlich Tensorkomponenten sind.

Man bezeichnet den entsprechenden Operator ∇_α als **kovariante Ableitung**, d.h.

$$W_\alpha^\beta = \nabla_\alpha V^\beta = \partial_\alpha V^\beta + \Gamma_{\alpha\gamma}^\beta V^\gamma. \quad (3.19.49)$$

Es ist nun wieder wichtig, dass wir die Christoffel-Symbole ohne Rückgriff auf kartesische Koordinaten und die Umrechnung auf generalisierte Koordinaten aus den Komponenten des Metrik-Tensors $g_{\alpha\beta}$ und $g^{\alpha\beta}$ berechnen können.

Dazu bemerken wir, dass gemäß (3.19.7) und gemäß der Definition von $(g^{\alpha\beta})$ als Inverse von $(g_{\alpha\beta})$

$$g_{\alpha\beta} = \vec{T}_\alpha \cdot \vec{T}_\beta = \delta_{ij} (\partial_\alpha x'^i) (\partial_\beta x'^j) \quad (3.19.50)$$

und

$$g^{\alpha\beta} = \delta^{ij} (\partial'_i q^\alpha) (\partial'_j q^\beta) \quad (3.19.51)$$

gilt. Letzteres folgt sofort aus der Kettenregel:

$$g^{\alpha\beta} g_{\beta\gamma} = \delta^{ij} (\partial'_i q^\alpha) (\partial'_j q^\beta) \delta_{kl} (\partial_\beta x'^k) (\partial_\gamma x'^l) = \delta^{ij} (\partial'_i q^\alpha) \delta_j^k (\partial_\gamma x'^l) \delta_{kl} = (\partial'_i q^\alpha) (\partial_\gamma x'^i) = \delta_\gamma^\alpha. \quad (3.19.52)$$

Damit ist in der Tat das durch (3.19.51) gegebene $g^{\alpha\beta}$ die inverse Matrix von $g_{\beta\gamma}$.

Es liegt nun nahe, dass die Christoffel-Symbole (3.19.48) durch Ableitungen der $g_{\alpha\beta}$ nach den q^γ ausgedrückt werden können. Dazu berechnen wir unter Verwendung von (3.19.50)

$$\partial_\alpha g_{\beta\gamma} = \delta_{ij} \partial_\alpha [(\partial_\beta x'^i) (\partial_\gamma x'^j)] = \delta_{ij} [(\partial_\alpha \partial_\beta x'^i) (\partial_\gamma x'^j) + (\partial_\beta x'^i) (\partial_\alpha \partial_\gamma x'^j)]. \quad (3.19.53)$$

Wir schreiben nun noch denselben Ausdruck in den beiden Varianten hin, die durch zyklische Permutation der Indizes α , β und γ aus (3.19.53) hervorgehen:

$$\partial_\alpha g_{\beta\gamma} = \delta_{ij} [(\partial_\beta \partial_\gamma x'^i) (\partial_\alpha x'^j) + (\partial_\gamma x'^i) (\partial_\beta \partial_\alpha x'^j)], \quad (3.19.54)$$

$$\partial_\beta g_{\gamma\alpha} = \delta_{ij} [(\partial_\gamma \partial_\alpha x'^i) (\partial_\beta x'^j) + (\partial_\alpha x'^i) (\partial_\gamma \partial_\beta x'^j)]. \quad (3.19.55)$$

Addieren wir nun (3.19.53) und (3.19.54) und subtrahieren (3.19.55) erhalten wir

$$\partial_\alpha g_{\beta\gamma} + \partial_\beta g_{\gamma\alpha} - \partial_\gamma g_{\alpha\beta} = 2\delta_{ij} (\partial_\alpha \partial_\beta x'^i) (\partial_\gamma x'^j) = 2G_{\gamma\alpha\beta}. \quad (3.19.56)$$

Wir wollen nun noch beweisen, dass sich durch Ziehen des Index γ nach oben aus $G_{\gamma\alpha\beta}$ tatsächlich das Christoffel-Symbol ergibt. In der Tat folgt mit (3.19.51) aus (3.19.57)

$$\begin{aligned} G_{\alpha\beta}^\delta &= g^{\delta\gamma} G_{\gamma\alpha\beta} = \delta^{kl} (\partial'_k q^\delta) (\partial'_l q^\gamma) (\partial_\alpha \partial_\beta x'^i) \partial_{\gamma x'^i} \\ &= (\partial_\alpha \partial_\beta x'^i) \delta^{kl} (\partial'_k q^\delta) \delta_l^j \delta_{ij} \\ &= (\partial_i q^\delta) (\partial_\alpha \partial_\beta x'^i) \stackrel{3.19.48}{=} \Gamma_{\alpha\beta}^\delta. \end{aligned} \quad (3.19.57)$$

Wir erhalten also insgesamt aus (3.19.56) und (3.19.57)

$$\Gamma_{\alpha\beta}^\delta = \frac{1}{2} g^{\delta\gamma} (\partial_\alpha g_{\beta\gamma} + \partial_\beta g_{\alpha\gamma} - \partial_\gamma g_{\alpha\beta}). \quad (3.19.58)$$

3. Vektoranalysis

Damit ist die kovariante Ableitung (3.19.49) in ihrer Wirkung auf kontravariante Vektorkomponenten definiert.

Wir wollen nun auch die Wirkung der kovarianten Ableitung auf kovariante Vektorkomponenten bestimmen. Durch Transformation von den entsprechenden Komponenten bzgl. kartesischer Koordinaten erhalten wir

$$\begin{aligned}\nabla_\alpha V_\beta &= (\partial_\alpha x'^i)(\partial_\beta x'^j)\partial'_i V'_j = (\partial_\alpha x'^i)(\partial_\beta x'^j)\partial'_i[(\partial'_j q^\gamma)V_\gamma] \\ &= \partial_\alpha V_\beta + (\partial'_i \partial'_j q^\gamma)(\partial_\alpha x'^i)(\partial_\beta x'^j)V_\gamma.\end{aligned}\quad (3.19.59)$$

Nun leiten wir

$$(\partial'_j q^\gamma)(\partial_\delta x'^j) = \delta^\gamma_\delta \quad (3.19.60)$$

nach x'^i ab:

$$(\partial'_i \partial'_j q^\gamma)(\partial_\delta x'^j) + (\partial'_j q^\gamma)(\partial'_i \partial'_\epsilon q^\epsilon)(\partial_\epsilon x'^j) = 0 \Rightarrow (\partial'_i \partial'_j q^\gamma)(\partial_\delta x'^j) = -(\partial'_j q^\gamma)(\partial'_i \partial'_\epsilon q^\epsilon)(\partial_\epsilon x'^j). \quad (3.19.61)$$

Setzen wir dies in (3.19.60) ein, finden wir

$$\begin{aligned}\nabla_\alpha V_\beta &= \partial_\alpha V_\beta - (\partial_\epsilon \partial_\beta x'^j)(\partial'_j q^\gamma) \underbrace{(\partial'_i q^\epsilon)(\partial_\alpha x'^i)}_{\delta^\epsilon_\alpha} \\ &= \partial_\alpha V_\beta - (\partial_\alpha \partial_\beta x'^j)(\partial'_j q^\gamma)V_\gamma.\end{aligned}\quad (3.19.62)$$

Mit (3.19.48) folgt schließlich

$$\nabla_\alpha V_\beta = \partial_\alpha V_\beta - \Gamma^\beta_{\alpha\gamma} V_\gamma. \quad (3.19.63)$$

Es ist nun klar, dass wir Komponenten von Tensoren bzgl. jeder Stufe kovariant ableiten können, wodurch wir Komponenten von Tensoren einer um eins höheren Stufe erhalten. Dabei sind die „Korrekturterme“ mit dem Christoffelsymbol für jeden der entsprechenden kontra- oder kovarianten Indizes anzuwenden. Z.B. ist

$$\nabla_\alpha T^\beta_\gamma = \partial_\alpha T^\beta_\gamma + \Gamma^\beta_{\alpha\delta} T^\delta_\gamma - \Gamma^\delta_{\alpha\gamma} T^\beta_\delta. \quad (3.19.64)$$

Wir können nun allgemeine Formeln für Tensorkomponenten in beliebigen Koordinaten berechnen und wissen, dass sie dann bzgl. der Komponenten desselben Tensorausdrucks auch in beliebigen Koordinaten gelten. Insbesondere können wir in kartesischen Koordinaten rechnen. Da für diese $g'_{ij} = \delta_{ij} = \text{const}$ und $g'^{ij} = \delta^{ij} = \text{const}$ gilt, verschwinden die Christoffelsymbole, und die kovarianten Ableitungen sind folglich durch partielle Ableitungen gegeben.

Insbesondere erhält man für die kovariante Ableitung der Metrikkomponenten

$$\vec{\nabla}'_i g'_{jk} = \partial'_i \delta_{jk} = 0. \quad (3.19.65)$$

Da dies ein kovarianter Ausdruck von Tensorkomponenten ist, gilt die Formel auch für allgemeine Koordinaten, d.h.

$$\vec{\nabla}_\alpha g_{\beta\gamma} = 0. \quad (3.19.66)$$

Das kann man freilich auch aus der Definition der kovarianten Ableitung mittels der Christoffelsymbole und (3.19.58) direkt nachrechnen (*Übung!*). Dieses Beispiel zeigt, dass man sehr komplizierte Rechnungen drastisch vereinfachen kann, indem man die entsprechenden Ausdrücke mittels kartesischer Koordinaten in allgemein kovarianter Form von Tensorkomponenten schreibt. Dann weiß man, dass diese Ausdrücke entsprechend auch in allgemeinen Koordinaten gelten.

So ist auch

$$\nabla_\alpha g^{\beta\gamma} = 0, \quad (3.19.67)$$

wie man sofort aus der entsprechenden Rechnung in kartesischen Koordinaten sieht

$$\nabla'_i g'^{jk} = \partial'_i \delta'^{jk} = 0. \quad (3.19.68)$$

Wir zeigen nun noch, dass die oben gefundenen Ausdrücke für div und rot sich auch aus den kovarianten Ableitungen der Vektorkomponenten ergeben. Demnach ist

$$\text{div } \vec{V} = \nabla_\alpha V^\alpha = \partial_\alpha V^\alpha + \Gamma^\alpha_{\alpha\beta} V^\beta. \quad (3.19.69)$$

Wir benötigen also die Kontraktion des Christoffel-Symbols. Mit (3.19.58) folgt nach einfachen Umformungen (*Nachrechnen!*)

$$\Gamma^\alpha_{\alpha\beta} = \frac{1}{2} g^{\alpha\alpha'} \partial_\beta g_{\alpha'\alpha}. \quad (3.19.70)$$

Wir können nun wieder die Darstellung der inversen Matrix $g^{\alpha\alpha'}$ von $g_{\alpha\beta}$ über Determinanten

$$g^{\alpha\alpha'} = \frac{1}{2g} e^{\alpha\delta\epsilon} e^{\alpha'\mu\nu} g_{\delta\mu} g_{\epsilon\nu} \quad (3.19.71)$$

verwenden. Damit wird aus (3.19.70)

$$\Gamma^\alpha_{\alpha\beta} = \frac{1}{4g} e^{\alpha\delta\epsilon} e^{\alpha'\mu\nu} g_{\delta\mu} g_{\epsilon\nu} \partial_\beta g_{\alpha'\alpha}. \quad (3.19.72)$$

Andererseits ist

$$\det(g_{\alpha\beta}) = g = \frac{1}{3!} e^{\alpha\beta\epsilon} e^{\alpha'\mu\nu} g_{\delta\mu} g_{\epsilon\nu} g_{\alpha'\alpha}. \quad (3.19.73)$$

Daraus folgt

$$\begin{aligned} \partial_\beta g &= \frac{1}{3!} e^{\alpha\beta\epsilon} e^{\alpha'\mu\nu} [(\partial_\beta g_{\delta\mu}) g_{\epsilon\nu} g_{\alpha'\alpha} + g_{\delta\mu} (\partial_\beta g_{\epsilon\nu}) g_{\alpha'\alpha} + g_{\delta\mu} g_{\epsilon\nu} (\partial_\beta g_{\alpha'\alpha})] \\ &= \frac{1}{2} e^{\alpha\beta\epsilon} e^{\alpha'\mu\nu} g_{\delta\mu} g_{\epsilon\nu} (\partial_\beta g_{\alpha'\alpha}). \end{aligned} \quad (3.19.74)$$

Damit ergibt sich für (3.19.72)

$$\Gamma^\alpha_{\alpha\beta} = \frac{1}{2g} \partial_\beta g = \frac{1}{\sqrt{g}} \partial_\beta \sqrt{g}. \quad (3.19.75)$$

Setzen wir dies in (3.19.69) ein, erhalten wir schließlich

$$\text{div } \vec{V} = \nabla_\alpha V^\alpha = \partial_\alpha V^\alpha + \frac{1}{\sqrt{g}} A^\gamma \partial_\gamma \sqrt{g} = \frac{1}{\sqrt{g}} \partial_\alpha (\sqrt{g} A^\alpha), \quad (3.19.76)$$

was mit (3.19.39) übereinstimmt.

Ebenso ergibt sich sofort die Formel für die Rotation. Dazu bilden wir zunächst

$$\nabla_\alpha A_\beta - \nabla_\beta A_\alpha = \partial_\alpha A_\beta - \Gamma^\gamma_{\alpha\beta} A_\gamma - (\partial_\beta A_\alpha - \Gamma^\gamma_{\beta\alpha} A_\gamma) = \partial_\alpha A_\beta - \partial_\beta A_\alpha, \quad (3.19.77)$$

und damit wird die Rotation, definiert durch Kontraktion mit dem Levi-Civita-Tensor

$$(\text{rot } \vec{A})^\gamma = \frac{1}{2} \epsilon^{\gamma\alpha\beta} (\nabla_\alpha A_\beta - \nabla_\beta A_\alpha) = \epsilon^{\gamma\alpha\beta} \partial_\alpha A_\beta, \quad (3.19.78)$$

was mit (3.19.43) übereinstimmt.

3.19.4 Alternierende Formen und Differentialformen

In diesem Abschnitt betrachten wir den wichtigen Spezialfall **vollständig antisymmetrischer Tensoren**. Dabei ist ein Tensor p -ter Stufe $\hat{T}$ antisymmetrisch, wenn für beliebige Vertauschungen zweier der p Argumente der Tensor das Vorzeichen ändert. Das bedeutet, dass die kovarianten Komponenten

$$\omega_{\alpha_1 \dots \alpha_p} = \hat{\omega}(\vec{T}_{\alpha_1}, \dots, \vec{T}_{\alpha_p}) \quad (3.19.79)$$

vollständig antisymmetrisch unter beliebigen Vertauschungen zweier Indexpaare sind. Sind insbesondere unter den Indizes α_j mit $j \in \{1, \dots, p\}$ wenigstens zwei gleich, verschwindet die entsprechende Tensorkomponente. Es gibt also in unserem Fall ($d = 3$) nur die Fälle $p \in \{0, 1, 2, 3\}$. Dabei ist der Fall $p = 0$ definitionsgemäß ein Skalar. Im Folgenden betrachten wir nun natürlich auch wieder die entsprechenden vollständig antisymmetrischen Tensorfelder. Man bezeichnet einen vollständig symmetrischen Tensor p -ter Stufe auch als **alternierende p -Form**.

Im Folgenden sind die **verallgemeinerten Kronecker-Symbole** $\delta_{\alpha_1 \dots \alpha_p}^{\beta_1 \dots \beta_p}$ wichtig. Wir betrachten sie gleich für den allgemeinen Fall, dass unser Vektorraum nicht 3 sondern allgemein d dimensionen besitzt. Sie sind definiert durch

$$\delta_{\alpha_1 \dots \alpha_p}^{\beta_1 \dots \beta_p} = \begin{cases} 1 & \text{falls } (\alpha_1, \dots, \alpha_p) \text{ gerade Permutation von } (\beta_1, \dots, \beta_p), \\ -1 & \text{falls } (\alpha_1, \dots, \alpha_p) \text{ ungerade Permutation von } (\beta_1, \dots, \beta_p), \\ 0 & \text{sonst.} \end{cases} \quad (3.19.80)$$

Es ist klar, dass sie nur für $p \in \{1, \dots, d\}$ nicht trivial sind, denn für $p > d$ verschwinden sie alle identisch. Wir können (3.19.80) auch in der Form

$$\delta_{\alpha_1 \dots \alpha_p}^{\beta_1 \dots \beta_p} = \det \begin{pmatrix} \delta_{\alpha_1}^{\beta_1} & \delta_{\alpha_2}^{\beta_1} & \dots & \delta_{\alpha_p}^{\beta_1} \\ \delta_{\alpha_1}^{\beta_2} & \delta_{\alpha_2}^{\beta_2} & \dots & \delta_{\alpha_p}^{\beta_2} \\ \vdots & \vdots & \ddots & \vdots \\ \delta_{\alpha_1}^{\beta_p} & \delta_{\alpha_2}^{\beta_p} & \dots & \delta_{\alpha_p}^{\beta_p} \end{pmatrix} \quad (3.19.81)$$

schreiben. Entwickeln wir diese Determinante nach der letzten Zeile, folgt

$$\delta_{\alpha_1 \dots \alpha_p}^{\beta_1 \dots \beta_p} = \sum_{j=1}^p (-1)^{j+p} \delta_{\alpha_j}^{\beta_p} \delta_{\alpha_1 \dots \hat{\alpha}_j \dots \alpha_p}^{\beta_1 \dots \beta_{p-1}}. \quad (3.19.82)$$

Dabei bedeutet der Index $\hat{\alpha}_j$, dass dieser aus der entsprechenden Indexliste wegzulassen ist. Kontrahieren wir dies über das letzte Indexpaar (α_p, β_p) , erhalten wir von dem Term mit $j = p$ wegen $\delta_{\alpha_p}^{\alpha_p} = d$ einen Beitrag $d \delta_{\alpha_1 \dots \alpha_{p-1}}^{\beta_1 \dots \beta_{p-1}}$. Die anderen Beiträge liefern $(p-1)$ -mal $-\delta_{\alpha_1 \dots \alpha_{p-1}}^{\beta_1 \dots \beta_{p-1}}$, d.h. es ist

$$\delta_{\alpha_1 \dots \alpha_p}^{\beta_1 \dots \beta_p} = (d - p + 1) \delta_{\alpha_1 \dots \alpha_{p-1}}^{\beta_1 \dots \beta_{p-1}} \quad (3.19.83)$$

und schließlich durch Iteration

$$\delta_{\alpha_1 \dots \alpha_s \alpha_{s+1} \dots \alpha_p}^{\beta_1 \dots \beta_s \alpha_{s+1} \dots \alpha_p} = (d-s)(d-s+1) \dots (d-p)(d-p+1) \delta_{\alpha_1 \dots \alpha_2}^{\beta_1 \dots \beta_2} = \frac{(d-s)!}{(d-p)!} \delta_{\alpha_1 \dots \alpha_2}^{\beta_1 \dots \beta_2}. \quad (3.19.84)$$

Es ist nun leicht zu zeigen, dass für die Komponenten alternierender p -Formen (*Übung!*)

$$\frac{1}{p!} \delta_{\alpha_1 \dots \alpha_p}^{\beta_1 \dots \beta_p} \omega_{\beta_1 \dots \beta_p} = \omega_{\alpha_1 \dots \alpha_p} \quad (3.19.85)$$

gilt. Schließlich ist (*warum?*) für $p = d$

$$\delta_{\alpha_1 \dots \alpha_d}^{\beta_1 \dots \beta_d} = \epsilon_{\alpha_1 \dots \alpha_d} \epsilon^{\beta_1 \dots \beta_d} = e_{\alpha_1 \dots \alpha_d} e^{\beta_1 \dots \beta_d}. \quad (3.19.86)$$

Das zeigt, dass das verallgemeinerte Kronecker-Symbol für $p = d$ Tensorkomponenten bilden. Wegen (3.19.84) können wir die Kronecker-Symbole für alle $p < d$ durch Kontraktion einer entsprechenden Anzahl von Indexpaaren erzeugen, d.h. dass auch für $p < d$ diese verallgemeinerten Kronecker-Symbole Tensorkomponenten sind.

Damit können wir nun **invariante Integrale** über p -Formen definieren, und zwar ohne auf den ϵ -Tensor oder andere mit den Komponenten der Metrik verknüpfte Größen zurückzugreifen, d.h. diese Integrale können auch auf Vektorräumen ohne Skalarprodukt definiert werden. Dabei handelt es sich um Integrale über alternierende p -Formenfelder über p -dimensionale Hyperflächen. Dabei ist eine p -dimensionale Hyperfläche H_p mittels beliebiger generalisierter Koordinaten u_i ($i \in \{1, \dots, p\}$) über

$$H_p: \quad q^\alpha = a^\alpha(u_1, \dots, u_p), \quad (u_1, \dots, u_p) \in G \subseteq \mathbb{R}^p, \quad (3.19.87)$$

definiert. Dabei ist natürlich $\alpha \in \{1, \dots, d\}$, denn die q^α sind ja generalisierte Koordinaten für Punkte in E^d . Dann definiert man für beliebige alternierende p -Formenfelder

$$\int_{H_p} d^p f^{\alpha_1 \dots \alpha_p} \omega_{\alpha_1 \dots \alpha_p}(q) := \int_G d^p u \frac{1}{p!} \delta_{\beta_1 \dots \beta_p}^{\alpha_1 \dots \alpha_p} \frac{\partial q^{\beta_1}}{\partial u^1} \dots \frac{\partial q^{\beta_p}}{\partial u^p} \omega_{\alpha_1 \dots \alpha_p}[q(u)]. \quad (3.19.88)$$

Nun ist aber

$$\delta_{\beta_1 \dots \beta_p}^{\alpha_1 \dots \alpha_p} \frac{\partial q^{\beta_1}}{\partial u^1} \dots \frac{\partial q^{\beta_p}}{\partial u^p} = \det \left(\frac{\partial(q^{\alpha_1}, \dots, q^{\alpha_p})}{\partial(u_1, \dots, u_p)} \right) = e_{j_1 \dots j_p} \frac{\partial q^{\alpha_1}}{\partial u^{j_1}} \dots \frac{\partial q^{\alpha_p}}{\partial u^{j_p}}. \quad (3.19.89)$$

Es ist nun leicht zu zeigen, dass (3.19.88) *parametrisierungsinvariant* ist. Seien nämlich $u'_k : G \rightarrow G'$ umkehrbar eindeutige stetig differenzierbare Funktionen der u_j , dann ist

$$\begin{aligned} \det \left(\frac{\partial(q^{\alpha_1}, \dots, q^{\alpha_p})}{\partial(u'_1, \dots, u'_p)} \right) &= \det \left[\frac{\partial(q^{\alpha_1}, \dots, q^{\alpha_p})}{\partial(u_1, \dots, u_p)} \frac{\partial(u_1, \dots, u_p)}{\partial(u'_1, \dots, u'_p)} \right] \\ &= \det \left(\frac{\partial(q^{\alpha_1}, \dots, q^{\alpha_p})}{\partial(u_1, \dots, u_p)} \right) \det \left(\frac{\partial(u_1, \dots, u_p)}{\partial(u'_1, \dots, u'_p)} \right). \end{aligned} \quad (3.19.90)$$

Damit ist nach dem allgemeinen Satz über die Substitution mehrdimensionaler Integrale in der Tat

$$\begin{aligned} \int_{G'} d^p u' \det \left(\frac{\partial(q^{\alpha_1}, \dots, q^{\alpha_p})}{\partial(u'_1, \dots, u'_p)} \right) \omega_{\alpha_1 \dots \alpha_p} &= \int_G d^p u \det \left(\frac{\partial(q^{\alpha_1}, \dots, q^{\alpha_p})}{\partial(u_1, \dots, u_p)} \right) \omega_{\alpha_1 \dots \alpha_p} \\ &= \int_{H_p} d^p f^{\alpha_1 \dots \alpha_p} \omega_{\alpha_1 \dots \alpha_p}. \end{aligned} \quad (3.19.91)$$

3.19.5 Alternierende Differentialformen und der Stokessche Integralsatz

Wir können nun einen sehr allgemeinen **Stokesschen Integralsatz** für p -Formen im E^d beweisen. Dazu definieren wir für eine alternierende $(p-1)$ -Form ω mit Komponenten $\omega_{\alpha_1 \dots \alpha_{p-1}}$ deren **alternierende Differentialform** $d\omega$ als die alternierende p -Form mit den Komponenten

$$(d\omega)_{\alpha_1 \dots \alpha_p} = \frac{1}{(p-1)!} \delta_{\alpha_1 \dots \alpha_p}^{\beta_1 \dots \beta_p} \partial_{\beta_1} \omega_{\beta_2 \dots \beta_p}. \quad (3.19.92)$$

Zuerst müssen wir beweisen, dass es sich hierbei tatsächlich um die kovarianten Komponenten eines Tensors p -ter Stufe handelt. Verwenden wir den Ausdruck auf der rechten Seite von (3.19.92) mit einer kovarianten Ableitung anstelle der partiellen Ableitung, erhalten wir in der Tat Tensorkomponenten, denn auch das verallgemeinerte Kronecker-Symbol bildet ja Tensorkomponenten. Nun ist aber

$$\nabla_{\beta_1} \omega_{\beta_2 \dots \beta_p} = \partial_{\beta_1} \omega_{\beta_2 \dots \beta_p} - \Gamma_{\beta_1 \beta_2}^{\beta'_2} \omega_{\beta'_2 \beta_3 \dots \beta_p} - \dots - \Gamma_{\beta_1 \beta_p}^{\beta'_p} \omega_{\beta_2 \dots \beta_{p-1} \beta'_p}. \quad (3.19.93)$$

Da nun die Christoffelsymbole unter Vertauschen ihrer unteren Indizes symmetrisch sind, fallen entsprechenden Beiträge beim Kontrahieren mit dem verallgemeinerten Kronecker-Symbol weg, weil es total antisymmetrisch unter Vertauschen seiner oberen Indizes ist.

Für alternierende Formen ist die antisymmetrisierte kovariante Ableitung identisch mit der entsprechenden antisymmetrisierten partiellen Ableitung, und folglich definiert (3.19.92) tatsächlich kovariante Tensorkomponenten.

Ist dann H_p eine p -dimensionale Hyperfläche und ∂p ihr Rand, also eine $(p-1)$ -dimensionale Hyperfläche, so gilt

$$\int_{H_p} d^p f^{\alpha_1 \dots \alpha_p} \frac{1}{p!} (d\omega)_{\alpha_1 \dots \alpha_p} = \int_{\partial H_p} d^{p-1} \frac{1}{(p-1)!} f^{\alpha_2 \dots \alpha_p} \omega_{\alpha_2 \dots \alpha_p}. \quad (3.19.94)$$

Zum Beweis nehmen wir an, dass sich H_p mit irgendwelchen Parametern $u = (u_1, \dots, u_p)$ parametrisieren lässt, so dass das entsprechende Parametergebiet ein Quader $G = (u_{1<}, u_{1>}) \times \dots \times (u_{p<}, u_{p>})$ ist. Dann folgt

$$\int_{H_p} d^p f^{\alpha_1 \dots \alpha_p} \frac{1}{p!} (d\omega)_{\alpha_1 \dots \alpha_p} = \int_G d^p u \frac{1}{(p-1)!} \det \left(\frac{(\partial q^{\alpha_1}, \dots, \partial q^{\alpha_p})}{\partial(u_1, \dots, u_p)} \right) \partial_{\alpha_1} \omega_{\alpha_2 \dots \alpha_p} [q(u)]. \quad (3.19.95)$$

Wir können nun den Integranden wie folgt umformen

$$\begin{aligned} \det \left(\frac{(\partial q^{\alpha_1}, \dots, \partial q^{\alpha_p})}{\partial(u_1, \dots, u_p)} \right) \partial_{\alpha_1} \omega_{\alpha_2 \dots \alpha_p} [q(u)] &= e^{j_1 \dots j_p} \frac{\partial q^{\alpha_1}}{\partial u^{j_1}} \dots \frac{\partial q^{\alpha_p}}{\partial u^{j_p}} \partial_{\alpha_1} \omega_{\alpha_2 \dots \alpha_p} [q(u)] \\ &= e^{j_1 \dots j_p} \frac{\partial q^{\alpha_2}}{\partial u^{j_2}} \dots \frac{\partial q^{\alpha_p}}{\partial u^{j_p}} \frac{\partial}{\partial u^{j_1}} \omega_{\alpha_2 \dots \alpha_p} [q(u)] \\ &= e^{j_1 \dots j_p} \frac{\partial}{\partial u^{j_1}} \left\{ \frac{\partial q^{\alpha_2}}{\partial u^{j_2}} \dots \frac{\partial q^{\alpha_p}}{\partial u^{j_p}} \omega_{\alpha_2 \dots \alpha_p} [q(u)] \right\}. \end{aligned} \quad (3.19.96)$$

Setzen wir dies in (3.19.95) ein, können wir für jedes $j_1 \in \{1, \dots, p\}$ die Integration nach u^{j_1} und die Kontrak-

tion mit dem Levi-Civita-Symbol bzgl. der Indizes $j_2, \dots, j_p$ ausführen:

$$\int_{H_p} d^p f^{\alpha_1 \dots \alpha_p} \frac{1}{p!} (d\omega)_{\alpha_1 \dots \alpha_p} = \sum_{j_1=1}^p \int_{\hat{G}} du^{j_2} \dots du^{j_p} \frac{1}{(p-1)!} e^{j_1 \dots j_p} \left\{ \frac{\partial q^{\alpha_2}}{\partial u^{j_2}} \dots \frac{\partial q^{\alpha_p}}{\partial u^{j_p}} \omega_{\alpha_2 \dots \alpha_p} [q(u)] \right\}_{q^{j_1}=q_{>}^{j_1}}^{q^{j_1}=q_{<}^{j_1}}. \quad (3.19.97)$$

Die rechte Seite ist aber gerade das invariante Integral über ∂H_p in (3.19.94), womit die Behauptung bewiesen ist.

Falls sich für H_p keine Parametrisierung mit einem Quader als Parametergebiet finden lässt, kann man es ggf. durch „infinitesimale“ solche Gebiete beliebig genau annähern. Wie schon oben beim Stokesschen Satz im E^3 heben sich dann die Beiträge der Innenhyperflächen in den Integralen über die Ränder der infinitesimalen Gebiete paarweise weg, und es bleibt nur der Beitrag vom Rand ∂H_p der ursprünglichen Hyperfläche übrig.

3.19.6 Hodge-Dualisierung und Integrale über kontravariante Tensorfeldkomponenten

Im vorigen Abschnitt haben wir gesehen, dass wir mit den alternierenden Formen und den dazugehörigen Differentialformen spezielle Objekte haben, die sich kovariant integrieren lassen, und für die sich die kovariante Ableitung zu den partiellen Ableitungen nach den verallgemeinerten Koordinaten vereinfachen. Dies war der Schlüssel für den Beweis des allgemeinen Stokesschen Integralsatzes.

Im euklidischen E^d haben wir aber auch ein Skalarprodukt als positiv definite Bilinearform zur Verfügung, und wie im Fall $d = 3$ können wir auch allgemein zeigen, dass

$$\epsilon_{\alpha_1 \dots \alpha_d} = \sqrt{g} e_{\alpha_1 \dots \alpha_d}, \quad \epsilon^{\alpha_1 \dots \alpha_d} = \frac{1}{\sqrt{g}} e^{\alpha_1 \dots \alpha_d} \quad \text{mit} \quad g = \det(g_{\alpha\beta}) \quad (3.19.98)$$

Tensorkomponenten bzgl. orientierungserhaltender Transformationen zu neuen generalisierten Koordinaten sind und den **Levi-Civita-Tensor** definieren (vgl. 3.19.1). Dabei ist $e_{\alpha_1 \dots \alpha_d} = e^{\alpha_1 \dots \alpha_d}$ das **Levi-Civita-Symbol**, das antisymmetrisch unter Vertauschungen beliebiger zweier Indizes mit $e_{12 \dots d} = 1$ ist.

Somit kann man aus den kontravarianten Komponenten von total antisymmetrischen Tensorfeldern k -ter Stufe mit Komponenten $\Omega^{\alpha_1 \dots \alpha_k}$ (mit $0 \leq k \leq d$) vermöge

$$\omega_{\alpha_{k+1} \dots \alpha_d} = {}^\dagger \Omega_{\alpha_{k+1} \dots \alpha_d} = \frac{1}{k!} \epsilon_{\alpha_1 \dots \alpha_k \alpha_{k+1} \dots \alpha_d} \Omega^{\alpha_1 \dots \alpha_k} \quad (3.19.99)$$

alternierende $p = d - k$ -Formen definieren. Man bezeichnet diese Operation als **Hodge-Dualisierung** (Sir William Vallance Douglas Hodge, 1903-1975). Wegen (3.19.86) und (3.19.84) kann man diese Operation auch eindeutig umkehren:

$$\begin{aligned} {}^\dagger \omega^{\alpha_1 \dots \alpha_k} \frac{1}{(d-k)!} \epsilon^{\alpha_1 \dots \alpha_k \alpha_{k+1} \dots \alpha_d} \omega_{\alpha_{k+1} \dots \alpha_d} &= \frac{1}{k!(d-k)!} \epsilon^{\alpha_1 \dots \alpha_k \alpha_{k+1} \dots \alpha_d} \epsilon_{\beta_1 \dots \beta_k \alpha_{k+1} \dots \alpha_d} \Omega^{\beta_1 \dots \beta_k} \\ &= \frac{1}{k!(d-k)!} \delta^{\alpha_1 \dots \alpha_k \alpha_{k+1} \dots \alpha_d}_{\beta_1 \dots \beta_k \alpha_{k+1} \dots \alpha_d} \Omega^{\beta_1 \dots \beta_k} \\ &= \frac{1}{k!} \delta^{\alpha_1 \dots \alpha_k}_{\beta_1 \dots \beta_k} \Omega^{\beta_1 \dots \beta_k} = \Omega^{\alpha_1 \dots \alpha_k}. \end{aligned} \quad (3.19.100)$$

Jetzt können wir die invarianten Integrale für alternierende Formen verwenden, um entsprechende invariante Integrale über kontravariante total antisymmetrische Tensorfelder zu definieren. Dazu verwenden wir einfach

3. Vektoranalysis

das Hodge-Dual dieses Tensorfeldes. Für ein Tensorfeld k -ter Stufe können wir also ein invariantes Integral über eine $(d-k)$ -dimensionale Hyperfläche definieren, d.h. unter Verwendung von (3.19.99)

$$\begin{aligned} \int_{H_{d-k}} d^{d-k} f^{\alpha_{k+1} \dots \alpha_d} \frac{1}{(d-k)!} \omega_{\alpha_{k+1} \dots \alpha_d} &= \int_{H_{d-k}} d^{d-k} f^{\alpha_{k+1} \dots \alpha_d} \frac{1}{(d-k)! k!} \epsilon_{\alpha_1 \dots \alpha_k \alpha_{k+1} \dots \alpha_d} \Omega^{\alpha_1 \dots \alpha_k} \\ &= \int_{H_{d-k}} (\dagger d^{d+k} f)_{\alpha_1 \dots \alpha_k} \Omega^{\alpha_1 \dots \alpha_k}. \end{aligned} \quad (3.19.101)$$

Jetzt können wir auch den **Gaußschen Integralsatz** auf beliebige Dimensionen erweitern. Dazu wenden wir den allgemeinen Stokesschen Integralsatz auf

$$\omega_{\alpha_2 \dots \alpha_d} = \epsilon_{\alpha_1 \dots \alpha_d} V^{\alpha_1} = \sqrt{g} e_{\alpha_1 \dots \alpha_d} V^{\alpha_1} \quad (3.19.102)$$

an. Nun ist

$$\begin{aligned} (d\omega)_{\alpha_1 \dots \alpha_d} &= \frac{1}{(d-1)!} \delta_{\alpha_1 \dots \alpha_d}^{\alpha'_1 \dots \alpha'_d} e_{\beta'_1 \alpha'_2 \dots \alpha'_d} \partial_{\alpha'_1} (\sqrt{g} V^{\beta'_1}) \\ &= \frac{1}{(d-1)!} e_{\alpha_1 \dots \alpha_d} e^{\alpha'_1 \dots \alpha'_d}_{\beta'_1 \alpha'_2 \dots \alpha'_d} \partial_{\alpha'_1} (\sqrt{g} V^{\beta'_1}) \\ &= e_{\alpha_1 \dots \alpha_d} \delta_{\beta'_1}^{\alpha'_1} \partial_{\alpha'_1} (\sqrt{g} V^{\alpha'_1}) \\ &= \epsilon_{\alpha_1 \dots \alpha_d} \operatorname{div} \vec{V}. \end{aligned} \quad (3.19.103)$$

Dabei haben wir im letzten Schritt (3.19.76) verwendet. Integrieren wir dies über ein d -dimensionales Volumen V , folgt mit dem allgemeinen Stokesschen Integralsatz

$$\begin{aligned} \int_V d^d f^{\alpha_1 \dots \alpha_d} \frac{1}{d!} (d\omega)_{\alpha_1 \dots \alpha_d} &= \int_V d^d f^{\alpha_1 \dots \alpha_d} \frac{1}{d!} \epsilon_{\alpha_1 \dots \alpha_d} \operatorname{div} \vec{V} \\ &= \int_V (\dagger d^d f) \operatorname{div} \vec{V} \\ &= \int_{\partial V} d^{d-1} f^{\alpha_2 \dots \alpha_d} \frac{1}{(d-1)!} \omega_{\alpha_2 \dots \alpha_d} \\ &= \int_{\partial V} d^{d-1} f^{\alpha_2 \dots \alpha_d} \frac{1}{(d-1)!} \epsilon_{\alpha_1 \dots \alpha_d} V^{\alpha_1} \\ &= \int_{\partial V} (\dagger d^{d-1} f)_{\alpha_1} V^{\alpha_1}, \end{aligned} \quad (3.19.104)$$

d.h.

$$\int_V (\dagger d^d f) \operatorname{div} \vec{V} = \int_{\partial V} (\dagger d^{d-1} f)_{\alpha_1} V^{\alpha_1}. \quad (3.19.105)$$

Kapitel 4

Komplexe Zahlen

Die Erweiterung von den reellen zu den **komplexen Zahlen** ist durch die Forderung nach der Lösbarkeit von **Polynomgleichungen** motiviert. Während sich reelle lineare Gleichungen der Form $ax + b = 0$ für $a \neq 0$ noch im Rahmen der reellen Zahlen eindeutig lösen lassen, denn offenbar wird die obige Gleichung dann durch $x = -b/a$ eindeutig gelöst, ist dies schon für quadratische Gleichungen nicht mehr der Fall. Dies wird durch die Einführung der **imaginären Einheit** behoben, wie wir gleich im nächsten Abschnitt sehen werden. Zugleich bilden die komplexen Zahlen in Analogie zu den reellen Zahlen algebraisch gesehen einen **Zahlenkörper**, und man kann Konvergenzfragen ebenfalls vollkommen analog behandeln wie für reellen Zahlen, und die komplexen Zahlen sind bzgl. Grenzwertbildung ebenso abgeschlossen wie die reellen Zahlen. In diesem Skript werden wir nur die wichtigsten algebraischen Eigenschaften der komplexen Zahlen begründen und die wichtigsten **elementaren Funktionen** über ihre Potenzreihen und die Bildung von Umkehrfunktionen definieren. Die eigentliche **Funktionentheorie** werden wir hier nicht behandeln. Der interessierte Leser sei dazu auf die Literatur verwiesen, z.B. [Rem92].

4.1 Definition der komplexen Zahlen

Bei der Lösung quadratischer Gleichungen der Form

$$x^2 + px + q = 0 \tag{4.1.1}$$

stoßen wir auf das Problem, dass für $x \in \mathbb{R}$ stets $x^2 \geq 0$ gilt, d.h. im Rahmen der reellen Zahlen können wir keine Quadratwurzeln aus negativen Zahlen ziehen. Die Lösungsstrategie für die Gleichung (4.1.1) besteht darin, eine **quadratische Ergänzung** auszuführen. Offenbar gilt nämlich

$$x^2 + px + q = \left(x + \frac{p}{2}\right)^2 - \frac{p^2}{4} + q. \tag{4.1.2}$$

Die Gleichung (4.1.1) ist also äquivalent zu der Gleichung

$$\left(x + \frac{p}{2}\right)^2 = \frac{p^2}{4} - q. \tag{4.1.3}$$

Wollen wir diese Gleichung nach x auflösen, müssen wir die Wurzel aus der rechten Seite ziehen können. Im Bereich der reellen Zahlen ist das offensichtlich nur möglich, wenn $p^2/4 - q \geq 0$ ist. Dann besitzt die Gleichung entweder eine (falls $p^2/4 - q = 0$ ist) oder (für $p^2/4 - q > 0$) zwei Lösungen. Dies schreiben wir dann kurz als

4. Komplexe Zahlen

$$x_{1,2} = -\frac{p}{2} \pm \sqrt{\frac{p^2}{4} - q}. \quad (4.1.4)$$

Wir versuchen nun, die reellen Zahlen einfach dadurch zu erweitern, dass wir eine neue zunächst rein symbolisch zu verstehende „Zahl“ i , die **imaginäre Einheit**, einführen, für die

$$i^2 = -1 \quad (4.1.5)$$

gelten soll. Dann hätte für $a > 0$ die Gleichung $x^2 = -a$ die beiden Lösungen $x = \pm i\sqrt{a}$, wobei wir voraussetzen, dass die **komplexen Zahlen**, die allgemein von der Form

$$z = x + iy, \quad x, y \in \mathbb{R} \quad (4.1.6)$$

sein sollen, die gewöhnlichen Rechenregeln wie für reelle Zahlen gelten, also die Axiome eines **Zahlenkörpers** erfüllen. Dabei soll eine komplexe Zahl definitionsgemäß durch ihren **Real- und Imaginärteil** x bzw. y eindeutig bestimmt sein. Wir schreiben

$$\operatorname{Re} z = x, \quad \operatorname{Im} z = y. \quad (4.1.7)$$

Nehmen wir dies an, so folgt für die **Addition** zweier komplexer Zahlen

$$z_1 + z_2 = (x_1 + iy_1) + (x_2 + iy_2) = (x_1 + x_2) + i(y_1 + y_2), \quad (4.1.8)$$

wobei wir mehrfach das Assoziativ-, Kommutativ- und Distributivgesetz verwendet haben und wir i wie eine gewöhnliche Variable behandelt haben. Die Menge aller komplexen Zahlen nennen wir $\mathbb{C}$.

Die Regel für die Multiplikation folgt ebenso durch formales Ausmultiplizieren:

$$z_1 z_2 = (x_1 + iy_1)(x_2 + iy_2) = (x_1 x_2 - y_1 y_2) + i(x_1 y_2 + y_1 x_2). \quad (4.1.9)$$

Dabei haben wir im zweiten Term des Realteils die definierende Eigenschaft (4.1.5) der imaginären Einheit benutzt.

Wir berechnen gleich noch die Potenzen von i :

$$i^2 := -1, \quad i^3 = (i^2)i = -i, \quad i^4 = (i^2)(i^2) = 1 \cdot 1 = 1, \dots \quad (4.1.10)$$

Als weitere Operation an einer einzelnen komplexen Zahl ist noch die **komplexe Konjugation** nützlich. Sie ist so definiert, dass die konjugiert komplexe Zahl z^* von z denselben Real- und den entgegengesetzt gleichen Imaginärteil wie z haben soll, d.h. durch

$$z^* = x - iy. \quad (4.1.11)$$

Offensichtlich ist $(z^*)^* = z$ für alle $z \in \mathbb{C}$. Weiter ist $\mathbb{R} \subset \mathbb{C}$, denn die komplexen Zahlen mit verschwindendem Imaginärteil sind umkehrbar eindeutig auf $\mathbb{R}$ abbildbar. Offenbar ist $z \in \mathbb{R}$ genau dann, wenn $\operatorname{Im} z = 0$, was zugleich $z^* = z$ impliziert. Außerdem gilt

$$\operatorname{Re} z = \frac{z + z^*}{2}, \quad \operatorname{Im} z = \frac{z - z^*}{2i}. \quad (4.1.12)$$

Weiter rechnet man leicht nach (*Übung!*), dass

$$(z_1 + z_2)^* = z_1^* + z_2^*, \quad (z_1 z_2)^* = z_1^* z_2^* \quad (4.1.13)$$

ist. Das Produkt einer komplexen Zahl mit ihrem konjugiert Komplexen ist

$$z z^* = (x + iy)(x - iy) = x^2 - (iy)^2 = x^2 - i^2 y^2 = x^2 + y^2 \geq 0. \quad (4.1.14)$$

Den Betrag der komplexen Zahl definieren wir als

$$|z| = \sqrt{x^2 + y^2} = \sqrt{z z^*}. \quad (4.1.15)$$

Schließlich können wir auch die Division im Bereich der komplexen Zahlen betrachten. Sei dazu $z_2 \neq 0$ und $z_1 \in \mathbb{C}$ beliebig. Dann gilt

$$\frac{z_1}{z_2} = \frac{z_1 z_2^*}{z_2 z_2^*} = \frac{(x_1 + iy_1)(x_2 - iy_2)}{x_2^2 + y_2^2} = \left(\frac{x_1 x_2 + y_1 y_2}{x_2^2 + y_2^2} \right) + i \left(\frac{x_2 y_1 - x_1 y_2}{x_2^2 + y_2^2} \right). \quad (4.1.16)$$

Da $z_2 \neq 0$ ist offenbar auch $x_2^2 + y_2^2 \neq 0$ und also die Division für $z_2 \neq 0$ durch die soeben berechneten Real- und Imaginärteile wohldefiniert.

Die reellen Zahlen können wir geometrisch durch eine Zahlengerade veranschaulichen. Entsprechend kann man die komplexen Zahlen geometrisch interpretieren, wenn man das Zahlenpaar $(x, y) = (\operatorname{Re} z, \operatorname{Im} z)$ als Komponenten bzgl. eines kartesischen Koordinatensystems in der Euklidischen Ebene interpretiert. Dies ist die **Gaußsche Zahlenebene**. Es ist klar, dass $|z|$ geometrisch die Länge des entsprechenden z repräsentierenden Ortsvektors in der Gaußschen Zahlenebene ist (s. Abb. 4.1).

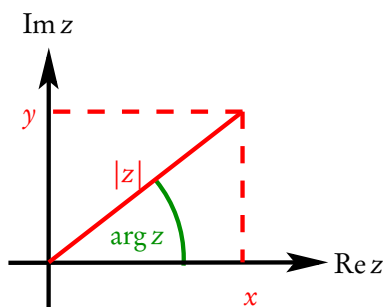


Abbildung 4.1: Zur Gaußschen Zahlenebene und Polarform einer komplexen Zahl.

Wir können nun diesen Vektor durch Polarkoordinaten (r, φ) darstellen. Offenbar ist $r = |z|$, und es gilt definitionsgemäß

$$z = x + iy = r(\cos \varphi + i \sin \varphi) = |z|(\cos \varphi + i \sin \varphi). \quad (4.1.17)$$

Definieren wir den Bereich für den Polarwinkel als $(-\pi, \pi]$, so errechnet sich dieser Winkel gemäß (s. Abschnitt 3.7.5)

$$\varphi = \arg z = \begin{cases} \operatorname{sign} y \arccos\left(\frac{x}{\sqrt{x^2 + y^2}}\right) & \text{falls } y \neq 0, \\ 0 & \text{falls } y = 0, x > 0, \\ \pi & \text{falls } y = 0, x < 0. \end{cases} \quad (4.1.18)$$

Man nennt φ auch das **Argument der komplexen Zahl** und schreibt $\varphi = \arg z$.

4.2 Potenzreihen

Die Konvergenz von Folgen und Reihen kann nun wörtlich wie für die entsprechenden Begriffe im Reellen definiert werden. Natürlich fallen alle Begriffe weg, die die Anordnungsrelationen von reellen Zahlen verwenden wie Monotoniekriterien usw. Andererseits sind natürlich alle Sätze über absolut konvergente Reihen

4. Komplexe Zahlen

anwendbar, und wegen der Dreiecksungleichung sind komplexe Folgen und Reihen genau dann konvergent, wenn ihr Real- und Imaginärteil konvergent sind.

Nun definieren wir noch einige elementare Funktionen, die wir schon aus der reellen Analysis kennen, auch für komplexe Zahlen. Dies geschieht am bequemsten über **Potenzreihen**. Die Potenzreihen weisen nun im Komplexen dieselben Konvergenzeigenschaften wie im Reellen auf, d.h. sie sind in jedem abgeschlossenen Gebiet in der komplexen Zahlenebene absolut konvergent, die ganz im Inneren des Konvergenzbereichs liegen, und das für die Potenzreihe

$$f(z) = \sum_{j=0}^{\infty} a_j (z-a)^j \quad (4.2.1)$$

ein Kreis um a mit dem **Konvergenzradius**

$$r = \lim_{j \rightarrow \infty} \left| \frac{a_j}{a_{j+1}} \right|, \quad (4.2.2)$$

falls der Limes existiert. Vgl. dazu die hier vollständig ins Komplexe Übertragbare Herleitung und Diskussion der Gl. (1.9.122).

Wir beginnen mit der Exponentialfunktion und übernehmen die entsprechende Potenzreihe einfach von der entsprechenden reellen Funktion als Definition für die Exponentialfunktion (1.9.101) im Komplexen

$$\exp z = \sum_{j=0}^{\infty} \frac{z^j}{j!}. \quad (4.2.3)$$

Auch sie konvergiert für alle $z \in \mathbb{C}$.

Weiter benötigen wir noch die trigonometrischen Funktionen. Auch ihre Potenzreihen übernehmen wir aus dem Reellen, d.h. mit (1.9.107) bzw. (1.9.109) folgt (*nachrechnen!*)

$$\cos z = 1 - \frac{z^2}{2!} + \frac{z^4}{4!} - \dots = \sum_{j=0}^{\infty} (-1)^j \frac{z^{2j}}{(2j)!} \quad (4.2.4)$$

$$\sin z = z - \frac{z^3}{3!} + \frac{z^5}{5!} - \dots = \sum_{j=0}^{\infty} (-1)^j \frac{z^{2j+1}}{(2j+1)!}. \quad (4.2.5)$$

Berechnen wir nun

$$\begin{aligned} \exp(iz) &= 1 + iz + \frac{(iz)^2}{2!} + \frac{(iz)^3}{3!} + \dots \\ &= \left(1 - \frac{z^2}{2!} + \dots \right) + i \left(z - \frac{z^3}{3!} + \dots \right) \\ &= \left(\sum_{j=0}^{\infty} (-1)^j \frac{z^{2j}}{(2j)!} \right) + i \left(\sum_{j=0}^{\infty} (-1)^j \frac{z^{2j+1}}{(2j+1)!} \right). \end{aligned} \quad (4.2.6)$$

Dabei haben wir die Reihe so umgeordnet, dass wir in einem Term den Faktor i ausklammern konnten. Das ist bei Potenzreihen erlaubt, da sie in jedem kompakten Bereich der komplexen Ebene absolut konvergiert.

Vergleichen wir nun die Reihen in den Klammern der Gleichung (4.2.6) mit (4.2.4) und (4.2.5), erhält man die **Eulersche Formel**

$$\exp(iz) = \cos z + i \sin z. \quad (4.2.7)$$

Für die Polardarstellung der komplexen Zahl (4.1.17) folgt damit

$$z = |z| \exp(i\varphi). \quad (4.2.8)$$

Da für die Exponentialfunktion auch im Komplexen die Formel

$$\exp(z_1 + z_2) = \exp(z_1) \exp(z_2) \quad (4.2.9)$$

gilt, wie man mit Hilfe der Reihe (4.2.3) beweisen kann (vgl. (1.9.104)), erleichtert dies die Rechnung mit trigonometrischen Funktionen erheblich. Z.B. folgt genau wie (4.2.7) auch die Gleichung

$$\exp(-iz) = \cos z - i \sin z. \quad (4.2.10)$$

Wir haben damit

$$\cos z = \frac{1}{2}[\exp(iz) + \exp(-iz)], \quad \sin z = \frac{1}{2i}[\exp(iz) - \exp(-iz)]. \quad (4.2.11)$$

Dies erinnert an die Definition der Hyperbelfunktionen

$$\cosh z = \frac{1}{2}[\exp(z) + \exp(-z)], \quad \sinh z = \frac{1}{2}[\exp(z) - \exp(-z)]. \quad (4.2.12)$$

Vergleicht man (4.2.11) mit diesen Definitionen folgt sofort, dass

$$\cosh(iz) = \cos z, \quad \sinh(iz) = i \sin z \quad (4.2.13)$$

gilt. Die trigonometrischen und Hyperbelfunktionen sind im Komplexen also bis auf Konstanten im wesentlichen die gleichen Funktionen, und beide sind durch die Exponentialfunktion definiert.

Genauso folgt aus (4.2.11)

$$\cos(iz) = \cosh z, \quad \sin(iz) = i \sinh z. \quad (4.2.14)$$

Als Anwendungsbeispiel leiten wir noch die Additionstheoreme für die trigonometrischen Funktionen für reelle Argumente aus (4.2.9) ab. Es gilt nämlich einerseits wegen der Eulerschen Formel (4.2.7) und (4.2.8)

$$\exp[i(\varphi_1 + \varphi_2)] = \cos(\varphi_1 + \varphi_2) + i \sin(\varphi_1 + \varphi_2) \quad (4.2.15)$$

und andererseits

$$\begin{aligned} \exp[i(\varphi_1 + \varphi_2)] &= \exp(i\varphi_1) \exp(i\varphi_2) = (\cos \varphi_1 + i \sin \varphi_1)(\cos \varphi_2 + i \sin \varphi_2) \\ &= (\cos \varphi_1 \cos \varphi_2 - \sin \varphi_1 \sin \varphi_2) + i(\sin \varphi_1 \cos \varphi_2 + \cos \varphi_1 \sin \varphi_2). \end{aligned} \quad (4.2.16)$$

4. Komplexe Zahlen

Vergleicht man nun Real- und Imaginärteil von (4.2.15) und (4.2.16), folgen die bekannten Additionstheoreme

$$\begin{aligned}\cos(\varphi_1 + \varphi_2) &= \cos \varphi_1 \cos \varphi_2 - \sin \varphi_1 \sin \varphi_2, \\ \sin(\varphi_1 + \varphi_2) &= \sin \varphi_1 \cos \varphi_2 + \cos \varphi_1 \sin \varphi_2.\end{aligned}\tag{4.2.17}$$

Es ist leicht zu zeigen, dass diese Additionstheoreme auch allgemein für beliebige komplexe Argumente gelten (*Übung*).

Kapitel 5

Gewöhnliche Differentialgleichungen

Die Entwicklung von Lösungsverfahren für **Differentialgleichungen** stellt eines der wichtigsten mathematischen Hilfsmittel für die Physik dar, denn die Naturgesetze werden durch eben solche Gleichungen beschrieben. Dabei stellen uns **gewöhnliche Differentialgleichungen** vor die Aufgabe, unbekannte Funktionen einer Veränderlichen (in der klassischen Mechanik ist das die Zeit t) aus Gleichungen zu bestimmen, die diese Funktion und Ableitungen dieser Funktion enthalten. Die höchste Ordnung der Ableitung bestimmt die **Ordnung der Differentialgleichung**. Die allgemeine Differentialgleichung n -ter Ordnung $n \in \mathbb{N}$ nimmt dann die Form

$$F(t, f, \dot{f}, \dots, f^{(n)}) = 0 \quad (5.0.1)$$

an, wobei $f^{(j)}$ die j -te Ableitung der Funktion f nach der Zeit bedeutet. Je nach Anwendung kann es erforderlich sein, alle Lösungen oder nur bestimmte, die durch die Anwendung vorgeschriebene Nebenbedingungen erfüllen, zu finden.

Systeme von Differentialgleichungen n -ter Ordnung sind entsprechend gekoppelte Differentialgleichungen für mehrere Funktionen $f_1, \dots, f_k$, wobei die die höchste Ableitung, die in diesen Gleichungen vorkommt, die n -te ist.

In der klassischen Mechanik haben wir es i.a. mit Differentialgleichungen oder Systemen von Differentialgleichungen 2. Ordnung zu tun, denn die **Newtonsche Bewegungsgleichung** für einen Massenpunkt, auf den irgendwelche vorgegebenen Kräfte wirken, lautet

$$m\ddot{\vec{x}} = \vec{F}(t, \vec{x}, \dot{\vec{x}}). \quad (5.0.2)$$

Ein Beispiel ist die Bewegung eines Planeten um die sehr viel schwerere Sonne, deren Bewegung wir näherungsweise vernachlässigen dürfen. Setzen wir die Sonne in den Ursprung eines Bezugssystems, ergibt sich mit dem Newtonschen Gravitationsgesetz die Gleichung

$$m\ddot{\vec{x}} = -\gamma m M \frac{\vec{x}}{|\vec{x}|^3}. \quad (5.0.3)$$

Hier hängt die Kraft nicht von der Geschwindigkeit ab. Dabei ist m die Masse des Planeten, M die der Sonne und γ die Newtonsche Gravitationskonstante.

Ein Beispiel für eine geschwindigkeitsabhängige Kraft ist die Bewegung eines geladenen Teilchens in einem elektromagnetischen Feld:

$$m\ddot{\vec{x}} = q \left[\vec{E}(t, \vec{x}) + \dot{\vec{x}} \times \vec{B}(t, \vec{x}) \right]. \quad (5.0.4)$$

Dabei ist m die Masse des Teilchens, q seine Ladung und $\vec{E}$ und $\vec{B}$ das durch irgendwelche anderen Ladungen und Ströme erzeugte elektrische bzw. magnetische Felder.

In der Physik suchen wir Lösungen solcher Bewegungsgleichungen unter Vorgabe bestimmter **Anfangsbedingungen**, d.h. man gibt zur Zeit $t = 0$ die Werte für Ort $\vec{x}(0) = \vec{x}_0$ und Geschwindigkeit $\dot{\vec{x}}(0) = \vec{v}_0$ des Teilchens vor und sucht Lösungen für die Differentialgleichung, die zusätzlich diese Anfangsbedingungen

$$\vec{x}(0) = \vec{x}_0, \quad \dot{\vec{x}}(0) = \vec{v}_0 \quad (5.0.5)$$

erfüllen.

Aus der physikalischen Fragestellung heraus erwarten wir, dass diese Differentialgleichungen stets eine Lösung besitzen (**Existenz**) und bei Vorgabe der Anfangsbedingungen (5.0.5) eindeutig (**Eindeutigkeit**) sind. Die Beweise entsprechender **Existenz- und Eindeutigkeitssätze** für Differentialgleichungen bei genauer Bestimmung der Eigenschaften der involvierten Funktionen (wie F in (5.0.1) oder die Kräfte in der Newtonschen Bewegungsgleichung) sind Klassiker der Analysis und finden sich in vielen mathematischen Lehrbüchern. In diesem Skript beschränken wir uns auf die Lösungsmethoden für die einfachsten Typen von Differentialgleichungen. Von der umfangreichen Spezialliteratur sei hier nur auf [Col90, Bro03] verwiesen.

5.1 Differentialgleichungen 1. Ordnung

Als einfachsten Typ von Differentialgleichungen betrachten wir Differentialgleichungen 1. Ordnung. Für sie gibt es mannigfaltige Lösungsverfahren, von denen wir hier nur einige der wichtigsten besprechen. Die allgemeinste Form des **Anfangswertproblems** ist

$$F(t, x, \dot{x}) = 0, \quad x(t_0) = x_0. \quad (5.1.1)$$

Natürlich lässt sich in dieser allgemeinsten Form wenig über die Lösungen der Differentialgleichung aussagen. Im Folgenden betrachten wir die einfachsten Fälle, in denen sich die Lösung des Anfangswertproblems zumindest in impliziter Form auf **Integrationen** zurückführen lässt.

5.1.1 Separierbare Differentialgleichungen

Man spricht von **separierbaren Differentialgleichungen 1. Ordnung**, wenn sie sich auf die Form

$$p(t) + \dot{x}q(x) = 0 \quad (5.1.2)$$

zurückführen lassen.

Dann genügt bereits eine beliebige Stammfunktion der Funktionen p und q (wobei im Einzelfall freilich stets auf deren Definitionsbereich und die damit verbundene Existenz der Integrale zu achten ist). Definieren wir nämlich die Stammfunktionen zu

$$P(t) = \int_{t_0}^t dt' p(t'), \quad Q(x) = \int_{x_0}^x dx' q(x), \quad (5.1.3)$$

können wir wegen des Hauptsatzes der Differential- und Integralrechnung für (5.1.2)

$$\dot{P}(t) + \dot{x} \frac{d}{dx} Q(x) = 0 \quad (5.1.4)$$

schreiben. Mit der Kettenregel wird dies zu

$$\frac{d}{dt} Q[x(t)] = -\dot{P}(t), \quad (5.1.5)$$

integrieren wir beide Gleichungen von $t = t_0$ bis t erhalten wir unter Berücksichtigung der Anfangsbedingung bereits die Lösung in impliziter Form

$$Q(x) = -P(t). \quad (5.1.6)$$

Diese Lösung müssen wir nun nur noch nach x auflösen, um die Lösung der Differentialgleichung zu erhalten. Da wir die Anfangsbedingung bereits eingearbeitet haben, ist auch diese dann automatisch erfüllt.

Beispiel: Gesetz vom radioaktiven Zerfall

Als *Beispiel* betrachten wir die Differentialgleichung des **radioaktiven Zerfalls**. Dazu gehen wir davon aus, dass die entsprechende Zerfallsrate (also die Anzahl der Atomkerne, die pro Zeiteinheit zerfällt) unabhängig von den äußeren Umständen (Temperatur, Druck usw.) und proportional zur Zahl $N(t)$ der zur Zeit t vorhandenen Kerne ist. Zur Zeit $t_0 = 0$ seien N_0 . In eine Gleichung gebracht, haben wir es also mit dem Anfangswertproblem einer Differentialgleichung 1. Ordnung zu tun,

$$\dot{N}(t) = -\lambda N(t), \quad \lambda = \text{const.} \quad N(0) = N_0. \quad (5.1.7)$$

Diese Gleichung ist tatsächlich vom separablen Typ, denn wir können sie in der Form

$$\lambda + \dot{N} \frac{1}{N} = 0 \quad (5.1.8)$$

schreiben. Bis auf die Bezeichnung der Variablen N statt x ist das in der Tat in der Form (5.1.2). Nach der allgemeinen Lösungsformel (5.1.6) benötigen wir die Stammfunktionen von $p(t) = \lambda$ und $q(N) = 1/N$:

$$P(t) = \int_0^t dt' \lambda = \lambda t, \quad N(t) = \int_{N_0}^N dN' \frac{1}{N'} = \ln\left(\frac{N}{N_0}\right). \quad (5.1.9)$$

Die Lösung lautet also gemäß (5.1.6)

$$\ln\left(\frac{N(t)}{N_0}\right) = -\lambda t \Rightarrow N(t) = N_0 \exp(-\lambda t). \quad (5.1.10)$$

Wir erhalten also das bekannte **Exponentialgesetz vom radioaktiven Zerfall**. Man bezeichnet als **Lebensdauer** τ des betreffenden Atomkerns die Zeit nach der die Anzahl der Kerne auf $N(\tau)/N_0 = e$ abgefallen ist. Aus unserer Lösung (5.1.10) folgt sofort, dass $\tau = 1/\lambda$ ist. Wir werden gleich noch ausrechnen, dass das in der Tat die mittlere Lebensdauer eines instabilen Kerns ist. Die **Halbwertszeit** $\tau_{1/2}$ gibt hingegen die Zeit an, nachdem die anfänglich vorhandene Menge an radioaktivem Material zur Hälfte zerfallen ist, d.h.

$$\ln\left(\frac{1}{2}\right) = -\ln 2 = -\lambda \tau_{1/2} \Rightarrow \tau_{1/2} = \frac{\ln 2}{\lambda} = \tau \ln 2 \approx 0,693 \tau. \quad (5.1.11)$$

Wir können das Zerfallsgesetz auch für einen einzelnen Atomkern als **Wahrscheinlichkeitsdichte** interpretieren. Dazu nehmen wir an, dass die Wahrscheinlichkeit, dass der Atomkern zur Zeit t in einem kleinen Zeitintervall $[t, t + dt]$ zerfällt unabhängig von der Vorgeschichte des Atomkerns ist, und dass die **Zerfallsrate** konstant λ ist, d.h. ist zur Zeit t der Atomkern mit Sicherheit vorhanden, ist die Wahrscheinlichkeit, dass er in dem kleinen Zeitintervall $[t, t + dt]$ zerfällt durch λdt gegeben. Sei nun $W(t)$ die Wahrscheinlichkeit, dass der Kern zur Zeit t vorhanden ist, so ist die Wahrscheinlichkeit, dass er innerhalb des Zeitintervalls $[t, t + dt]$ zerfällt offenbar $W(t)\lambda dt$, d.h. es gilt

$$W(t)\lambda dt = -dt \dot{W}(t) \Rightarrow \dot{W} = -\lambda W. \quad (5.1.12)$$

Dies ist natürlich dieselbe Differentialgleichung wie oben für N , d.h. es ist

$$W(t) = W_0 \exp(-\lambda t). \quad (5.1.13)$$

Soll nun der Kern zur Zeit $t = 0$ sicher vorhanden gewesen sein, gilt $W_0 = 1$. Es ist also

$$W(t) = \exp(-\lambda t). \quad (5.1.14)$$

Daraus ergibt sich die Wahrscheinlichkeitsdichte $w(t)$, d.h. die Wahrscheinlichkeit, dass der Kern zur Zeit t in einem kleinen Zeitintervall $[t, t + dt]$ zerfällt, sei $w(t)dt$. Aufgrund der obigen Überlegung gilt

$$w(t) = -\dot{W}(t) = \lambda \exp(-\lambda t). \quad (5.1.15)$$

Daraus können wir nun einige charakteristische Größen ausrechnen. Zum Ersten prüfen wir nach, dass w korrekt normiert ist. Ein anfangs vorhandenes Teilchen muss nämlich mit Sicherheit irgendwann einmal zerfallen, d.h. es sollte

$$\int_0^\infty dt w(t) = 1 \quad (5.1.16)$$

gelten. In der Tat ist

$$\int_0^\infty dt w(t) = \int_0^\infty dt \lambda \exp(-\lambda t) = -\exp(-\lambda t) \Big|_{t=0}^{t=\infty} = 0 - (-\exp 0) = 1. \quad (5.1.17)$$

Weiter wollen wir die mittlere Lebensdauer bestimmen, also

$$\langle t \rangle = \tau = \int_0^\infty dt t w(t) = \lambda \int_0^\infty dt t \exp(-\lambda t). \quad (5.1.18)$$

Hierzu nutzen wir einen Trick aus und definieren zunächst die Funktion (*Nachrechnen!*)

$$F(\lambda) = \int_0^\infty dt \exp(-\lambda t) = \frac{1}{\lambda}. \quad (5.1.19)$$

Nun darf man die Ableitung nach λ mit der Integration nach t vertauschen, d.h. es gilt

$$\frac{d}{d\lambda} F(\lambda) = \int_0^\infty dt \frac{d}{d\lambda} \exp(-\lambda t) = -\lambda \int_0^\infty dt t \exp(-\lambda t). \quad (5.1.20)$$

Vergleicht man dies mit (5.1.19), folgt

$$\tau = -\lambda F'(\lambda) = -\lambda \frac{d}{d\lambda} \lambda^{-1} = \frac{1}{\lambda}. \quad (5.1.21)$$

5.1.2 Homogene Differentialgleichungen 1. Ordnung

Angenommen, die Differentialgleichung besitzt die spezielle Form

$$\dot{x} = f(t, x), \quad (5.1.22)$$

wobei die Funktion f die **Homogenitätseigenschaft**

$$f(\lambda t, \lambda x) = f(t, x) \quad (5.1.23)$$

5.1. Differentialgleichungen 1. Ordnung

besitzt. Dann können wir durch die Substitution

$$x(t) = t\alpha(t) \quad (5.1.24)$$

die Differentialgleichung in eine separable Gleichung transformieren. In der Tat ist dann wegen der Homogenitätseigenschaft (5.1.23)

$$\dot{x} = \alpha + t\dot{\alpha} = f(t, t\alpha) = f(1, \alpha) = \tilde{f}(\alpha). \quad (5.1.25)$$

Etwas umgeformt erhalten wir die folgende Gleichung für α :

$$-t + \frac{\dot{\alpha}}{\tilde{f}(\alpha) - \alpha} = 0, \quad (5.1.26)$$

die tatsächlich von der separablen Form (5.1.2) ist.

5.1.3 Exakte Differentialgleichung 1. Ordnung

Angenommen die Differentialgleichung sei von der Form

$$\alpha(x, t)\dot{x} + \beta(x, t) = 0, \quad (5.1.27)$$

und es existiere eine Funktion $g(x, t)$, so dass

$$\alpha(x, t) = -\frac{\partial}{\partial t}g(x, t), \quad \beta(x, t) = -\frac{\partial}{\partial x}g(x, t). \quad (5.1.28)$$

Betrachten wir also $\vec{A} = (\alpha, \beta)$ als zweidimensionales Vektorfeld mit $\vec{\xi} = (x, t)$ als unabhängige kartesische Koordinaten, verlangen wir, dass g ein Potential von $\vec{A}$ ist. Nach dem **Lemma von Poincaré** (vgl. Abschnitt 3.11) ist das in einfach zusammenhängenden Gebieten der Fall, wenn $\partial_x \alpha = \partial_t \beta$ gilt, und das Potential ist dann durch das entsprechende Wegintegral gegeben (vgl. Abschnitt 3.9). Jedenfalls können wir dann (5.1.27) in der Form

$$\frac{d\vec{\xi}}{dt} \cdot \vec{\nabla} g = \dot{x}\partial_x g + t\partial_t g = -[\dot{x}\alpha(x, t) + \beta(t, x)] = \frac{d}{dt}g[x(t), t] \stackrel{(5.1.27)}{=} 0 \quad (5.1.29)$$

schreiben. Die allgemeine Lösung ist demnach in impliziter Form sofort durch

$$g[x(t), t] = C = \text{const.} \quad (5.1.30)$$

gegeben. Die Integrationskonstante C bestimmt sich aus der Anfangsbedingung $x(t_0) = x_0$.

5.1.4 Homogene lineare Differentialgleichung 1. Ordnung

Eine Differentialgleichung 1. Ordnung heißt **linear**, wenn sie von der Form

$$\dot{x} + p(t)x = q(t) \quad (5.1.31)$$

ist und **homogen**, wenn $q(t) = 0$ ist. Wir beschäftigen uns zuerst mit dem homogenen Fall

$$\dot{x} + p(t)x = 0. \quad (5.1.32)$$

Dividiert man diese Gleichung durch x , erkennt man, dass sie vom separablen Typ ist. Jedenfalls können wir sie auf die Form

$$\frac{d}{dt} \ln\left(\frac{x}{x_0}\right) = -p(t) \quad (5.1.33)$$

bringen. Berücksichtigen wir die Anfangsbedingung $x(t_0) = x_0$ und integrieren diese Gleichung bzgl. t , erhalten wir nach einer einfachen Umformung die Lösung in der Form

$$x(t) = x_0 \exp \left[- \int_{t_0}^t dt' p(t') \right]. \quad (5.1.34)$$

5.1.5 Inhomogene lineare Differentialgleichung 1. Ordnung

Kommen wir nun auf (5.1.31) mit $q \neq 0$ zurück. Zuerst bemerken wir, dass wegen der Linearität der Differentialgleichung die Differenz zweier Lösungen x_1 und x_2 die homogene Gleichung (5.1.33) löst:

$$\dot{x}_1(t) + p(t)x_1(t) = q(t), \quad \dot{x}_2(t) + p(t)x_2(t) = q(t) \Rightarrow \frac{d}{dt}(x_1 - x_2) + p(t)(x_1 - x_2) = 0. \quad (5.1.35)$$

Die allgemeine Lösung der inhomogenen Gleichung ist also durch die Summe aus der allgemeinen Lösung der homogenen Gleichung und einer beliebigen Lösung der inhomogenen Gleichung gegeben. Sei also $\tilde{x} \neq 0$ eine Lösung der homogenen Gleichung. Dann können wir eine spezielle Lösung der inhomogenen Gleichung finden, indem wir den

Ansatz der Variation der Konstanten

$$x(t) = y(t)\tilde{x}(t) \quad (5.1.36)$$

vornehmen. Wegen $\dot{\tilde{x}} + p\tilde{x} = 0$ folgt dann nämlich aus der Produktregel

$$\dot{x} + p(t)x = \dot{y}\tilde{x} + y\dot{\tilde{x}} + p(t)y\tilde{x} = \dot{y}\tilde{x} \stackrel{!}{=} q. \quad (5.1.37)$$

Damit finden wir die Lösung für y durch eine einfache Integration. Da wir nur irgendeine Lösung der inhomogenen Gleichung benötigen, können wir zusätzlich $y(t_0) = 0$ fordern. Dann folgt

$$y(t) = \int_{t_0}^t dt' \frac{q(t')}{\tilde{x}(t')}. \quad (5.1.38)$$

Dann wird nach der Überlegung oben das Anfangswertproblem der inhomogenen Gleichung mit $x(t_0) = x_0$ durch

$$x(t) = \tilde{x}(t) + \tilde{x}(t) \int_{t_0}^t dt' \frac{q(t')}{\tilde{x}(t')}, \quad (5.1.39)$$

wobei gemäß (5.1.34)

$$\tilde{x}(t) = x_0 \exp \left[- \int_{t_0}^t dt' p(t') \right]. \quad (5.1.40)$$

die Lösung der homogenen Gleichung (5.1.32) ist, die die Anfangsbedingung $\tilde{x}(t_0) = x_0$ erfüllt.

5.2 Lineare Differentialgleichungen 2. Ordnung

Lineare Differentialgleichungen besitzen eine besonders einfache Lösungsstruktur. In der Physik kann man auch oft kompliziertere Probleme durch lineare Differentialgleichungen nähern (für ein Beispiel s.u. den Ab-

schnitt zum harmonischen Oszillator als Näherung für das mathematische Pendel). Hier betrachten wir kurz die ganz allgemeine Struktur der allgemeinen Lösungen für lineare DGLn 2. Ordnung.

Die allgemeine lineare Differentialgleichung (DGL) 2. Ordnung lautet

$$\ddot{x}(t) + A(t)\dot{x}(t) + B(t)x(t) = C(t). \quad (5.2.1)$$

Dabei sind A , B und C vorgegebene Funktionen der unabhängigen Variablen t und $x(t)$ die gesuchte Funktion. Hier besprechen wir nur die wichtigsten Grundlagen über die Struktur der Lösungen solcher linearer Differentialgleichungen. Konkrete Beispiele liefern die in den nächsten Abschnitten behandelten harmonischen Oszillatoren verschiedener Art.

Man nennt die obige Differentialgleichung (5.2.1) **homogen**, wenn $C(t) = 0$ und entsprechend **inhomogen**, wenn $C(t) \neq 0$. Wir betrachten zuerst die Lösungsstruktur der

homogenen Gleichung

$$\ddot{x}(t) + A(t)\dot{x}(t) + B(t)x(t) = 0. \quad (5.2.2)$$

Da der Ableitungsoperator linear ist, d.h. für irgendwelche zwei Funktionen $x_1(t)$ und $x_2(t)$, die mindestens zweimal differenzierbar sind,

$$\frac{d}{dt}[C_1x_1(t) + C_2x_2(t)] = C_1\dot{x}_1 + C_2\dot{x}_2, \quad \frac{d^2}{dt^2}[C_1x_1(t) + C_2x_2(t)] = C_1\ddot{x}_1 + C_2\ddot{x}_2 \quad (5.2.3)$$

gilt, ist für zwei Lösungen x_1 und x_2 von (5.2.2) auch die **Linearkombination**

$$x(t) = C_1x_1(t) + C_2x_2(t) \quad (5.2.4)$$

mit $C_1, C_2 = \text{const}$ eine weitere Lösung. Es ist weiter klar, dass wir zur eindeutigen Festlegung der Lösung **Anfangsbedingungen** fordern müssen, d.h. wir verlangen von der Lösung x der DGL zusätzlich, dass sie und ihre erste Ableitung bei $t = 0$ bestimmte Werte annimmt:

$$x(0) \stackrel{!}{=} x_0, \quad \dot{x}(0) \stackrel{!}{=} v_0. \quad (5.2.5)$$

Nehmen wir an, wir hätten zwei Lösungen x_1 und x_2 gefunden, können wir versuchen, die Konstanten C_1 und C_2 in der Linearkombination (5.2.4) so zu bestimmen, dass diese Anfangsbedingungen (5.2.5) gelten, d.h. wir müssen das lineare Gleichungssystem

$$\begin{aligned} C_1x_1(0) + C_2x_2(0) &= x_0 \\ C_1\dot{x}_1(0) + C_2\dot{x}_2(0) &= v_0 \end{aligned} \quad (5.2.6)$$

nach C_1 und C_2 auflösen. Multiplizieren wir die erste Gleichung mit $\dot{x}_2(0)$ und die zweite mit $x_2(0)$ und subtrahieren die beiden entstehenden Gleichungen, finden wir

$$C_1[x_1(0)\dot{x}_2(0) - x_2(0)\dot{x}_1(0)] = \dot{x}_2(0)x_0 - x_2(0)v_0. \quad (5.2.7)$$

Wenn die eckige Klammer nicht verschwindet, können wir nach C_1 auflösen:

$$C_1 = \frac{\dot{x}_2(0)x_0 - x_2(0)v_0}{x_1(0)\dot{x}_2(0) - x_2(0)\dot{x}_1(0)}. \quad (5.2.8)$$

Unter derselben Voraussetzung können wir auf ähnliche Weise auch C_2 berechnen:

$$C_2 = \frac{v_0 x_1(0) - x_0 \dot{x}_1(0)}{x_1(0) \dot{x}_2(0) - x_2(0) \dot{x}_1(0)}. \quad (5.2.9)$$

Wir untersuchen nun noch, wann die Bedingung, dass der Nenner in (5.2.8) und (5.2.9) für die Lösungen x_1 und x_2 der DGL nicht verschwindet, erfüllt ist. Es handelt sich um die **Determinante der Koeffizientenmatrix** des linearen Gleichungssystems (5.2.6).

Um diesen Ausdruck näher zu untersuchen, definieren wir für die beiden Lösungen x_1 und x_2 die **Wronski-Determinante** genannte Größe

$$W(t) = \det \begin{pmatrix} x_1(t) & x_2(t) \\ \dot{x}_1(t) & \dot{x}_2(t) \end{pmatrix} = x_1(t) \dot{x}_2(t) - x_2(t) \dot{x}_1(t). \quad (5.2.10)$$

Berechnen wir die Zeitableitung, finden wir mit Hilfe der Produktregel nach einiger Rechnung

$$\dot{W}(t) = x_1(t) \ddot{x}_2(t) - x_2(t) \ddot{x}_1(t). \quad (5.2.11)$$

Jetzt verwenden wir, dass x_1 und x_2 Lösungen der DGL (5.2.2) sind und setzen

$$\ddot{x}_j(t) = -A(t) \dot{x}_j(t) - B(t) x_j(t), \quad j \in \{1, 2\} \quad (5.2.12)$$

in (5.2.11) ein. Das ergibt

$$\dot{W}(t) = -A(t) [x_1(t) \dot{x}_2(t) - x_2(t) \dot{x}_1(t)] = -A(t) W(t). \quad (5.2.13)$$

Dies ist eine lineare Differentialgleichung 1. Ordnung, die sich durch Trennen der Variablen lösen lässt. Teilen wir also (5.2.13) durch $W(t)$, erhalten wir

$$\frac{\dot{W}(t)}{W(t)} = -A(t). \quad (5.2.14)$$

Integration dieser Gleichung bzgl. t von $t = 0$ bis t , liefert

$$\ln \left(\frac{W(t)}{W(0)} \right) = - \int_0^t dt' A(t'). \quad (5.2.15)$$

Lösen wir dies nach $W(t)$ auf, finden wir schließlich

$$W(t) = W(0) \exp \left[- \int_0^t dt' A(t') \right]. \quad (5.2.16)$$

Das bedeutet aber, dass entweder $W(t) = 0 = \text{const}$ ist (nämlich wenn $W(0) = 0$) oder $W(t) \neq 0$ für alle $t \in \mathbb{R}$ gilt.

Untersuchen wir deshalb weiter, was es für die Lösungen x_1 und x_2 der DGL bedeutet, wenn $W(t) = 0$ für alle t gilt. Aus der Definition der Wronski-Determinante (5.2.10) folgt dann

$$x_1(t) \dot{x}_2(t) - x_2(t) \dot{x}_1(t) = 0 \Rightarrow \frac{\dot{x}_1(t)}{x_1(t)} = \frac{\dot{x}_2(t)}{x_2(t)}. \quad (5.2.17)$$

Auch diese Gleichung können wir wieder bzgl. t von 0 bis t integrieren, und das ergibt

$$\ln\left(\frac{x_1(t)}{x_1(0)}\right) = \ln\left(\frac{x_2(t)}{x_2(0)}\right) \Rightarrow x_2(t) = \frac{x_2(0)}{x_1(0)} x_1(t). \quad (5.2.18)$$

Das bedeutet aber, dass $W(0) = 0$ genau dann, wenn $x_2(t) = C x_1(t)$ mit $C = x_2(0)/x_1(0) = \text{const}$ ist.

Die Anfangsbedingungen (5.2.5) sind also durch die Linearkombination (5.2.4) genau dann immer erfüllbar, wenn die Lösungen x_1 und x_2 **linear unabhängig** sind und daher $W(0) \neq 0$ ist. Um also die **allgemeine Lösung der homogenen DGL** zu finden, müssen wir nur irgendwelche zwei linear unabhängigen Lösungen finden. Die allgemeine Lösung ist dann durch die allgemeine Linearkombination (5.2.4) gegeben.

Kommen wir nun auf die inhomogene Gleichung (5.2.1) zurück. Nehmen wir wieder an, dass $x_1^{(\text{inh})}(t)$ und $x_2^{(\text{inh})}(t)$ Lösungen dieser inhomogenen Gleichung ist. Wegen der Linearität der Ableitungsoperation und der Linearität der linken Seite der inhomogenen Gleichung erfüllt dann offenbar $x_1^{(\text{inh})}(t) - x_2^{(\text{inh})}(t)$ die *homogene DGL*. Das bedeutet aber, dass bei Kenntnis von zwei linear unabhängigen Lösungen $x_1^{(\text{hom})}(t)$ und $x_2^{(\text{hom})}(t)$ der *homogenen DGL* diese Differenz durch eine Linearkombination dieser Lösungen gegeben sein muss. Es ist also die *allgemeine* Lösung der inhomogenen Gleichung durch

$$x(t) = C_1 x_1^{(\text{hom})}(t) + C_2 x_2^{(\text{hom})}(t) + x_1^{(\text{inh})}(t) \quad (5.2.19)$$

gegeben. Haben wir also die allgemeine Lösung der homogenen Gleichung gefunden, genügt es, nur eine einzige spezielle Lösung der inhomogenen DGL zu kennen, um alle Lösungen in der Form (5.2.19) angeben zu können. Es ist klar, dass auch hier die Anfangsbedingungen (5.2.5) durch die entsprechende Berechnung der Integrationskonstanten C_1 und C_2 stets erfüllbar sind, wenn nur $x_1^{(\text{hom})}(t)$ und $x_2^{(\text{hom})}(t)$ linear unabhängig sind und also die Wronski-Determinante $W(0) \neq 0$ ist.

5.3 Der ungedämpfte harmonische Oszillator

Bei vielen typischen Bewegungsgleichungen der klassischen Mechanik tritt der Fall auf, dass ein Massepunkt sich in einem **Kräftepotential** bewegt. Wir betrachten eindimensionale Bewegungen entlang der x -Achse eines kartesischen Koordinatensystems. Dann ist die Kraft durch die Ableitung des Potentials gegeben:

$$F(x) = -V'(x). \quad (5.3.1)$$

Die Newtonsche Bewegungsgleichung für solch einen Massenpunkt lautet demnach

$$m\ddot{x} = -V'(x). \quad (5.3.2)$$

Um die Bahn der Bewegung als Funktion der Zeit zu erhalten, müssen wir also eine **Differentialgleichung** zweiter Ordnung lösen, d.h. wir suchen die Ortskoordinate x als Funktion von t . Dabei ergibt sich eine ganze Schar von Lösungen. Um die Bewegung des Massenpunktes eindeutig festzulegen, müssen wir noch **Anfangsbedingungen** fordern, d.h. wir müssen zu einem vorgegebenen Zeitpunkt, den wir bequemlichkeitshalber bei $t = 0$ wählen, **Ort und Geschwindigkeit** des Massenpunktes vorgeben. Wir verlangen also von der Lösung der Bewegungsgleichung (5.3.2), dass die Anfangsbedingungen

$$x(0) = x_0, \quad \dot{x}(0) = v_0 \quad (5.3.3)$$

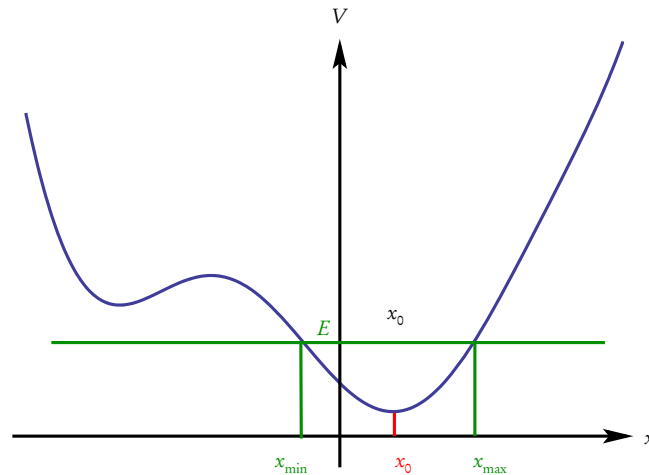


Abbildung 5.1: Bewegung in einer Potentialmulde.

erfüllt sind. Wir wissen bereits, dass für die Bewegungsgleichung (5.3.2) der **Satz von der Energieerhaltung** gilt, denn multiplizieren wir (5.3.2) mit $\dot{x}$ und bringen alle Ausdrücke auf die linke Seite der Gleichung, erhalten wir

$$m\dot{x}\ddot{x} + \dot{x}V'(x) = 0. \quad (5.3.4)$$

Es ist aber leicht zu sehen, dass dies eine totale Zeitableitung ist, denn es gilt

$$\frac{d}{dt}(\dot{x}^2) = 2\dot{x}\ddot{x}, \quad \frac{d}{dt}V(x) = \dot{x}V'(x). \quad (5.3.5)$$

Wir können also (5.3.4) in der Form

$$\frac{d}{dt} \left[\frac{m}{2} \dot{x}^2 + V(x) \right] = 0 \quad (5.3.6)$$

schreiben. Das bedeutet aber, dass der Ausdruck in den eckigen Klammern, die **Gesamtenergie** des Massenpunktes, für alle Lösungen der Bewegungsgleichung (5.3.2) zeitlich konstant ist:

$$E = \frac{m}{2} \dot{x}^2 + V(x) = \frac{m}{2} v_0^2 + V(x_0) = \text{const.} \quad (5.3.7)$$

Dabei haben wir die Anfangsbedingung (5.3.4) eingesetzt, um den Wert der Gesamtenergie zu bestimmen.

Nun ist die **kinetische Energie**

$$E_{\text{kin}} = \frac{m}{2} \dot{x}^2 \geq 0. \quad (5.3.8)$$

In Abb. 5.1 haben wir ein beliebiges Potential als Funktion der Ortskoordinate aufgezeichnet und den konstanten Wert der Gesamtenergie als horizontale Linie eingetragen. Da $E_{\text{kin}} \geq 0$ muss für alle Zeiten

$$E \geq V(x) \quad (5.3.9)$$

gelten. D.h. für eine durch die Anfangsbedingungen gegebene Gesamtenergie E kann sich das Teilchen nur dort aufhalten, wo das Potential unterhalb der Linie für die Gesamtenergie verläuft.

Nun kommt es oft vor, dass das Potential bei einer Stelle x_0 ein **lokales Minimum** aufweist. In diesem Minimum ist $V'(x_0) = 0$. Ist dann die Energie E so gewählt, dass die Linie $E = \text{const}$ das Potential an zwei Stellen $x_{\min}$ und $x_{\max}$ schneidet, muss das Teilchen in dem Bereich $[x_{\min}, x_{\max}]$ bleiben, denn aufgrund der Differentialgleichung muss die Ortskoordinate als Funktion der Zeit mindestens zweimal differenzierbar sein

und ist daher stetig. Der Massenpunkt kann also nicht über eine Potentialbarriere einfach in einen anderen Bereich springen, wo wieder $E > V(x)$ gilt. Das Teilchen ist also in dem besagten Intervall gefangen. Ist dieser Bereich nicht zu groß, reicht es weiter aus, das Potential um x_0 in eine **Potenzreihe** zu entwickeln und nur die Terme bis zur zweiten Ordnung mitzunehmen. Angenommen, das Potential ist mindestens dreimal stetig differenzierbar, können wir schreiben (**Taylor-Entwicklung**)

$$V(x) = V(x_0) + \frac{1}{2} V''(x_0)(x - x_0)^2 + \mathcal{O}[(x - x_0)^3]. \quad (5.3.10)$$

Dabei bedeutet das **Landau-Symbol** $\mathcal{O}[(x - x_0)^3]$, dass der nächste Term in der Potenzreihenentwicklung von der Größenordnung $(x - x_0)^3$ ist. Es ist klar, dass der Term linear zu $(x - x_0)$ verschwindet, weil voraussetzungsgemäß V an der Stelle x_0 ein Minimum besitzt. Außerdem nehmen wir an, dass $V''(x_0) > 0$ ist.

Für die Kraft folgt dann

$$F(x) = -V'(x) = -V''(x_0)(x - x_0) + \mathcal{O}[(x - x_0)^2]. \quad (5.3.11)$$

Für nicht zu große Abweichungen der Lage des Massenpunktes von x_0 können wir also näherungsweise die vereinfachte Bewegungsgleichung

$$m\ddot{x} = -D(x - x_0) \quad \text{mit} \quad D = V''(x_0) > 0. \quad (5.3.12)$$

betrachten.

Wählen wir das Koordinatensystem noch so, dass $x_0 = 0$ ist, erhalten wir die relativ einfache Bewegungsgleichung eines **harmonischen Oszillators**:

$$m\ddot{x} = -Dx. \quad (5.3.13)$$

Wir können diese Situation durch ein reales System in sehr guter Näherung realisieren, indem wir einen Massenpunkt an eine Feder hängen. Für nicht zu große Auslenkungen der Feder aus ihrer Gleichgewichtslage ist die von der Feder ausgeübte Kraft proportional zur Auslenkung ($|F_{\text{Feder}}| = D\Delta x$, wo Δx die Dehnung der Feder aus ihrer Ruhelage ist). Der Gleichgewichtspunkt $x_0 = 0$ ist dann dadurch gegeben, dass dort die Federkraft die Schwerkraft mg gerade kompensiert. Die Feder wirkt immer der Auslenkung entgegen, und die gesamte Kraft auf den Massenpunkt ist dann durch $F_x = -Dx$ gegeben. Dabei rührt das Vorzeichen in dieser Gleichung daher, dass die Feder immer der Auslenkung aus der Gleichgewichtslage entgegenwirkt.

Zur Lösung dieser Gleichung beachten wir, dass es sich um eine lineare homogene Differentialgleichung (DGL) zweiter Ordnung handelt. Wir können uns also auf die allgemeinen Betrachtungen im vorigen Abschnitt stützen. Aus den dortigen Überlegungen folgt, dass die allgemeine Lösung von der Form

$$x(t) = C_1 x_1(t) + C_2 x_2(t) \quad (5.3.14)$$

ist, wobei x_1 und x_2 irgendwelche zwei linear unabhängige Lösungen der Gleichung sind, d.h. es muss $x_1/x_2 \neq \text{const}$ sein, und beide Funktionen müssen die DGL lösen. Ein Blick auf (5.3.13) zeigt, dass ein Ansatz mit trigonometrischen Funktionen

$$x_1(t) = C_1 \cos(\omega_0 t), \quad x_2(t) = C_2 \sin(\omega_0 t) \quad (5.3.15)$$

erfolgsversprechend ist, denn es gilt

$$\dot{x}_1 = -C_1 \omega_0 \sin(\omega_0 t), \quad \ddot{x}_1 = -C_1 \omega_0^2 \cos(\omega_0 t) \quad (5.3.16)$$

und

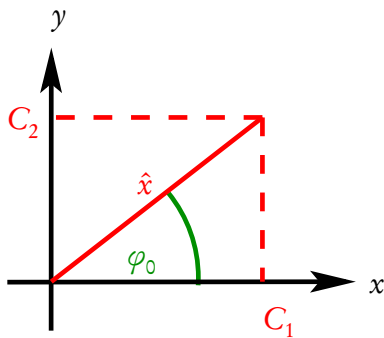
$$\dot{x}_2 = C_2 \omega_0 \cos(\omega_0 t), \quad \ddot{x}_2 = -C_2 \omega_0^2 \sin(\omega_0 t). \quad (5.3.17)$$

Setzt man diese Ansätze in (5.3.13) ein, erkennt man sofort, dass beides Lösungen der Differentialgleichung sind, und zwar für

$$\omega_0 = \sqrt{\frac{D}{m}}. \quad (5.3.18)$$

Die allgemeine Lösung der DGL (5.3.13) lautet also

$$x(t) = C_1 \cos(\omega_0 t) + C_2 \sin(\omega_0 t). \quad (5.3.19)$$



Um diese Lösung etwas einfacher analysieren zu können, bringen wir sie noch in eine etwas einfachere Form, und zwar versuchen wir Konstanten $\hat{x} \geq 0$ und φ_0 so zu bestimmen, dass

$$x(t) = \hat{x} \cos(\omega_0 t - \varphi_0) \quad (5.3.20)$$

gilt. Ausnutzen des Additionstheorems für den Kosinus liefert

$$x(t) = \hat{x} [\cos \varphi_0 \cos(\omega_0 t) + \sin \varphi_0 \sin(\omega_0 t)]. \quad (5.3.21)$$

Vergleicht man dies mit (5.3.19) folgt, dass dann

$$C_1 = \hat{x} \cos \varphi_0, \quad C_2 = \hat{x} \sin \varphi_0 \quad (5.3.22)$$

gelten muss. Es ist klar, dass man dies als Gleichung für die Komponenten eines Vektors (C_1, C_2) in der Ebene, ausgedrückt durch seine Polarkoordinaten $(\hat{x}, \varphi_0)$ ansehen kann (s. nebenstehende Abbildung). Quadriert man jedenfalls diese beiden Gleichungen, erhält man

$$\hat{x}^2 (\cos^2 \varphi_0 + \sin^2 \varphi_0) = \hat{x}^2 = C_1^2 + C_2^2 \Rightarrow \hat{x} = \sqrt{C_1^2 + C_2^2}. \quad (5.3.23)$$

Aus dem Bild liest man weiter ab, dass

$$\varphi_0 = \text{sign } C_2 \arccos\left(\frac{C_1}{\hat{x}}\right) \in (-\pi, \pi) \quad (5.3.24)$$

gegeben ist. Das einzige Problem mit dieser Formel ist, dass für $C_2 = 0$ und $C_1 \neq 0$ ein unbestimmtes Ergebnis herauskommt. Man hat dann aber $\cos \varphi_0 = C_1/|C_1| = \pm 1$. Für $C_1 > 0$ erhält man dann immer noch eindeutig $\varphi_0 = 0$. Für $C_1 < 0$ wären aber zwei Lösungen $\varphi_0 = \pm \pi$ korrekt. Man kann in diesem Fall einfach eine von beiden Möglichkeiten wählen, z.B. $\varphi_0 = +\pi$. Diese Gleichung zur Berechnung des Polarwinkels liefert im Gegensatz zu der in der Literatur oft zu findenden Formel " $\varphi_0 = \arctan(C_2/C_1)$ " stets den korrekten Winkel, ohne dass man sich genauere Gedanken machen muss, in welchem Quadranten der gerade betrachtete Punkt liegt.

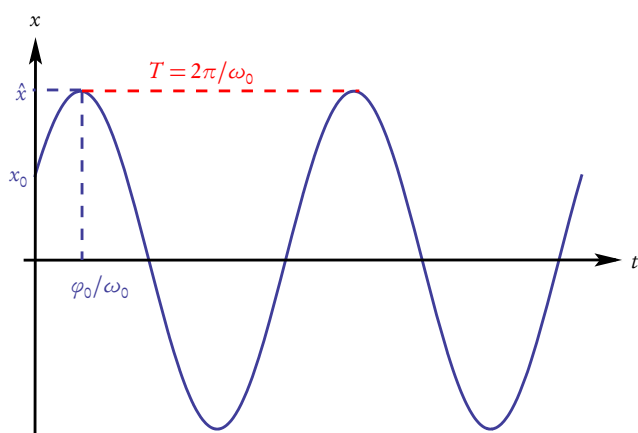


Abbildung 5.2: Lösung zum harmonischen Oszillator mit den Kenngrößen $\hat{x}$ (Amplitude), φ_0 (Anfangsphase) und T (Periodendauer).

Die Konstanten C_1 und C_2 in (5.3.19) lassen sich aus den Anfangsbedingungen (5.3.3) bestimmen. Wir verlangen also

$$\begin{aligned} x(0) &= C_1 = x_0, \\ \dot{x}(0) &= C_2 \omega_0 = v_0 \Rightarrow C_2 = \frac{v_0}{\omega_0}. \end{aligned} \quad (5.3.25)$$

Die Lösung für das Anfangswertproblem lautet also

$$x(t) = x_0 \cos(\omega_0 t) + \frac{v_0}{\omega_0} \sin(\omega_0 t). \quad (5.3.26)$$

Für die Lösungsform (5.3.20) ergibt sich aus (5.3.23) und (5.3.24)

$$\begin{aligned} \hat{x} &= \sqrt{x_0^2 + \frac{v_0^2}{\omega^2}}, \\ \varphi_0 &= \text{sign } v_0 \arccos\left(\frac{x_0}{\hat{x}}\right). \end{aligned} \quad (5.3.27)$$

Es ergibt sich also insgesamt eine um $\Delta t = \varphi_0/\omega_0$ entlang der t -Achse verschobene cos-Funktion mit der **Periodendauer** T , wobei

$$\omega_0 T = 2\pi \quad T = \frac{2\pi}{\omega_0}. \quad (5.3.28)$$

Die **Frequenz**, also die Anzahl der Schwingungen pro Zeiteinheit ist durch

$$f = \frac{1}{T} = \frac{\omega_0}{2\pi} \quad (5.3.29)$$

gegeben. Der Massepunkt schwingt zwischen den Werten $\pm \hat{x}$ hin und her. Diese Maximalabweichung von der Ruhelage $\hat{x}$ heißt **Amplitude** der Schwingung (vgl. Abb. 5.2). Wir bemerken, dass die Periodendauer der Schwingung unabhängig von der Amplitude ist. Man bezeichnet solche Schwingungen als **harmonische Schwingungen**. Schwingungen sind nur dann strikt harmonisch, wenn die Kraft exakt proportional zur Auslenkung von der Ruhelage ist. Für allgemeineren Kraftgesetze liegt dieser Fall nur näherungsweise für **kleine Amplituden** vor.

5.4 Der gedämpfte harmonische Oszillator

Im Allgemeinen wird die Bewegung eines Massepunktes auch irgendwelchen Reibungsprozessen unterliegen. Um zu sehen, welche Auswirkungen die Reibung besitzt, untersuchen wir den besonders einfachen Fall der **Stokesschen Reibung**, wo die Reibungskraft proportional zur Geschwindigkeit $v = \dot{x}$ ist und der Geschwindigkeit entgegengerichtet ist. Die Bewegungsgleichung lautet dann

$$m\ddot{x} = -Dx - \beta\dot{x}. \quad (5.4.1)$$

In die Normalform gebracht ergibt sich wieder eine lineare homogene Differentialgleichung:

$$\ddot{x} + 2\gamma\dot{x} + \omega_0^2 x = 0 \quad \text{mit} \quad \gamma = \frac{\beta}{2m}, \quad \omega_0^2 = \frac{D}{m}. \quad (5.4.2)$$

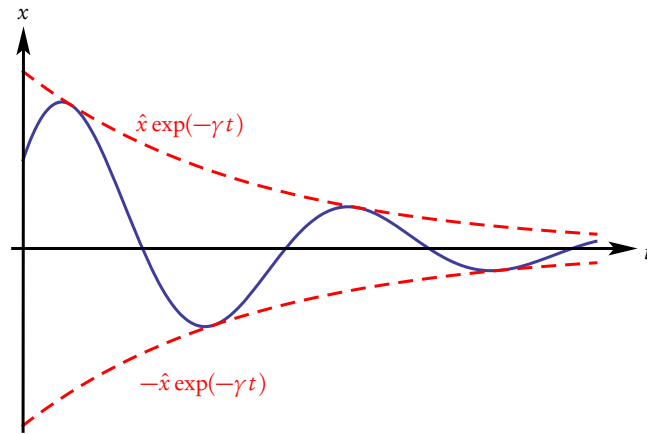


Abbildung 5.3: Lösung zum gedämpften harmonischen Oszillator. Für $\omega_0 > \gamma$ schwingt der Massenpunkt wieder sinusförmig auf Dämpfungsrate ist γ .

Wir werden gleich sehen, dass die willkürlich erscheinende Einführung des Faktors 2 im Reibungsterm einige Formeln ein wenig übersichtlicher macht.

Um die Gleichung zu vereinfachen, versuchen wir, den Term mit $\dot{x}$ zu eliminieren. Dazu machen wir den Ansatz

$$x(t) = \exp(\lambda t)y(t) \quad (5.4.3)$$

mit einer unbekannten Konstanten λ und einer neuen unbekannten Funktion $y(t)$. Mit (Nachrechnen)

$$\dot{x} = (\lambda y + \dot{y})\exp(\lambda t), \quad \ddot{x} = (\lambda^2 y + 2\lambda \dot{y} + \ddot{y})\exp(\lambda t) \quad (5.4.4)$$

folgt durch Einsetzen in (5.4.2)

$$\exp(\lambda t)[\ddot{y} + 2\lambda \dot{y} + \lambda^2 y + 2\gamma(\lambda y + \dot{y}) + \omega_0^2 y] = 0. \quad (5.4.5)$$

Da $\exp(\lambda t) \neq 0$ für alle $t \in \mathbb{R}$, muss der Ausdruck in der eckigen Klammer verschwinden. Weiter hebt sich der Term mit $\dot{y}$ weg, wenn wir $\lambda = -\gamma$ setzen. Dann vereinfacht sich (5.4.6) zu

$$\ddot{y} + (\omega_0^2 - \gamma^2)y = 0. \quad (5.4.6)$$

Um diese Gleichung zu lösen, unterscheiden wir drei Fälle:

1. $\omega_0 > \gamma$ (Schwingfall),
2. $\omega_0 = \gamma$ (aperiodischer Grenzfall),
3. $\omega_0 < \gamma$ (Kriechfall).

5.4.1 Schwingfall ($\omega_0 > \gamma$)

In diesem Fall ist die Gleichung (5.4.6) offenbar die Bewegungsgleichung eines harmonischen Oszillators mit der Kreisfrequenz

$$\omega = \sqrt{\omega_0^2 - \gamma^2} > 0, \quad (5.4.7)$$

und die allgemeine Lösung lautet

$$y(t) = C_1 \cos(\omega t) + C_2 \sin(\omega t), \quad (5.4.8)$$

wobei C_1 und C_2 zwei Integrationskonstanten sind.

Die allgemeine Lösung unseres ursprünglichen Problems lautet demnach gemäß (5.4.3) mit $\lambda = -\gamma$

$$x(t) = \exp(-\gamma t) [C_1 \cos(\omega t) + C_2 \sin(\omega t)]. \quad (5.4.9)$$

Wir wollen nun noch die Integrationskonstanten durch die **Anfangsbedingungen**

$$x(0) = x_0, \quad \dot{x}(0) = v_0 \quad (5.4.10)$$

ausdrücken. Es gilt (*nachrechnen!*)

$$\dot{x}(t) = \exp(-\gamma t) [-\gamma(C_1 \cos(\omega t) + C_2 \sin(\omega t)) - C_1 \omega \sin(\omega t) + C_2 \omega \cos(\omega t)]. \quad (5.4.11)$$

Damit folgt aus den Anfangsbedingungen

$$x_0 = C_1, \quad v_0 = -\gamma C_1 + \omega C_2 \Rightarrow C_2 = \frac{v_0 + \gamma x_0}{\omega} \quad (5.4.12)$$

und die Lösung des Anfangswertproblems ergibt sich gemäß (5.4.9) zu

$$x(t) = \exp(-\gamma t) \left[x_0 \cos(\omega t) + \frac{v_0 + \gamma x_0}{\omega} \sin(\omega t) \right]. \quad (5.4.13)$$

Analog zu unserem Vorgehen oben beim ungedämpften Oszillator können wir die eckige Klammer auch in Form einer einzelnen Kosinus-Funktion gemäß

$$x(t) = \exp(-\gamma t) \hat{x} \cos(\omega t - \varphi_0) \quad (5.4.14)$$

schreiben. Dabei ist

$$\hat{x} = \frac{\sqrt{x_0^2 \omega^2 + (v_0 + \gamma x_0)^2}}{\omega}, \quad \varphi_0 = \text{sign}(v_0 + \gamma x_0) \arccos\left(\frac{x_0}{\hat{x}}\right). \quad (5.4.15)$$

Wir haben also insgesamt eine periodische Schwingung mit der durch die Dämpfung verringerten Kreisfrequenz $\omega = \sqrt{\omega_0^2 - \gamma^2}$, deren Amplitude exponentiell abfällt. In der **Dämpfungszeit** $t_d = 1/\gamma$ verringert sich die Amplitude um einen Faktor $1/e = \exp(-1) \approx 1/2,718$. Der Massepunkt bewegt sich stets innerhalb der Einhüllenden $\pm \hat{x} \exp(-\gamma t)$ (vgl. Abb 5.3).

5.4.2 Kriechfall ($\omega_0 < \gamma$)

In diesem Fall lautet die Differentialgleichung (5.4.6)

$$\ddot{y} = +\Gamma^2 y \quad \text{mit} \quad \Gamma = \sqrt{\gamma^2 - \omega_0^2} > 0. \quad (5.4.16)$$

Die allgemeine Lösung der ursprünglichen Bewegungsgleichung ergibt sich in enger Analogie zur Rechnung beim Schwingfall dann zu (*nachrechnen!*)

$$x(t) = \exp(-\gamma t) \left[x_0 \cosh(\Gamma t) + \frac{v_0 + \gamma x_0}{\Gamma} \sinh(\Gamma t) \right]. \quad (5.4.17)$$

Dabei haben wir die Integrationskonstanten bereits durch die Anfangsbedingungen (5.4.10) ausgedrückt. Da $\Gamma < \gamma$ ist, ist die Bewegung für große Zeiten ($t \rightarrow \infty$) stets mit dem Faktor $\exp[-(\gamma - \Gamma)t]$ gedämpft (*warum*).

5.4.3 Aperiodischer Grenzfall ($\omega_0 = \gamma$)

Hier vereinfacht sich die Differentialgleichung (5.4.6) zu

$$\ddot{y} = 0 \Rightarrow y(t) = C_1 + C_2 t. \quad (5.4.18)$$

Die Lösung der ursprünglichen Gleichung ist also gemäß (5.4.9)

$$x(t) = \exp(-\gamma t)(C_1 + C_2 t). \quad (5.4.19)$$

Es ist weiter

$$\dot{x}(t) = \exp(-\gamma t)[- \gamma(C_1 + C_2 t) + C_2]. \quad (5.4.20)$$

Die Anfangsbedingungen (5.4.10) verlangen

$$x(0) = x_0 = C_1, \quad \dot{x}(0) = v_0 = C_2 - \gamma C_1 \Rightarrow C_2 = \frac{v_0 + \gamma x_0}{\gamma}, \quad (5.4.21)$$

und die Lösung des Anfangswertproblems lautet demnach

$$x(t) = \exp(-\gamma t) \left[x_0 + \frac{v_0 + \gamma x_0}{\gamma} t \right]. \quad (5.4.22)$$

5.5 Der getriebene gedämpfte Oszillator

Wir schließen unsere Betrachtung der harmonischen Oszillatoren mit der Behandlung der Bewegungsgleichung für den Fall, dass zusätzlich zur Reibungs- und harmonischen Kraft noch eine zeitabhängige äußere **treibende Kraft** an dem Massenpunkt angreift. Dabei beschränken wir uns auf den Fall einer **harmonischen Zeitabhängigkeit der äußeren Kraft**. Die Bewegungsgleichung lautet dann

$$m\ddot{x} = -m\omega_0^2 x - 2m\gamma\dot{x} + mA \cos(\Omega t). \quad (5.5.1)$$

Dies in die Normalform für lineare Differentialgleichungen 2. Ordnung gebracht liefert die **inhomogene Gleichung**

$$\ddot{x} + 2\gamma\dot{x} + \omega_0^2 x = A \cos(\Omega t). \quad (5.5.2)$$

Wir bemerken als erstes, dass wegen der Linearität des Differentialoperators auf der linken Seite die Differenz zweier Lösungen dieser inhomogenen Gleichung wieder die homogene Gleichung löst. Die allgemeine Lösung der inhomogenen Gleichung ist also durch die Summe aus der allgemeinen Lösung der homogenen Gleichung und einer beliebigen speziellen Lösung der inhomogenen Gleichung gegeben:

$$x(t) = C_1 x_1^{(\text{hom})}(t) + C_2 x_2^{(\text{hom})}(t) + x^{(\text{inh})}(t). \quad (5.5.3)$$

Dabei sind x_1 und x_2 beliebige zueinander linear unabhängige Lösungen der homogenen Differentialgleichung, die wir im vorigen Abschnitt für die drei Fälle (Schwingfall, Kriechfall, aperiodischer Grenzfall) gefunden haben:

$$\begin{cases} x_1(t) = \exp(-\gamma t) \cos(\omega t), & x_2(t) = \exp(-\gamma t) \sin(\omega t) & \text{für } \omega_0 > \gamma, \\ x_1(t) = \exp(-\gamma t) \cosh(\Gamma t), & x_2(t) = \exp(-\gamma t) \sinh(\Gamma t) & \text{für } \omega_0 < \gamma, \\ x_1(t) = \exp(-\gamma t), & x_2(t) = t \exp(-\gamma t) & \text{für } \omega_0 = \gamma. \end{cases} \quad (5.5.4)$$

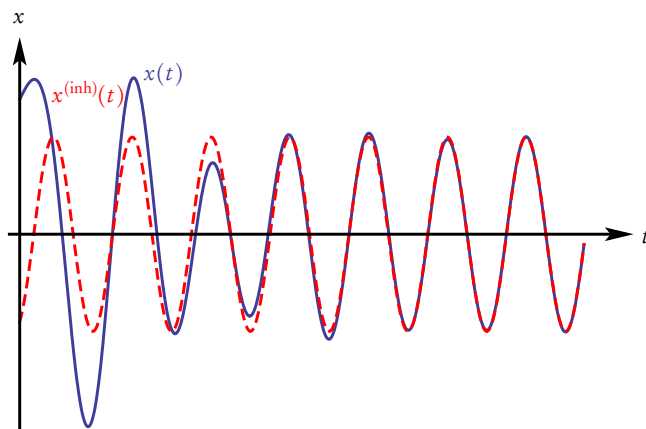


Abbildung 5.4: Lösung zum getriebenen gedämpften harmonischen Oszillator im Schwingfall $\omega_0 > \gamma$. Für $t \gg 1/\gamma$ werden die Eigenschwingungen, also der Anteil der Lösung der homogenen Gleichung in (5.5.3) merklich gedämpft, und die Bewegung geht in den durch die spezielle Lösung der inhomogenen Schwingung eingeschwungenen Zustand über.

Dabei ist im Schwingfall $\omega = \sqrt{\omega_0^2 - \gamma^2}$ und im Kriechfall $\Gamma = \sqrt{\gamma^2 - \omega_0^2}$.

Diese Lösungen werden allesamt für $t \gg 1/\gamma$ (bzw. im Kriechfall ($\omega_0 < \gamma$) für $t \gg 1/(\gamma - \Gamma)$) exponentiell weggedämpft werden. Für große Zeiten wird die Lösung also durch die spezielle Lösung der inhomogenen Gleichung dominiert. Man spricht vom **eingeschwungenen Zustand**, und wir interessieren uns im Folgenden für diesen Zustand.

5.5.1 Spezielle Lösung der inhomogenen Gleichung

Da also die „Eigenbewegungen“ des Massenpunktes für alle Fälle nach hinreichend langer Zeit abgeklungen sind, verbleibt für große Zeiten nur die spezielle Lösung der inhomogenen Gleichung. Es liegt nahe, anzunehmen, dass diese allein von der äußeren treibenden Kraft bestimmte Bewegung eine harmonische Bewegung mit der Kreisfrequenz der äußeren Kraft Ω sein wird. Dabei wird i.a. eine Phasenverschiebung zwischen der äußeren Kraft und der Bewegung des Massenpunktes vorliegen. Dies wird aber durch den Ansatz

$$x(t) = C_1 \cos(\Omega t) + C_2 \sin(\Omega t), \quad C_1, C_2 = \text{const} \quad (5.5.5)$$

erfüllt. Um die Konstanten C_1 und C_2 zu finden, setzen wir den Ansatz in (5.5.2) ein. Dies liefert (nachrechnen)

$$\cos(\Omega t)[(\omega_0^2 - \Omega^2)C_1 + 2\gamma\Omega C_2] + \sin(\Omega t)[-2\gamma\Omega C_1 + (\omega_0^2 - \Omega^2)C_2] = A \cos(\Omega t). \quad (5.5.6)$$

Es ist klar, dass in der Tat der Ansatz (5.5.5) eine Lösung liefert, wenn die Koeffizienten der Winkelfunktionen auf der linken und rechten Seite von (5.5.6) übereinstimmen, was auf das lineare Gleichungssystem

$$(\omega_0^2 - \Omega^2)C_1 + 2\gamma\Omega C_2 = A, \quad (5.5.7)$$

$$-2\gamma\Omega C_1 + (\omega_0^2 - \Omega^2)C_2 = 0 \quad (5.5.8)$$

zur Bestimmung der Konstanten C_1 und C_2 führt. Wir können das Gleichungssystem leicht lösen, indem wir (5.5.7) mit $2\gamma\Omega$ und (5.5.8) mit $(\omega_0^2 - \Omega^2)$ multiplizieren und die entstehenden Gleichungen addieren. Dies liefert dann

$$C_2 = \frac{2\gamma\Omega}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} A. \quad (5.5.9)$$

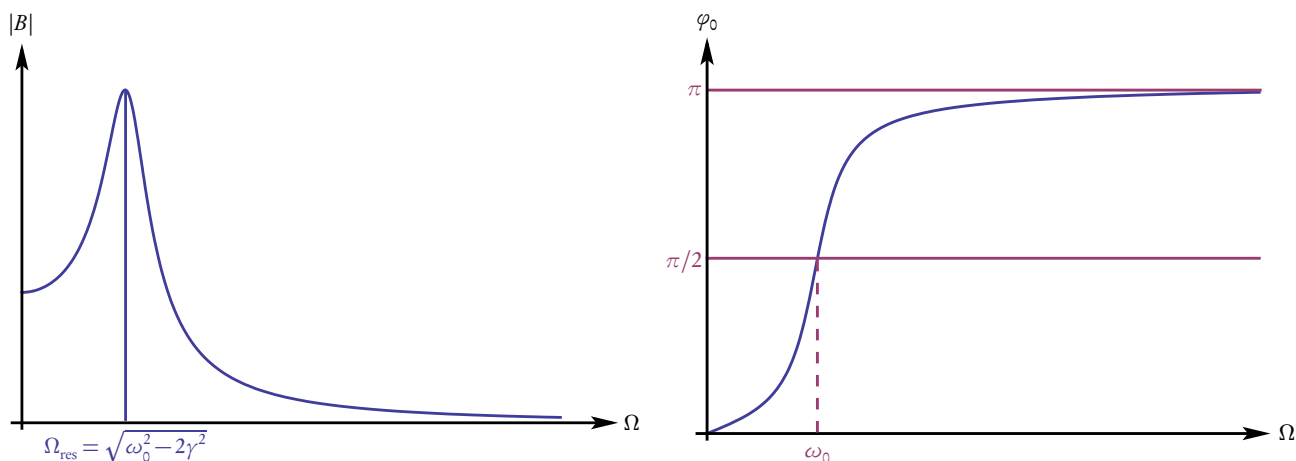


Abbildung 5.5: Amplitudenfaktor (links) $|B|$ und Phasenverschiebung φ_0 für den eingeschwungenen Zustand (5.5.12).

Um auch C_1 zu finden, multiplizieren wir (5.5.7) mit $(\omega_0^2 - \Omega^2)$ und (5.5.8) mit $(-2\gamma\Omega)$ und addieren wiederum die entstehenden Gleichungen, woraus sofort

$$C_1 = \frac{\omega_0^2 - \Omega^2}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} A \quad (5.5.10)$$

resultiert. Damit haben wir die gewünschte spezielle Lösung der inhomogenen Differentialgleichung gefunden.

Wir können (5.5.5) wieder in die alternative Form

$$x(t) = AB \cos(\Omega t - \varphi) \quad (5.5.11)$$

bringen. Dieselbe Methode wie beim ungedämpften Oszillator liefert

$$B = \sqrt{C_1^2 + C_2^2} = \frac{1}{\sqrt{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2}}, \quad \varphi = + \arccos\left(\frac{\omega_0^2 - \Omega^2}{\sqrt{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2}}\right). \quad (5.5.12)$$

Im eingeschwungenen Zustand schwingt also der Massepunkt mit derselben Frequenz wie die äußere harmonische Kraft, und zwischen der Kraft und der Schwingung des Massenpunktes besteht eine *stets positive* Phasenverschiebung φ_0 (vgl. Abb 5.5, rechts). Als Funktion von Ω ist φ_0 monoton wachsend. Für $\Omega = \omega_0$ wird $\varphi_0 = \pi/2$, und für $\Omega \rightarrow \infty$ strebt $\varphi_0 \rightarrow \pi$.

Bei vorgegebener Amplitude der äußeren Kraft mA ist die Amplitude des eingeschwungenen Zustands durch B gemäß (5.5.12) gegeben. In Abb. 5.5 (links) haben wir diesen Proportionalitätsfaktor als Funktion der Kreisfrequenz der antreibenden Kraft Ω geplottet. Er weist in dem gewählten Fall ein ausgeprägtes Maximum bei einer **Resonanzfrequenz** Ω_{res} auf, die wir im nächsten Abschnitt ausrechnen wollen.

5.5.2 Amplitudenresonanzfrequenz

Die Lösung (5.5.12) zeigt, dass die Amplitude der Schwingung im eingeschwungenen Zustand zur Amplitude der äußeren Kraft proportional ist und der Proportionalitätsfaktor B gemäß (5.5.12) durch die Parameter des

5.5. Der getriebene gedämpfte Oszillator

gedämpften freien Oszillators, also ω_0 und γ und die Kreisfrequenz Ω der äußeren Kraft allein bestimmt ist (s. Abb 5.5, links).

Wir untersuchen nun, für welches Ω dieser Proportionalitätsfaktor maximal wird, d.h. für welche Frequenz der äußeren Kraft die Amplitude der Schwingung bei festgehaltener Amplitude der äußeren Kraft am größten wird.

Dazu müssen wir das Minimum des Ausdrucks unter der Wurzel in (5.5.12) suchen, d.h. wir müssen die Funktion

$$f(\Omega) = (\Omega^2 - \omega_0^2)^2 + 4\gamma^2\Omega^2 \quad (5.5.13)$$

untersuchen. Dazu bilden wir die Ableitung

$$\frac{d}{d\Omega}f(\Omega) = f'(\Omega) = 4\Omega(\Omega^2 - \omega_0^2 + 2\gamma^2). \quad (5.5.14)$$

Mögliche Extrema ergeben sich als die Nullstellen dieser Ableitung. Offenbar ist die Lösung entweder $\Omega = 0$, d.h. es wirkt eine zeitlich konstante äußere Kraft. Dieser Fall interessiert uns hier aber nicht. Offenbar ergibt sich eine positive Nullstelle von (5.5.14) nur für $\omega_0^2 > 2\gamma^2$. Dann ist die Nullstelle bei der **Resonanzfrequenz**

$$\Omega_{\text{res}} = \sqrt{\omega_0^2 - 2\gamma^2} \quad (5.5.15)$$

Weiter gilt

$$f''(\Omega) = 4(2\gamma^2 - \omega_0^2 + 3\Omega^2) \quad (5.5.16)$$

und $f''(\Omega_{\text{res}}) = 8(\omega_0^2 - 2\gamma^2) > 0$, und es liegt demnach ein Minimum für f und also ein Maximum der Amplitude vor.

Man nennt daher Ω_{res} die **Amplitudenresonanzfrequenz**, weil dies die Frequenz ist, bei der die Amplitude der eingeschwungenen Bewegung maximal wird. Interessanterweise ist sie weder durch die Eigenfrequenz des ungedämpften Oszillators ω_0 noch durch die Eigenfrequenz der gedämpften Schwingung $\omega = \sqrt{\omega_0^2 - \gamma^2}$ im Schwingfall $\omega_0 > \gamma$ gegeben sondern durch die kleinere Amplitudenresonanzfrequenz (5.5.15).

5.5.3 Energieresonanz

Eine andere Frage ist, bei welcher Frequenz der äußeren Kraft, diese die **größte mittlere Leistung** aufbringen muss, um den Massenpunkt in dem erzwungenen eingeschwungenen Schwingungszustand zu halten. Dazu berechnen wir die über eine Periode $T = 2\pi/\Omega$ gemittelte Leistung der äußeren Kraft

$$\bar{P} = \frac{1}{T} \int_0^T dt m A \cos(\Omega t) \dot{x}^{(\text{inh})}(t). \quad (5.5.17)$$

Nun ist gemäß (5.5.11)

$$P(t) = m A \cos(\Omega t) \dot{x}^{(\text{inh})}(t) = -m A^2 |B| \Omega \cos(\Omega t) \sin(\Omega t - \varphi_0). \quad (5.5.18)$$

Mit dem Additionstheorem für den Sinus ist

$$P(t) = -m A^2 |B| \Omega \cos(\Omega t) [\sin(\Omega t) \cos \varphi_0 - \cos(\Omega t) \sin \varphi_0]. \quad (5.5.19)$$

Zur einfacheren Integration bemerken wir, dass

$$\cos(\Omega t) \sin(\Omega t) = \frac{1}{2} \sin(2\Omega t), \quad \cos^2(\omega t) = \frac{1}{2} [1 + \cos(2\Omega t)] \quad (5.5.20)$$

gilt, was man sofort durch Anwendung der Additionstheoreme nachweist. Da weiter

$$\int_0^T dt \sin(2\Omega t) = -\frac{\cos(2\Omega t)}{2\Omega} \Big|_{t=0}^{t=T} = 0, \quad \int_0^T dt \cos(2\Omega t) = \frac{\sin(2\Omega t)}{2\Omega} \Big|_{t=0}^{t=T} = 0 \quad (5.5.21)$$

ist, erhalten wir also durch Einsetzen von (5.5.19) in (5.5.17) für die mittlere Leistung

$$\bar{P} = \frac{1}{2} m A^2 B \Omega \sin \varphi_0. \quad (5.5.22)$$

Aus (5.5.12) folgt, dass $\varphi_0 \in [0, \pi]$ und also $\sin \varphi_0 > 0$ und damit

$$\sin \varphi_0 = \sqrt{1 - \cos^2 \varphi_0} = \frac{2\gamma\Omega}{\sqrt{(\Omega^2 - \omega_0^2)^2 + 4\gamma^2\Omega^2}} = 2\gamma\Omega B \quad (5.5.23)$$

ist. Somit wird (5.5.22)

$$\bar{P} = m A^2 B^2 \gamma \Omega^2. \quad (5.5.24)$$

Wir fragen nun, bei welcher Kreisfrequenz Ω der äußeren Kraft bei vorgegebenen Parametern des Oszillators und der Amplitude A der äußeren Kraft maximal wird. Man spricht in diesem Fall von **Energieresonanz**, denn dort ist die mittlere Leistungsaufnahme des Massenpunktes aus der äußeren Kraft maximal. Wir haben also diesmal das Maximum der Funktion

$$g(\Omega) = \Omega^2 |B|^2 = \frac{\Omega^2}{(\Omega^2 - \omega_0^2)^2 + 4\gamma^2\Omega^2} \quad (5.5.25)$$

zu suchen. Nach einiger Rechnung findet man für die Ableitung

$$g'(\Omega) = \frac{2\Omega(\Omega^4 - \omega_0^4)}{[(\Omega^2 - \omega_0^2)^2 + 4\gamma^2\Omega^2]^2}. \quad (5.5.26)$$

Offenbar liegt bei $\Omega = 0$ ein Minimum vor, denn dort ist g' lokal monoton fallend. Ein Maximum erhalten wir bei $\Omega = \omega_0$, denn dort ist g' offenbar lokal monoton wachsend.

Die größte mittlere Leistung wird also vom Oszillator aufgenommen, wenn $\Omega = \omega_0$ ist, also die Schwingungsfrequenz der äußeren Kraft der Eigenfrequenz des *ungedämpften* Oszillators entspricht. Wie (5.5.12) zeigt, ist dort gerade die Phasenverschiebung der eingeschwungenen Bewegung gegenüber der Phase der antreibenden Kraft $\varphi_0 = \pi/2$.

Wir bemerken noch, dass für verschwindende Dämpfung $\gamma = 0$ bei $\Omega = \omega_0$ der Faktor $|B|$ unendlich wird. Das ist die sogenannte **Resonanzkatastrophe**.

5.5.4 Lösung des Anfangswertproblems

Wir kommen schließlich auf die Lösung des Anfangswertproblems für den getriebenen harmonischen Oszillators zurück, die zur vollständigen Beschreibung der Bewegung bei beliebig vorgegebenen Anfangsbedingungen (5.3.3) dient und nicht nur den eingeschwungenen Zustand liefert. Man spricht auch vom **Einschwingvorgang**.

Wir müssen nur für die allgemeine Lösung (5.5.3) die Integrationskonstanten C_1 und C_2 aus den Anfangsbedingungen (5.4.1) durch Lösung des entsprechenden linearen Gleichungssystems zu bestimmen. Wir geben das Ergebnis für die oben diskutierten drei Fälle an

Schwingfall ($\omega_0 > \gamma$):

$$\begin{aligned} x(t) = & \left[x_0 - \frac{A(\omega_0^2 - \Omega^2)}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} \cos(\omega t) \right] \exp(-\gamma t) \\ & + \frac{1}{\omega} \left[v_0 + \gamma x_0 - \frac{A\gamma(\Omega^2 + \omega_0^2)}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} \right] \sin(\omega t) \exp(-\gamma t) \\ & + A \frac{\omega_0^2 - \Omega^2}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} \cos(\Omega t) + \frac{2A\gamma\Omega}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} \sin(\Omega t). \end{aligned} \quad (5.5.27)$$

Kriechfall ($\omega_0 < \gamma$): Setzen wir zur Abkürzung $\alpha = \sqrt{\gamma^2 - \omega_0^2}$, erhalten wir für diesen Fall

$$\begin{aligned} x(t) = & \left[x_0 - \frac{A(\omega_0^2 - \Omega^2)}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} \cosh(\alpha t) \right] \exp(-\gamma t) \\ & + \frac{1}{\alpha} \left[v_0 + \gamma x_0 - \frac{A\gamma(\Omega^2 + \omega_0^2)}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} \right] \sinh(\alpha t) \exp(-\gamma t) \\ & + A \frac{\omega_0^2 - \Omega^2}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} \cos(\Omega t) + \frac{2A\gamma\Omega}{(\omega_0^2 - \Omega^2)^2 + 4\gamma^2\Omega^2} \sin(\Omega t). \end{aligned} \quad (5.5.28)$$

Zu dieser Lösung können wir auch wieder gelangen, indem wir in (5.5.27) $\omega = i\alpha$ setzen und die in Formeln (4.2.14)

$$\cos(iz) = \cosh z, \quad \sin(iz) = i \sinh z \quad (5.5.29)$$

verwenden.

Aperiodischer Grenzfall ($\omega_0 = \gamma$):

$$\begin{aligned} x(t) = & \exp(-\gamma t) \left[x_0 - \frac{A(\gamma^2 - \Omega^2)}{(\gamma^2 + \Omega^2)^2} + \left(v_0 + \gamma x_0 - \frac{A\gamma}{\gamma^2 + \Omega^2} \right) t \right] \\ & + A \frac{(\gamma^2 - \Omega^2) \cos(\Omega t) + 2\gamma\Omega \sin(\Omega t)}{(\Omega^2 + \gamma^2)^2}. \end{aligned} \quad (5.5.30)$$

5.5.5 Resonant angetriebener ungedämpfter Oszillator

Wie wir oben gesehen haben, müssen wir den Fall des angetriebenen ungedämpften harmonischen Oszillators gesondert behandeln, wenn die Frequenz Ω der antreibenden Kraft der Eigenfrequenz ω_0 des Oszillators entspricht, da dann B wegen $\gamma = 0$ für $\Omega = \omega_0$ aufgrund von (5.5.12) nicht definiert ist, also der Ansatz (5.5.5) nicht zum Ziel führt.

Wir suchen also eine beliebige Lösung der inhomogenen Gleichung

$$\ddot{x} + \omega_0^2 x = a \cos(\omega_0 t). \quad (5.5.31)$$

Dazu machen wir den modifizierten Ansatz

$$x(t) = y(t) [C_1 \cos(\omega_0 t) + C_2 \sin(\omega_0 t)] \quad (5.5.32)$$

Die Ableitungen sind

$$\begin{aligned} \dot{x}(t) &= \dot{y}(t) [C_1 \cos(\omega_0 t) + C_2 \sin(\omega_0 t)] + y(t) [-C_1 \omega_0 \sin(\omega_0 t) + C_2 \omega_0 \cos(\omega_0 t)], \\ \ddot{x}(t) &= [\ddot{y}(t) - \omega_0^2 y(t)] [C_1 \cos(\omega_0 t) + C_2 \sin(\omega_0 t)] + 2\dot{y}(t) [-C_1 \omega_0 \sin(\omega_0 t) + C_2 \omega_0 \cos(\omega_0 t)]. \end{aligned} \quad (5.5.33)$$

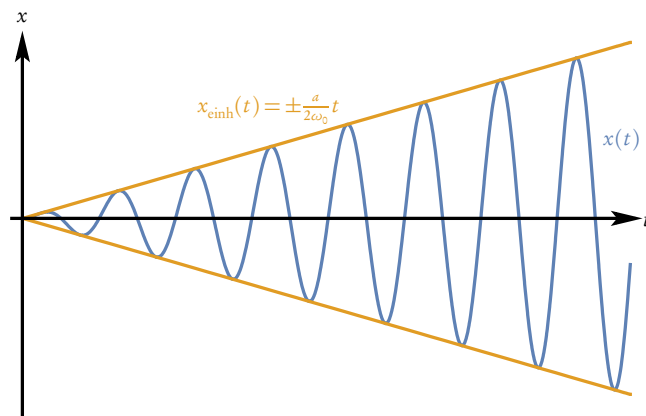


Abbildung 5.6: Lösung zum resonant angetriebenen, ungedämpften harmonischen Oszillator mit den Anfangsbedingungen $x(0) = 0$, $\dot{x}(0) = 0$. Dann ergibt sich die Lösung zu (5.5.36). Die Einhüllende ist durch die Geraden $x_{\text{einh}}(t) = \pm \frac{a}{2\omega_0}t$ gegeben.

Setzen wir dies in (5.5.31) ein, erhalten wir

$$2\dot{y}(t)[-C_1\omega_0\sin(\omega_0 t) + C_2\omega_0\cos(\omega_0 t)] = a\cos(\omega_0 t). \quad (5.5.34)$$

Um diese Gleichung zu erfüllen, muss offenbar $C_1 = 0$ und

$$2C_2\omega_0\dot{y} = a \Rightarrow \dot{y} = \frac{a}{2C_2\omega_0} \Rightarrow y(t) = \frac{a}{2\omega_0 C_2}t + C_3 \quad (5.5.35)$$

sein. Damit erhalten wir durch Einsetzen in (5.5.32) als eine spezielle Lösung der inhomogenen Gleichung

$$x_{\text{inh}}(t) = \frac{a}{2\omega_0}t\sin(\omega_0 t). \quad (5.5.36)$$

Dabei haben wir $C_3 = 0$ gesetzt, weil wir ja nur eine spezielle Lösung der inhomogenen Gleichung benötigen. Die allgemeine Lösung von (5.5.31) ergibt sich also aus der Überlagerung der allgemeinen Lösung der entsprechenden homogenen Gleichung, deren Lösung wir oben in (5.3.19) gefunden haben, und der speziellen Lösung (5.5.36):

$$x(t) = \frac{a}{2\omega_0}t\sin(\omega_0 t) + C_1\cos(\omega_0 t) + C_2\sin(\omega_0 t). \quad (5.5.37)$$

Die Integrationskonstanten C_1 und C_2 bestimmen sich wieder aus den Anfangsbedingungen. Unabhängig von diesen wächst die Amplitude beständig an. Man spricht daher auch von der **Resonanzkatastrophe** (vgl. Abb. 5.6).

Natürlich kommt so etwas in der Natur nicht wirklich vor, da i.a. die Kräfte für große Auslenkungen nicht mehr harmonisch sind und daher die Beschreibung als harmonischer Oszillator für große Zeiten ungültig wird.

5.5.6 Allgemeine äußere Kräfte und die δ -Distribution

Zum Abschluss des Kapitels über den harmonischen Oszillator wollen wir noch die Frage beantworten, wie man den Fall beliebig vorgegebener äußerer Kräfte behandeln kann. Wir benötigen dazu nur eine spezielle Lösung der inhomogenen Gleichung

$$\ddot{x} + 2\gamma\dot{x} + \omega_0^2 x = a(t), \quad (5.5.38)$$

wobei a eine vorgegebene Funktion ist (die äußere Kraft ist $F(t) = ma(t)$). Die Idee zu einer Lösung ist, dass sich die Auslenkung $x(t)$ zur Zeit t als eine Art „kontinuierlicher Linearkombination“ der rechten Seite der Gleichung ergibt. Das physikalische Bild hinter dieser Idee ist, dass zu jedem Zeitpunkt die äußere Kraft den harmonischen Oszillator neu anstößt. Diese Überlegung führt uns auf den Ansatz

$$x(t) = \int_{-\infty}^{\infty} dt' G(t, t') a(t'). \quad (5.5.39)$$

Das **Kausalitätsprinzip** verlangt nun, dass die Auslenkung zur Zeit t nur von der äußeren Kraft zu früheren Zeitpunkten $t' < t$ abhängen darf. MaW. die äußere Kraft zu einer späteren Zeit kann nicht in die Vergangenheit zurückwirken. Dies verlangt, dass

$$G(t, t') = 0 \quad \text{für} \quad t' > t. \quad (5.5.40)$$

Die untere Integrationsgrenze, also den Anfangszeitpunkt der Bewegung, haben wir bei $t_0 \rightarrow -\infty$ gewählt, damit wir genau die Lösung erhalten, für die die von spezifischen Anfangsbedingungen abhängigen Einschwingvorgänge zu jedem endlichen Zeitpunkt bereits abgeklungen sind.

Wir setzen nun diesen Ansatz in (5.5.38) ein. Das führt auf die Gleichung

$$\int_{-\infty}^{\infty} dt' [\partial_t^2 G(t, t') + 2\gamma \partial_t G(t, t') + \omega_0^2 G(t, t')] a(t') = a(t). \quad (5.5.41)$$

Dies verlangt nun, dass der Ausdruck in den eckigen Klammern eine „Funktion“ $\delta(t, t')$ ist, die für beliebige stetige Funktionen a stets¹

$$\int_{-\infty}^{\infty} dt' \delta(t, t') a(t') = a(t) \quad (5.5.42)$$

gilt. Man kann zeigen, dass es eine solche Funktion im strengen Sinne nicht gibt. Wir können aber Näherungen für eine solche Funktion definieren. Die einfachste Näherung ist,

$$\delta_{\epsilon}(t, t') = \begin{cases} \frac{1}{\epsilon} & \text{falls } t' \in (t - \epsilon/2, t + \epsilon/2), \\ 0 & \text{sonst,} \end{cases} \quad (5.5.43)$$

wobei $\epsilon > 0$ beliebig ist. Mit dem Mittelwertsatz der Integralrechnung folgt dann

$$\int_{-\infty}^{\infty} dt' \delta_{\epsilon}(t, t') a(t') = \frac{1}{\epsilon} \int_{t-\epsilon/2}^{t+\epsilon/2} dt' a(t') = \frac{1}{\epsilon} a(\tau) \epsilon = a(\tau), \quad \tau \in (t - \epsilon/2, t + \epsilon/2). \quad (5.5.44)$$

Lassen wir nun $\epsilon \rightarrow 0^+$ gehen, erhalten wir wegen der Stetigkeit von a im Limes für das Integral tatsächlich $a(t)$; die Funktionen δ_{ϵ} streben allerdings nicht gegen eine gewöhnliche Funktion. Anschaulich ergibt sich eine Funktion, die überall verschwindet außer an der Stelle $t' = t$, wo sie unendlich wird. Der Grenzwert $\epsilon \rightarrow 0^+$ ergibt nur im Sinne des Integrals (5.5.42) einen Sinn, d.h. man führt *zuerst* die Integration in (5.5.42) mit der regularisierten Funktion δ_{ϵ} aus und lässt *dann* $\epsilon \rightarrow 0^+$ streben.

¹Im mathematisch strikten Sinne muss man an die möglichen Funktionen a noch weitere Forderungen stellen, insbesondere, dass sie im Unendlichen hinreichend schnell verschwinden, damit die uneigentlichen Integrale in unserer Betrachtung existieren. Wir gehen darauf hier nicht genauer ein.

Dies nennt man einen **Grenzwert im schwachen Sinne** und den Grenzwert $\delta(t, t')$ eine **verallgemeinerte Funktion oder Distribution**^a.

^aDiese Idee wurde u.a. von P. A. M. Dirac (1902-1984) im Zusammenhang mit seinen Pionierarbeiten zur Quantenmechanik in die theoretische Physik eingeführt. Die Mathematiker kritisierten anfangs diese Idee heftig, weil sie mathematisch nicht streng begründet war. Die Nützlichkeit der Idee in der praktischen Anwendung führte schließlich die Mathematiker dazu, diese Idee zu formalisieren. Daraus ist ein ganzer Zweig der modernen Mathematik entstanden, die **Funktionalanalysis**.

Offensichtlich können wir die Schreibweise für die δ -Distribution noch etwas vereinfachen. Dazu definieren wir zunächst

$$\delta_\epsilon(t) = \begin{cases} \frac{1}{\epsilon} & \text{für } t \in (-\epsilon/2, \epsilon/2) \\ 0 & \text{sonst.} \end{cases} \quad (5.5.45)$$

Dann gilt offensichtlich (*nachprüfen!*)

$$\delta_\epsilon(t, t') = \delta_\epsilon(t - t'), \quad (5.5.46)$$

und entsprechend schreiben wir auch

$$\delta(t, t') = \delta(t - t'). \quad (5.5.47)$$

Dann verlangt offenbar (5.5.41), dass

$$\partial_t^2 G(t, t') + 2\gamma \partial_t G(t, t') + \omega_0^2 G(t, t') = \delta(t - t'). \quad (5.5.48)$$

Daraus folgt, dass man auch mit dem Ansatz $G(t, t') = G(t - t')$ auskommen sollte und dass dann

$$\ddot{G}(t) + 2\gamma \dot{G}(t) + \omega_0^2 G(t) = \delta(t) \quad (5.5.49)$$

gelten sollte.

Wir suchen nun also eine Funktion, die die folgenden Eigenschaften besitzt:

- $G(t) = 0$ für $t < 0$ (**Kausalitäts- oder Retardierungsbedingung**)
- G erfüllt für $t > 0$ die homogene Differentialgleichung

Betrachten wir der Einfachheit halber den Schwingfall. Dann lautet die allgemeine Lösung für $t > 0$ gemäß (5.4.9)

$$G(t) = \begin{cases} \exp(-\gamma t)[A \cos(\omega t) + B(\sin \omega t)] & \text{für } t > 0, \\ 0 & \text{für } t \leq 0. \end{cases} \quad (5.5.50)$$

Wir müssen nun die beiden Integrationskonstanten A und B bestimmen, so dass die Singularität der δ -Funktion bei $t = 0$ in (5.5.49) auftritt. Wir benötigen offenbar zwei Bedingungen, um beide Konstanten eindeutig bestimmen zu können. Dazu überlegen wir uns, dass Ableitungen von Funktionen i.a. mehr Singularitäten besitzen als die abgeleitete Funktion. Eine stetige Funktion, die an der Stelle $t = 0$ einen Knick besitzt, ist dort sicher nicht differenzierbar, und wenn sie überall sonst differenzierbar ist, macht die Ableitungsfunktion bei $t = 0$ einen Sprung, weil die Steigung der Tangenten an der Knickstelle sich abrupt ändert. Wir erwarten also, dass die Singularität in (5.5.49) bei $t = 0$ in der höchsten vorkommenden Ableitung, also $\ddot{G}$ auftreten sollte, während G dort lediglich einen endlichen Sprung und G sogar stetig sein sollte (aber einen Knick bei $t = 0$ aufweisen sollte).

Als erste Bedingung zur Bestimmung der Integrationskonstanten in (5.5.50) verlangen wir also, dass G stetig bei $t = 0$ sein soll. Das führt auf

$$A = 0 \Rightarrow G(t) = B \exp(-\gamma t) \sin(\omega t) \quad \text{für } t > 0. \quad (5.5.51)$$

5.5. Der getriebene gedämpfte Oszillator

Schließlich benötigen wir noch eine zweite Bedingung. Dazu integrieren wir (5.5.49) über ein kleines Intervall $(-\epsilon, \epsilon)$. Für $t < 0$ ist $G(t) = \dot{G}(t) = \ddot{G}(t) = 0$. Daraus ergibt sich beim Integrieren

$$\dot{G}(\epsilon) + 2\gamma G(\epsilon) + \omega_0^2 \int_{-\epsilon}^{\epsilon} dt G(t) = 1. \quad (5.5.52)$$

Lassen wir nun $\epsilon \rightarrow 0^+$ streben, verschwindet im Limes sowohl $G(\epsilon)$ als auch das Integral wegen der angenommenen Stetigkeit von G bei $t = 0$. Es ergibt sich zur Bestimmung der verbliebenen Integrationskonstante also die Bedingung

$$\dot{G}(0^+) = 1. \quad (5.5.53)$$

Nun folgt aus (5.5.52) für $\epsilon \rightarrow 0^+$ für $t > 0$

$$\dot{G}(t) = B \exp(-\gamma t) [-\gamma \sin \omega t + \omega \cos \omega t] \rightarrow B\omega \quad \text{für } t \rightarrow 0^+. \quad (5.5.54)$$

Aus (5.5.53) folgt damit $B = 1/\omega$, und

$$G(t) = \begin{cases} \frac{\exp(-\gamma t)}{\omega} \sin(\omega t) & \text{für } t > 0, \\ 0 & \text{für } t \leq 0. \end{cases} \quad (5.5.55)$$

Diese Funktion nennt man die **retardierte Greensche Funktion** für den linearen Differentialoperator

$$\hat{L} = \frac{d^2}{dt^2} + 2\gamma \frac{d}{dt} + \omega_0^2. \quad (5.5.56)$$

Die ursprüngliche Differentialgleichung (5.5.38) können wir dann in der Form

$$\hat{L}x(t) = a(t) \quad (5.5.57)$$

und die allgemeine Lösung gemäß (5.5.39)

$$x(t) = \exp(-\gamma t) [C_1 \cos(\omega t) + C_2 \sin(\omega t)] + \int_{-\infty}^t dt' G(t-t') a(t') \quad (5.5.58)$$

schreiben. Dabei ist der erste Term die allgemeine Lösung der homogenen Differentialgleichung $\hat{L}x(t) = 0$ und das Integral die spezielle Lösung der inhomogenen Gleichung.

Wir können nun zeigen, dass dies tatsächlich eine Lösung ist. Dazu bemerken wir, dass das Integral effektiv nur bis t läuft, da für $t' > t$ die Green-Funktion konstruktionsgemäß wegen der Retardierungsbedingung verschwindet. Wir haben also für die Lösung der inhomogenen Gleichung

$$\begin{aligned} x_{\text{inh}}(t) &= \int_{-\infty}^t dt' G(t-t') a(t') \\ \dot{x}_{\text{inh}}(t) &= G(0^+) a(t) + \int_{-\infty}^t dt' \partial_t G(t-t') a(t') = \int_{-\infty}^t dt' \partial_t G(t-t') a(t'), \\ \ddot{x}_{\text{inh}}(t) &= \partial_t G(0^+) a(t) + \int_{-\infty}^t dt' \partial_t^2 G(t-t') a(t'). \end{aligned} \quad (5.5.59)$$

Konstruktionsgemäß ist nun $\partial_t G(0^+) = 1$, und damit finden wir schließlich

$$\hat{L}x(t) = a(t) + \int_{-\infty}^t dt' [\hat{L}G(t-t')] a(t') = a(t), \quad (5.5.60)$$

denn für $t' < t$ erfüllt $G(t - t')$ als Funktion von t die homogene Differentialgleichung $\hat{L}G(t - t') = 0$. Dies zeigt, dass unsere mathematisch recht unscharf begründeten Manipulationen mit der δ -Funktion tatsächlich zum Ziel führen, die inhomogene Differentialgleichung für beliebige rechte Seiten zu lösen.

Dieses Konzept der Green-Funktionen spielt eine weit größere Rolle bei der Lösung partieller Differentialgleichungen wie sie in der Feldtheorie (z.B. der Elektrodynamik im 3. Semester) auftreten als in der klassischen Mechanik.

5.6 Differentialgleichungen der Fuchsschen Klasse

In der theoretischen Physik sind auch oft Differentialgleichungen der sog. **Fuchsschen Klasse** (Lazarus Fuchs 1833-1902) wichtig. Sie tauchen z.B. bei bestimmten Eigenwertproblemen der Quantentheorie (Energieeigenzustände für den harmonischen Oszillator, des Wasserstoffatoms, Kugelflächenfunktionen usw.) auf. Solche Differentialgleichungen lassen sich mit einem (verallgemeinerten) Potenzreihenansatz lösen. Diese **Frobenius-Methode** (Ferdinand Georg Frobenius, 1849-1917) wollen wir hier besprechen.

Diese Differentialgleichungen sind homogene lineare Differentialgleichungen 2. Ordnung von der Form

$$z^2 u''(z) + z p(z) u'(z) + q(z) u(z) = 0. \quad (5.6.1)$$

Dabei verwenden wir als unabhängige Integrationsvariable jetzt z , um anzudeuten, dass wir allgemein Differentialgleichungen mit komplexen Funktionen betrachten. Natürlich gelten alle Überlegungen auch für reelle Differentialgleichungen. Wie wir wissen, ergibt sich die allgemeine Lösung einer solchen DGL durch Superposition zweier beliebiger voneinander linear unabhängiger Lösungen. Im Folgenden zeigen wir, wie wir immer zwei solcher Lösungen mittels verallgemeinerter Potenzreihenansätze finden können.

Die Funktionen p und q sollen dabei die Eigenschaft besitzen, dass sie Potenzreihenentwicklungen um 0 mit endlichem Konvergenzradius besitzen. Dies impliziert, dass diese Funktionen in einer Umgebung der 0 in der komplexen Zahlenebene beliebig oft differenzierbar (analytisch) sind, d.h. es gilt

$$p(z) = \sum_{k=0}^{\infty} p_k z^k, \quad q(z) = \sum_{k=0}^{\infty} q_k z^k. \quad (5.6.2)$$

Um eine Lösung zu finden, liegt es nahe, dass auch die Lösung von (5.6.1) die Form einer verallgemeinerten Potenzreihe besitzt, d.h.

$$u(z) = z^\lambda \sum_{k=0}^{\infty} u_k z^k. \quad (5.6.3)$$

Dabei dürfen wir ohne Beschränkung der Allgemeinheit annehmen, dass der Koeffizient $u_0 \neq 0$ ist, weil man ansonsten durch Umnummerieren des Summationsindex k und entsprechende Verschiebung von λ immer diese Annahme erfüllen kann (*nachrechnen!*).

Als erstes benötigen wir die Ableitungen dieses Ansatzes, um diese in (5.6.1) einsetzen zu können (*Nachrechnen!*):

$$\begin{aligned} u'(z) &= z^{\lambda-1} \sum_{k=0}^{\infty} (\lambda + k) u_k z^k, \\ u''(z) &= z^{\lambda-2} \sum_{k=0}^{\infty} (\lambda + k)(\lambda + k - 1) u_k z^k. \end{aligned} \quad (5.6.4)$$

Setzen wir dies nun in (5.6.1) ein, folgt nach Division durch z^λ

$$\sum_{k=0}^{\infty} (\lambda + k)(\lambda + k - 1) u_k z^k + p(z) \sum_{j=0}^{\infty} (\lambda + j) u_j z^j + q(z) \sum_{j=0}^{\infty} u_j z^j = 0. \quad (5.6.5)$$

Voraussetzungsgemäß kann man nun die Funktionen p und q um $z = 0$ in Potenzreihen (5.6.2) entwickeln. Um (5.6.5) verwenden zu können, um die gesuchten Koeffizienten u_k zu finden, müssen wir die linke Seite der Gl. (5.6.5) in eine Potenzreihe in Potenzen von z entwickeln.

Dazu betrachten wir zwei Potenzreihen

$$a(z) = \sum_{i=0}^{\infty} a_i z^i, \quad b(z) = \sum_{i=0}^{\infty} b_i z^i. \quad (5.6.6)$$

Diese mögen beide von 0 verschiedene Konvergenzradien r_a und r_b besitzen. Dann konvergieren beide Reihen nach bekannten Sätzen über Potenzreihen in jedem abgeschlossenen Intervall $|z| \leq r$ mit $r < \min(r_a, r_b)$ absolut. Demnach dürfen wir für diese z die Reihen im Produkt $a(z)b(z)$ beliebig umordnen, d.h. wir können einfach

$$a(z)b(z) = \sum_{i=0}^{\infty} \sum_{j=0}^{\infty} a_i b_j z^{i+j} \quad (5.6.7)$$

schreiben. Dies ist offenbar eine Potenzreihe in z . Wir müssen nur die Reihenglieder abermals umordnen, so dass sie nach Potenzen von z sortiert ist. Dazu muss man offenbar nur den neuen Summationsindex $k = i + j$ einführen. Da $i, j \in \mathbb{N}_0$ ist auch $k \in \mathbb{N}_0$. Wir erhalten demnach die gewünschte Form des Produkts durch

$$a(z)b(z) = \sum_{k=0}^{\infty} z^k \sum_{j=0}^k a_{j-k} b_j. \quad (5.6.8)$$

Es ist also

$$c(z) = a(z)b(z) = \sum_{k=0}^{\infty} c_k z^k \quad \text{mit} \quad c_k = \sum_{j=0}^k a_{j-k} b_j. \quad (5.6.9)$$

Da $c(z)$ für $|z| < \min(r_a, r_b)$ analytisch ist, ist der Konvergenzradius der Reihe (5.6.9) $r_c = \min(r_a, r_b)$. Man nennt die Formel (5.6.9) für das Produkt zweier Potenzreihen den **Cauchy-Produktsatz für Potenzreihen**.

Wir können diesen Satz nun für die beiden Produkte in (5.6.5) anwenden. Nach einigen Rechnungen (*Übung!*) folgt

$$\sum_{k=0}^{\infty} C_k z^k = \sum_{k=0}^{\infty} z^k \left\{ [(\lambda + k)(\lambda + k - 1) + (\lambda + k)p_0 + q_0]u_k + \sum_{j=0}^{k-1} [(\lambda + j)p_{k-j} + q_{k-j}]u_j \right\} = 0. \quad (5.6.10)$$

Eine Potenzreihe ist nun nur 0, wenn die durch die geschweifte Klammer definierten Koeffizienten C_k alle verschwinden.

Damit können wir wenigstens eine von zwei linear unabhängigen Lösungen der Differentialgleichung in Form einer verallgemeinerten Potenzreihe anzugeben. Dazu betrachten wir zunächst die Bedingung, dass $C_0 = 0$ sein muss. Es folgt

$$C(\lambda) = \lambda(\lambda - 1) + p_0 \lambda + q_0 = 0, \quad (5.6.11)$$

D.h. die Potenz des Vorfaktors z^λ muss eine Nullstelle der quadratischen Funktion $C(\lambda)$ sein. Im Folgenden seien λ_1 und λ_2 diese Nullstellen. Dabei wählen wir die Bezeichnungen so, dass $\operatorname{Re} \lambda_1 \geq \operatorname{Re} \lambda_2$ ist.

Nun betrachten wir weiter die Bedingung $C_k = 0$ für $k \in \mathbb{N}$. Wir können sie offenbar in der Form

$$C(\lambda + k)u_k + \sum_{j=0}^{k-1} [(\lambda + j)p_{k-j} + q_{k-j}]u_j = 0 \quad (5.6.12)$$

schreiben. Da voraussetzungsgemäß $\operatorname{Re} \lambda_1 \geq \operatorname{Re} \lambda_2$ ist, ist $C(\lambda_1 + k) \neq 0$ für alle $k \in \mathbb{N}$, und daraus ergeben sich nach der willkürlichen Wahl $u_0 = 1$ eindeutig alle u_j mit $j \in \mathbb{N}$, d.h. man erhält eine von zwei linear

unabhängigen Lösungen durch die Rekursion

$$u_1(z) = z^{\lambda_1} \sum_{k=0}^{\infty} u_k^{(1)} z^k, \quad u_0^{(1)} = 1, \quad u_k^{(1)} = -\frac{1}{C(\lambda_1 + k)} \sum_{j=0}^{k-1} [(\lambda_1 + j)p_{k-j} + q_{k-j}] u_j^{(1)}. \quad (5.6.13)$$

Man kann also (zumindest im Prinzip) $u_k^{(1)}$ für alle $k \in \mathbb{N}$ beginnend mit $k = 1$ iterativ lösen, womit eine Lösung der Differentialgleichung gefunden ist.

Wir benötigen jedoch für die allgemeine Lösung noch eine zweite von $u_1(z)$ linear unabhängige Lösung. Hier müssen wir nun eine Fallunterscheidung vornehmen.

Fall 1

Offenbar ergibt sich eine solche zweite Lösung genau wie eben für die Nullstelle λ_2 , wenn $\lambda_1 - \lambda_2 \notin \mathbb{N}_0$ ist, denn dann ist auch $C(\lambda_2 + k) \neq 0$ für alle $k \in \mathbb{N}$, und wir finden eine zweite Lösung ganz analog wie eben für die Nullstelle λ_1 :

$$u_2(z) = z^{\lambda_2} \sum_{k=0}^{\infty} u_k^{(2)} z^k, \quad u_0^{(2)} = 1, \quad u_k^{(2)} = -\frac{1}{C(\lambda_2 + k)} \sum_{j=0}^{k-1} [(\lambda_2 + j)p_{k-j} + q_{k-j}] u_j^{(2)}. \quad (5.6.14)$$

Wegen $\lambda_1 \neq \lambda_2$ sind u_1 und u_2 offenbar linear unabhängig, und die allgemeine Lösung der Differentialgleichung ist durch die Linearkombinationen der Funktionen $u_1(z)$ und $u_2(z)$ gegeben.

Fälle 2 und 3

In dem Fall, dass $\lambda_1 - \lambda_2 = m \in \mathbb{N}_0$ ist, lässt sich i.a. die zweite Lösung nicht durch einen einfachen verallgemeinerten Potenzreihenansatz finden. Um eine Idee für einen Ansatz für eine zweite linear unabhängige Lösung zu finden, betrachten wir das einfachste Beispiel für den Fall $m = 0$, nämlich die Differentialgleichung

$$z^2 u''(z) + z u'(z) = 0. \quad (5.6.15)$$

In dem Fall ist einfach $p(z) = 1 = p_0$ und $q(z) = 0 = q_0$. Das charakteristische Polynom ist also

$$C(\lambda) = \lambda(\lambda - 1) + \lambda = \lambda^2 \quad (5.6.16)$$

Man erhält also nur die eine Nullstelle $\lambda_1 = \lambda_2 = 0$, und es ist in diesem Beispiel $m = 0$. Offensichtlich ist die eine Lösung u_1 einfach $u_1 = 1$. Um die zweite Lösung zu finden, setzen wir $v(z) = u'(z)$. Dann erhalten wir die Differentialgleichung 1. Ordnung

$$z^2 v'(z) + z v(z) = 0 \Rightarrow \frac{v'(z)}{v(z)} = \frac{d}{dz} \ln[v(z)] = -\frac{1}{z} \Rightarrow v(z) = \frac{1}{z}. \quad (5.6.17)$$

Eine weitere Integration liefert

$$u_2(z) = \ln z. \quad (5.6.18)$$

Die allgemeine Lösung ist also in diesem Fall

$$u(z) = A + B \ln z \quad (5.6.19)$$

mit beliebigen Integrationskonstanten A und B .

Für unseren allgemeinen Fall mit $\lambda_1 - \lambda_2 \in \mathbb{N}_0$ legt dies den Ansatz

$$u_2(z) = v(z) + \gamma u_1(z) \ln z, \quad \gamma = \text{const.} \quad (5.6.20)$$

nahe. Wir wollen dies wieder in die Differentialgleichung (5.6.1) einsetzen. Die benötigten Ableitungen sind

$$u_2'(z) = v'(z) + \gamma u_1'(z) \ln z + \frac{\gamma}{z} u_1(z), \quad u_2''(z) = v''(z) + \frac{2\gamma}{z} u_1'(z) - \frac{\gamma}{z^2} u_1(z). \quad (5.6.21)$$

Da u_1 die Differentialgleichung löst, liefert das Einsetzen des Ansatzes (5.6.20) nach einiger Reibung (*Übung!*)

$$z^2 v''(z) + z p(z) v'(z) + q(z) v(z) = \gamma \{ [1 - p(z)] u_1(z) - 2z u_1'(z) \}. \quad (5.6.22)$$

Da voraussetzungsgemäß $p(z)$ um $z = 0$ in eine Potenzreihe entwickelt werden kann und $u_1(z)$ durch die verallgemeinerte Potenzreihe (5.6.13) gegeben ist folgt aus dem obigen Cauchy-Produktsatz für Potenzreihen, dass die rechte Seite ebenfalls durch eine verallgemeinerte Potenzreihe mit dem Vorfaktor z^{λ_1} gegeben sein muss, d.h.

$$z^2 v''(z) + z p(z) v'(z) + q(z) v(z) = \gamma z^{\lambda_1} \sum_{k=0}^{\infty} a_k z^k. \quad (5.6.23)$$

Demnach sollte auch v durch eine solche verallgemeinerte Potenzreihe gegeben sein. Dies legt es nahe, es mit dem Ansatz

$$v(z) = z^{\lambda_2} \sum_{k=0}^{\infty} v_k z^k \quad (5.6.24)$$

mit der zweiten Nullstelle λ_2 des charakteristischen Polynoms (5.6.11) zu versuchen. Voraussetzungsgemäß gilt im jetzt betrachteten Fall $\lambda_1 - \lambda_2 = m \in \mathbb{N}_0$. Für die gesuchten Koeffizienten v_k und die Konstante γ finden wir nun eine Gleichung, indem wir den Ansatz (5.6.23) in (5.6.22) einsetzen. Dazu können wir die entsprechende Rechnung für u (5.6.10) verwenden, wobei zu beachten ist, dass (5.6.10) durch die Division der linken Seite der Differentialgleichung durch z^{λ_1} entstanden ist. Man muss also entsprechend für v auch die rechte Seite durch den Faktor z^{λ_2} dividieren. Man erhält dann wegen $\lambda_1 - \lambda_2 = m$ schließlich

$$\sum_{k=0}^{\infty} z^k \left\{ C(\lambda_2 + k) v_k + \sum_{j=0}^{k-1} [(\lambda_2 + j) p_{k-j} + q_{k-j}] v_j \right\} = \gamma z^m \sum_{k=0}^{\infty} a_k z^k. \quad (5.6.25)$$

Weiter erhält man aus (5.6.22) und (5.6.23) mit Hilfe von (5.6.13) und (5.6.4) sowie der Cauchy-Produktformel

$$a_k = \left\{ [1 - 2(\lambda_1 + k)] - \sum_{j=0}^k p_{k-j} \right\} u_j^{(1)}. \quad (5.6.26)$$

Insbesondere ergibt sich für $k = 0$

$$a_0 = 1 - p_0 + 2\lambda_1. \quad (5.6.27)$$

Aus der pq -Formel zur Lösung für die Nullstellen des charakteristische Polynoms (5.6.16) folgt $1 - p_0 = \lambda_1 + \lambda_2$ (*Nachrechnen!*) und damit

$$a_0 = \lambda_2 - \lambda_1 = -m. \quad (5.6.28)$$

Nun müssen wir noch die Fälle $m = 0$ und $m \in \mathbb{N}$ unterscheiden.

Fall 2 ($m = \lambda_1 - \lambda_2 = 0$)

Für $m = 0$ ergibt der Koeffizientenvergleich der beiden Seiten der Gl. (5.6.25)

$$C(\lambda_2 + k) v_k + \sum_{j=0}^{k-1} [(\lambda_2 + j) p_{k-j} + q_{k-j}] v_j = \gamma a_k \quad \text{für } m = 0. \quad (5.6.29)$$

Für $k = 0$ steht wegen $C(\lambda_2 + k) = 0$ auf der linken Seite der Gleichung 0, was mit der rechten Seite kompatibel ist, da ja jetzt gemäß (5.6.28) $a_0 = -m = 0$ gilt. Wir können also sowohl v_0 als auch γ beliebig wählen. Der Einfachheit halber setzen wir $v_0 = 0$ und $\gamma = 1$. Da weiter $C(\lambda_2 + k) = C(\lambda_1 + k) \neq 0$ für alle $k \in \mathbb{N}$ gilt, sind dann die Koeffizienten v_k für $k \in \mathbb{N}$ eindeutig über (5.6.29) zu

$$v_k = \frac{1}{C(\lambda_2 + k)} \left\{ a_k - \sum_{j=0}^{k-1} [(\lambda_2 + j) p_{k-j} + q_{k-j}] v_j \right\} \quad (5.6.30)$$

bestimmt.

Fall 3 ($m = \lambda_1 - \lambda_2 \in \mathbb{N}$)

Falls $m \in \mathbb{N}$ ist, ergibt der Koeffizientenvergleich von Gl. (5.6.25) die Gleichung

$$C(\lambda_2 + k)v_k + \sum_{j=0}^{k-1} [(\lambda_2 + j)p_{k-j} + q_{k-j}]v_j = \begin{cases} 0 & \text{für } k \in \{0, \dots, (m-1)\}, \\ \gamma a_{k-m} & \text{für } k \in \{m, (m+1), \dots\}. \end{cases} \quad (5.6.31)$$

Um die Koeffizienten v_k sowie γ zu bestimmen, bemerken wir, dass $C(\lambda_2) = C(\lambda_1) = C(m + \lambda_2) = 0$ und $D(\lambda_2 + k) \neq 0$ für $k \in \{1, 2, \dots, (m-1)\}$ und $k \in \{(m+1), (m+2), \dots\}$.

Für $k = 0$ ist dann (5.6.31) für beliebiges v_0 erfüllt. Wir können daher $v_0 = 1$ wählen. Dann folgt aus (5.6.31) für $k \in \{1, 2, \dots, (m-1)\}$

$$v_k = -\frac{1}{C(\lambda_2 + k)} \sum_{j=0}^{k-1} [(\lambda_2 + j)p_{k-j} + q_{k-j}]v_j. \quad (5.6.32)$$

Für $k = m$ folgt aus (5.6.31) eindeutig γ , und zwar unabhängig von v_m (das wir daher willkürlich $v_m = 0$ setzen dürfen) mit der obigen Wahl $v_0 = 1$

$$\gamma = -\frac{1}{m} \sum_{j=0}^{m-1} [(\lambda_2 + j)p_{m-j} + q_{m-j}]v_j. \quad (5.6.33)$$

Für $k \in \{(m+1), (m+2), \dots\}$ liefert dann (5.6.31) eindeutig die übrigen Koeffizienten durch Rekursion

$$v_k = \frac{1}{C(\lambda_2 + k)} \left\{ \gamma a_{k-m} - \sum_{j=0}^{k-1} [(\lambda_2 + j)p_{k-j} + q_{k-j}]v_j \right\}. \quad (5.6.34)$$

Man macht sich leicht klar, dass sich durch diese Lösungsstrategie der Ansatz (5.6.20) für u_2 eine von u_1 linear unabhängige Lösung ergibt, so dass die allgemeine Lösung der DGL (5.6.31) durch beliebige Superpositionen von u_1 und u_2 gegeben ist.

Dass u_2 auch für $m \in \mathbb{N}_0$ von u_1 linear unabhängig ist, macht man sich wie folgt klar. Für $|z| \ll 1$ ist $u_1(z) \simeq z^{\lambda_1}$. Für den Fall $\lambda_1 = \lambda_2$, also $m = 0$, ist dann $u_2(z) \simeq u_1(z) \ln z \simeq z^{\lambda_1} \ln z$ und damit u_2 linear unabhängig von u_1 . Für $\lambda_1 - \lambda_2 = m \in \mathbb{N}$ ist $u_2(z) \simeq z^{\lambda_2} = z^{\lambda_1 - m} \simeq z^{-m} u_1(z)$ und folglich auch hier u_2 linear unabhängig von u_1 .

Damit haben wir eine Lösungsstrategie für die Gleichung (5.6.1) durch verallgemeinerte Potenzreihenansätze gefunden.

Betrachten wir als Beispiel die Gleichung

$$z^2 u''(z) + (2 - z)z u'(z) - zu = 0. \quad (5.6.35)$$

Offenbar ist diese Gleichung von der Art (5.6.1), und hier sind die Koeffizientenfunktionen glücklicherweise

$$p(z) = 2 - z, \quad q(z) = -z. \quad (5.6.36)$$

Wir haben also in diesem Fall nicht die komplizierten Summen zu berechnen, wenn $p(z)$ und $q(z)$ echte Potenzreihenentwicklungen erfordern.

Unsere allgemeinen Betrachtungen zeigen, dass wir mit dem Ansatz (5.6.3) beginnen können. Das charakteristische Polynom in diesem Fall ist gemäß (5.6.11) wegen $p_0 = p(0) = 2$ und $q_0 = q(0) = 0$

$$C(\lambda) = \lambda(\lambda - 1) + 2\lambda = \lambda(\lambda + 1). \quad (5.6.37)$$

Die Nullstellen sind (in der oben verwendeten Konvention mit $\operatorname{Re} \lambda_1 \geq \operatorname{Re} \lambda_2$ durch $\lambda_1 = 0$ und $\lambda_2 = -1$ gegeben. Es liegt also der Fall vor, dass $\lambda_1 - \lambda_2 = m = 1 \in \mathbb{N}$ ist. Wir müssen also die etwas aufwändigere

Lösungsstrategie des Falles 3 anwenden. Die erste Lösung u_1 , die sich für $\lambda = \lambda_1 = 0$ ergibt, folgt in allen drei Fällen aus (5.6.13). Da $p_0 = 2$ und $p_1 = -1$ und $q_0 = 0$ und $q_1 = -1$ und alle übrigen p_j und $q_j = 0$ ist, trägt in der Summe in (5.6.13) nur der Term mit $j = k - 1$ bei:

$$u_0^{(1)} = 1, \quad u_k^{(1)} = -\frac{1}{k(k+1)}[(1-k)-1]u_{k-1}^{(1)} = \frac{1}{k+1}u_{k-1}^{(1)}. \quad (5.6.38)$$

Damit folgt durch Iteration

$$u_0^{(1)} = 1, \quad u_1^{(1)} = \frac{1}{2}, \quad u_2^{(1)} = \frac{1}{2 \cdot 3}, \dots \Rightarrow u_k^{(1)} = \frac{1}{(k+1)!}. \quad (5.6.39)$$

Dies liefert also die Reihe

$$u_1(z) = \sum_{k=0}^{\infty} \frac{1}{(k+1)!} z^k = \sum_{k=1}^{\infty} \frac{1}{k!} z^{k-1} = \frac{1}{z}(\exp z - 1). \quad (5.6.40)$$

Dabei haben wir im letzten Schritt die Exponentialreihe

$$\exp z = \sum_{k=0}^{\infty} \frac{1}{k!} z^k \quad (5.6.41)$$

verwendet (*Nachprüfen!*). Unsere erste Lösung ist also überall analytisch.

Wir wissen weiter, dass für die zweite Lösung der Ansatz (5.6.20) mit der verallgemeinerten Potenzreihe (5.6.24) für v mit $v_0 = 1$ und $v_1 = 0$ zum Ziel führt. Hier ist $m = 1$. Damit erhalten wir sofort aus (5.6.33) $\gamma = 0$. Für $k \geq 2$ ergibt (5.6.34) $v_k = 0$. Die zweite Lösung ist also einfach

$$u_2(z) = \frac{1}{z}. \quad (5.6.42)$$

Man bestätigt durch Einsetzen dieser Lösung in (5.6.35), dass $u_1(z)$ und $u_2(z)$ tatsächlich zwei linear unabhängige Lösungen dieser DGL sind.

5. Gewöhnliche Differentialgleichungen

Anhang A

Vollständige Induktion

Das Beweisprinzip der **vollständigen Induktion** bezieht sich auf Aussagen, die von natürlichen Zahlen abhängen.

A.1 Das Beweisprinzip

Das Beweisprinzip der vollständigen Induktion beruht darauf, dass die natürlichen Zahlen abzählbar sind, d.h. beginnend mit 1 (oder manchmal auch 0) kann man die Menge der natürlichen Zahlen $\mathbb{N} = \{1, 2, 3, \dots\}$ (bzw. $\mathbb{N}_0 = \{0, 1, 2, \dots\}$) dadurch definieren, dass es eben eine kleinste solche Zahl gibt, die wir 1 (bzw. 0) nennen und zu jeder Zahl $n \in \mathbb{N}$ (oder $n \in \mathbb{N}_0$) einen eindeutig bestimmten und von allen vorherigen natürlichen Zahlen verschiedenen Nachfolger gibt.

Diese Idee lässt sich formal durch die **Peano-Axiome** ausdrücken (Giuseppe Peano 1858-1932). Die Elemente der Menge der natürlichen Zahlen $\mathbb{N}_0$ besitzen folgende Eigenschaften:

1. 0 ist eine natürliche Zahl.
2. Zu jedem $n \in \mathbb{N}_0$ existiert ein Nachfolger $n' \in \mathbb{N}_0$.
3. 0 ist kein Nachfolger einer natürlichen Zahl.
4. Natürliche Zahlen mit gleichem Nachfolger sind gleich, d.h. ist $m, n \in \mathbb{N}_0$ und gilt $n' = m'$, so ist notwendig $n = m$.
5. Sei X eine beliebige Menge und sei $0 \in X$ und folgt aus $n \in X$ auch $n' \in X$, so ist $\mathbb{N}_0 \subseteq X$.

Im Folgenden benutzen wir wieder die übliche Notation, dass $n' = n + 1$ ist. Man kann übrigens die Rechenregeln für das Addieren und Multiplizieren natürlicher Zahlen mit Hilfe der Peano-Axiome genauer begründen. Dies ist aber eher Gegenstand formaler Mathematikvorlesungen.

Auf dem letzten Axiom beruht das Beweisprinzip der vollständigen Induktion. Sei dazu $A(n)$ eine beliebige Aussage, die von $n \in \mathbb{N}_0$ abhängt. Dann kann man diese Aussage beweisen, wenn man zeigt

1. $A(0)$ ist wahr (**Induktionsanfang**).
2. Aus der Annahme, dass $A(n)$ wahr sei, folgt, dass auch $A(n + 1)$ gilt (**Induktionsschritt**).

Man macht sich leicht klar, dass dies direkt aus den Peano-Axiomen folgt. Dazu müssen wir nur die Menge X als die Menge definieren, für die $A(n)$ mit $n \in \mathbb{N}_0$ wahr ist. Dann zeigt die Anwendung des 5. Peanoaxioms sofort, dass $\mathbb{N}_0 \subseteq X$ ist, wenn sowohl der Induktionsanfang als auch der Induktionsschritt erfüllt sind.

Es sei noch erwähnt, dass das Beweisprinzip der vollständigen Induktion einerseits sehr mächtig ist, denn es ermöglicht die Beweise von *Behauptungen*, die von natürlichen Zahlen abhängen, auf meist recht einfache Weise. Andererseits ist der Nachteil, dass es i.a., keine Möglichkeit bietet, die entsprechenden Formeln auch *herzuleiten*.

A.2 Beispiele

Wir wollen das Beweisprinzip der vollständigen Induktion anhand einiger einfacher aber nichttrivialer Beispiele demonstrieren.

A.2.1 Der „kleine Gauß“

Der Legende nach hat ein genervter Schullehrer seiner Klasse die vermeintlich zeitraubende Aufgabe gestellt, die ersten 50 natürlichen Zahlen zu addieren, also die Summe $1+2+3+\dots+50$ auszurechnen. In dieser Klasse befand sich auch der später berühmte Mathematiker Carl-Friedrich Gauß, und dieser soll die Aufgabe schon nach kurzer Zeit gelöst haben. Sein Trick war, einfach die Summen in umgekehrter Reihenfolge untereinander zu schreiben:

$$x = 1 + 2 + 3 + 4 + \dots + 50 \quad (\text{A.2.1})$$

$$x = 50 + 49 + 48 + 47 + \dots + 1. \quad (\text{A.2.2})$$

Addiert man nun erst die übereinanderstehenden Zahlen, erhält man immer dieselbe Summe 51, und das insgesamt 50-mal. Damit erhielt Gauß sofort das Ergebnis $x + x = 2x = 50 \cdot 51$ oder $x = 50 \cdot 51/2 = 25 \cdot 51 = 1275$.

Man macht sich klar, dass das für jede Summe natürlicher Zahlen bis n funktioniert, d.h.

$$1 + 2 + 3 + \dots + n = \sum_{i=1}^n i = \frac{1}{2}(n+1)n. \quad (\text{A.2.3})$$

Nachdem wir die Formel gefunden haben, lässt sie sich auch leicht mittels vollständiger Induktion beweisen. Hier haben wir es mit einer Aussage für die Menge $\mathbb{N} = \{1, 2, \dots\}$ zu tun. Auch hier funktioniert das Beweisprinzip der vollständigen Induktion, nur dass der Induktionsanfang bei $n = 1$ beginnt. Die Aussage für $n = 1$ ist trivial, denn es gilt $1 = (1+1) \cdot 1/2 = 1$. Nehmen wir nun an, die Aussage sei für $n \in \mathbb{N}$ wahr. Dann gilt

$$\begin{aligned} \sum_{i=1}^{n+1} i &= \left(\sum_{i=1}^n i \right) + (n+1) \\ &\stackrel{\text{IA}}{=} \frac{1}{2}(n+1)n + (n+1) \\ &= (n+1) \left(\frac{n}{2} + 1 \right) \\ &= (n+1) \frac{n+2}{2} \\ &= \frac{1}{2}(n+2)(n+1), \end{aligned} \quad (\text{A.2.4})$$

und das ist die Behauptung für $(n+1)$. Demnach ist die Formel aufgrund des Beweisprinzips der vollständigen Induktion für alle $n \in \mathbb{N}$ bewiesen. Dabei steht „IA“ über dem einen Gleichheitszeichen für „Induktionsannahme“, d.h. in diesem Schritt wurde die Annahme verwendet, dass die Formel für n wahr ist.

A.2.2 Höhere Summenformeln

Wir beweisen noch zwei Summenformeln mit dem Beweisprinzip der Vollständigen Induktion:

$$\sum_{i=1}^n i^2 = \frac{1}{6}n(n+1)(2n+1), \quad (\text{A.2.5})$$

$$\sum_{i=1}^n i^3 = \left(\frac{n(n+1)}{2} \right)^2. \quad (\text{A.2.6})$$

Beweisen wir als erstes (A.2.5):

Induktionsanfang: Für $n = 1$ ist die Behauptung $1 = 1(1+1)(2 \cdot 1 + 1)/6 = 1$ wahr.

Induktionsschritt: Nehmen wir an, die Behauptung stimmt für n , dann folgt

$$\begin{aligned} \sum_{i=1}^{n+1} i^2 &= \left(\sum_{i=1}^n i^2 \right) + (n+1)^2 \\ &\stackrel{\text{IA}}{=} \frac{1}{6}n(n+1)(2n+1) + (n+1)^2 \\ &= (n+1) \frac{n(2n+1) + 6(n+1)}{6} \\ &= (n+1) \frac{2n^2 + 7n + 6}{6} \\ &= \frac{1}{6}(n+1)(n+2)[2(n+1) + 1], \end{aligned} \quad (\text{A.2.7})$$

und das ist die Behauptung für $(n+1)$, und damit ist (A.2.5) nach dem Beweisprinzip der vollständigen Induktion für alle $n \in \mathbb{N}$ bewiesen.

Genauso zeigt man (A.2.6):

Induktionsanfang: In der Tat ist für $n = 1$ die Gleichung erfüllt, denn $1 = [1(1+1)/2]^2 = 1$.

Induktionsschritt: Angenommen, die Formel sei korrekt für n . Dann ist

$$\begin{aligned} \sum_{i=1}^{n+1} i^3 &= \left(\sum_{i=1}^n i^3 \right) + (n+1)^3 \\ &\stackrel{\text{IA}}{=} \frac{[n(n+1)]^2}{4} + (n+1)^3 \\ &= (n+1)^2 \left[\frac{n^2}{4} + (n+1) \right] \\ &= (n+1)^2 \frac{n^2 + 4(n+1)}{4} \\ &= (n+1)^2 \frac{(n+2)^2}{4} = \left[\frac{(n+1)(n+2)}{2} \right]^2, \end{aligned} \quad (\text{A.2.8})$$

und das ist die Behauptung für $(n+1)$, und damit ist (A.2.6) nach dem Beweisprinzip der vollständigen Induktion für alle $n \in \mathbb{N}$ bewiesen.

A.2.3 Geometrische Reihe

Es sei $q \in \mathbb{R} \setminus \{1\}$ beliebig. Dann gilt für $n \in \mathbb{N}_0$

$$\sum_{i=0}^n q^i = 1 + q + q^2 + \cdots + q^n = \frac{q^{n+1} - 1}{q - 1}. \quad (\text{A.2.9})$$

A. Vollständige Induktion

Wir beweisen diese Behauptung wieder mit dem Beweisprinzip der vollständigen Induktion.

Induktionsanfang: Für $n = 0$ ist $q^0 = (q^{0+1} - 1)/(q - 1) = 1$ und damit die Formel richtig.

Induktionsschritt: Angenommen die Formel gilt für n . Dann ist

$$\begin{aligned} \sum_{i=0}^{n+1} q^i &= \left(\sum_{i=0}^n q^i \right) + q^{n+1} \\ &\stackrel{\text{IA}}{=} \frac{q^{n+1} - 1}{q - 1} + q^{n+1} \\ &= \frac{(q^{n+1} - 1) + (q - 1)q^{n+1}}{q - 1} \\ &= \frac{q^{n+1} - 1 + q^{n+2} - q^{n+1}}{q - 1} \\ &= \frac{q^{n+2} - 1}{q - 1}, \end{aligned} \tag{A.2.10}$$

und das ist die Behauptung für $(n + 1)$. Damit ist (A.2.9) nach dem Beweisprinzip der vollständigen Induktion für alle $n \in \mathbb{N}_0$ bewiesen.

A.2. Beispiele

A. Vollständige Induktion

Literaturverzeichnis

- [AHK⁺18] T. Arens, et al., *Mathematik*, 4. Aufl., Springer-Verlag, Berlin, Heidelberg (2018).
URL <https://doi.org/10.1007/978-3-662-56741-8>
- [BFK⁺18a] M. Bartelmann, et al., *Theoretische Physik 1 – Mechanik*, Springer-Verlag, Berlin, Heidelberg (2018).
URL <https://doi.org/10.1007/978-3-662-56115-7>
- [BFK⁺18b] M. Bartelmann, et al., *Theoretische Physik 2 – Elektrodynamik*, Springer-Verlag, Berlin, Heidelberg (2018).
URL <https://doi.org/10.1007/978-3-662-56117-1>
- [BFK⁺18c] M. Bartelmann, et al., *Theoretische Physik 3 – Quantenmechanik*, Springer-Verlag, Berlin, Heidelberg (2018).
URL <https://doi.org/10.1007/978-3-662-56072-3>
- [BFK⁺18d] M. Bartelmann, et al., *Theoretische Physik 4 – Thermodynamik und Statistische Physik*, Springer-Verlag, Berlin, Heidelberg (2018).
URL <https://doi.org/10.1007/978-3-662-56113-3>
- [BK88] D. Bourne, P. Kendall, *Vektoranalysis*, 2. Aufl., B. G. Teubner, Stuttgart (1988).
- [Bro03] R. Bronson, *Differential Equations*, Schaum's Easy Outlines, McGraw-Hill, New York, Chicago, San Francisco (2003).
- [Col90] L. Collatz, *Differentialgleichungen*, Teubner (1990).
- [For13] O. Forster, *Analysis 1*, 11. Aufl., Springer Spektrum, Wiesbaden (2013).
- [Gro12] S. Großmann, *Mathematischer Einführungskurs für die Physik*, 10. Aufl., Springer Verlag, Berlin, Heidelberg (2012).
URL <https://doi.org/10.1007/978-3-8348-8347-6>
- [Hef18] K. Hefft, *Mathematischer Vorkurs zum Studium der Physik*, 2. Aufl., Springer Spektrum, Berlin (2018).
URL <https://doi.org/10.1007/978-3-662-53831-9>
- [HH19a] K. Hölbig, J. Hörner, *Aufgaben und Lösungen zur Höheren Mathematik 1*, 2. Aufl., Springer Spektrum, Berlin (2019).
URL <https://doi.org/10.1007/978-3-662-58445-3>
- [HH19b] K. Hölbig, J. Hörner, *Aufgaben und Lösungen zur Höheren Mathematik 2*, 2. Aufl., Springer Spektrum, Berlin (2019).
URL <https://doi.org/10.1007/978-3-662-58667-9>

- [HH19c] K. Hölbig, J. Hörner, *Aufgaben und Lösungen zur Höheren Mathematik 3*, 2. Aufl., Springer Spektrum, Berlin (2019).
URL <https://doi.org/10.1007/978-3-662-58723-2>
- [Joo89] G. Joos, *Lehrbuch der theoretischen Physik*, 15. Aufl., Aulaverlag, Wiesbaden (1989).
- [Nol18] W. Nolting, *Grundkurs Theoretische Physik 1: Klassische Mechanik*, 11. Aufl., Springer Spektrum, Berlin, Heidelberg (2018).
URL <https://doi.org/10.1007/978-3-662-57584-0>
- [Rem92] R. Remmert, *Funktionentheorie 1*, 3. Aufl., Springer-Verlag, Berlin, Heidelberg, New York (1992).
URL <https://doi.org/10.1007/978-3-642-97397-0>
- [Sau73] F. Sauter, *Becker/Sauter Theorie der Elektrizität 1*, 21. Aufl., B. G. Teubner, Stuttgart (1973).
URL <https://doi.org/10.1007/978-3-322-96789-3>
- [Sch84] N. Schaumberger, The Derivatives of Sin x and Cos x, *The Two-Year College Mathematics Journal* **15**, 143 (1984).
URL <https://doi.org/10.2307/2686521>
- [SH99] M. R. Spiegel, C. Hipp, *Einführung in die höhere Mathematik*, Schaum's Outline - Überblicke/Aufgaben, McGraw-Hill Book Company (1999).
- [Som92] A. Sommerfeld, *Vorlesungen über Theoretische Physik II, Mechanik der deformierbaren Medien*, Verlag Harri Deutsch, Frankfurt/M. (1992).
- [Wei80] J. Weidmann, *Linear Operators in Hilbert Space*, Springer Verlag, New York, Berlin, Heidelberg (1980).