

Einführung in die Methoden der Theoretischen Physik

Thomas Filk

Skript zur Vorlesung

Wintersemester 2005/6 an der Universität Freiburg
Wintersemester 2011/12 an der Universität Freiburg

(Version vom 11. August 2011)

Inhaltsverzeichnis

1	Abbildungen, Ableitungen und Integrale	7
1.1	Abbildungen	7
1.2	Folgen und Grenzwerte	8
1.3	Ableitungen	10
1.3.1	Definitionen	10
1.3.2	Ableitungsregeln	11
1.3.3	Notation	11
1.3.4	Taylor-Entwicklung	13
1.4	Integrale	14
1.4.1	Definitionen	14
1.4.2	Integrationsregeln	15
2	Vektoren	17
2.1	Gruppen und Körper	17
2.2	Vektorräume	19
2.2.1	Beispiel 1: Verschiebungen in einer Ebene	20
2.2.2	Beispiel 2: n -Tupel von Zahlen	20
2.3	Dimension, Basis, lineare Unabhängigkeit	21
2.4	Der duale Vektorraum	22
2.4.1	Index-Schreibweise und Summenkonvention	23
2.5	Das Skalarprodukt	25
2.6	Tensoren	27
2.7	Der euklidische Vektorraum	28
2.7.1	Das kartesische Skalarprodukt	28
2.7.2	Das Kreuzprodukt	30
2.7.3	Der ϵ -Tensor	32
2.7.4	Das Spatprodukt	34
2.7.5	Determinanten	35

3	Lineare Abbildungen	39
3.1	Lineare Abbildungen	39
3.2	Basistransformationen	42
3.3	Spezielle lineare Abbildungen	45
3.3.1	Multiplikation von Matrizen	45
3.3.2	Die inverse Matrix	46
3.3.3	Die transponierte Matrix	46
3.3.4	Orthogonale Matrizen	47
3.3.5	Eigenwerte und Eigenvektoren	49
3.3.6	Selbstadjungierte Matrizen	50
3.4	Skalare - Pseudoskalare - Vektoren - Pseudovektoren	51
4	Komplexe Zahlen	53
4.1	Definition und einfache Relationen	53
4.2	Analytische Funktionen	56
4.3	Die komplexe Exponentialfunktion	57
4.4	Die Exponentialfunktion und 2-dimensionale Drehungen	59
5	Kurven und Flächen	63
5.1	Darstellung von Wegen im $\mathbf{R}^n$	63
5.2	Der Tangentialvektor - Die Geschwindigkeit	64
5.3	Die Beschleunigung	66
5.4	Zwischenabschnitt: der Gradient	68
5.5	Das mitgeführte Dreibein	72
5.6	Flächen	74
5.6.1	Parameterdarstellung von Flächen	74
5.6.2	Charakterisierung durch Bedingungen	76
5.7	Parameterdarstellung d -dimensionaler Räume im $\mathbf{R}^n$	78
5.8	Kegelschnitte	79
5.9	Koordinatensysteme	82
5.9.1	Die Polarkoordinaten im $\mathbf{R}^2$	82
5.9.2	Zylinderkoordinaten im $\mathbf{R}^3$	83
5.9.3	Kugelkoordinaten	83
5.9.4	Allgemeine Koordinatentransformationen	84
5.10	Das metrische Feld	85
5.10.1	Allgemeine Definition	86

5.10.2	Polar- und Kugelkoordinaten	87
5.10.3	Die Kugeloberfläche	88
5.10.4	Intrinsische Bedeutung der Metrik	89
6	Elemente der Vektoranalysis	93
6.1	Tensorfelder	93
6.1.1	Skalare Felder	94
6.1.2	Vektorfelder	95
6.2	Besondere Ableitungen von Feldern	96
6.2.1	Gradient, Divergenz und Rotation	96
6.2.2	Der Laplace-Operator	97
6.2.3	Beziehungen zwischen zweiten Ableitungen von Feldern . .	98
6.3	Kurvenintegrale	99
6.3.1	Die Bogenlänge	99
6.3.2	Kurvenintegration über skalare Felder	100
6.3.3	Berechnung der Arbeit	101
6.4	Flächenintegrale und Stokes'scher Satz	102
6.4.1	Flächenelement und Flächenintegral im $\mathbf{R}^3$	102
6.4.2	Beispiel: Die Kugeloberfläche	103
6.4.3	Der Stokes'sche Satz	104
6.5	Volumenintegrale und der Satz von Gauß	106
6.5.1	Volumenelement und Volumenintegration	106
6.5.2	Beispiel: Die Kugel vom Radius R	107
6.5.3	Der Satz von Gauß – Die Divergenz als Quellendichte . . .	108
6.6	Die Kontinuitätsgleichung	109
6.7	Jacobi-Determinanten und Metrik	110
7	Bewegungsgleichungen, Symmetrien und Erhaltungsgrößen	113
7.1	Die Newton'schen Gesetze	113
7.2	Die Newton'schen Bewegungsgleichungen	115
7.2.1	Anfangsbedingungen	117
7.2.2	Randbedingungen	118
7.2.3	Zustandsraum bzw. Phasenraum	119
7.2.4	Beispiel: Die freie Bewegungsgleichung	120
7.2.5	Beispiel: Der schiefe Wurf	120
7.3	Symmetrien und Erhaltungsgrößen	122

7.3.1	Erhaltungsgrößen	122
7.3.2	Die Energieerhaltung	122
7.3.3	Die Drehimpulserhaltung	123
7.3.4	Symmetrien	125
8	Oszillatoren	127
8.1	Der freie harmonische Oszillator	127
8.1.1	Die Bewegungsgleichung und die „geratenen“ Lösungen . .	127
8.1.2	Lösung durch den Exponentialansatz	130
8.1.3	Lösung durch Integration der Erhaltungsgröße	131
8.2	Der gedämpfte harmonische Oszillator	133
8.2.1	Bewegungsgleichung und Exponentialansatz	133
8.2.2	Die gedämpfte Schwingung: $4km > \mu^2$	134
8.2.3	Die reine Dämpfung: $4km < \mu^2$	135
8.2.4	Der aperiodische Grenzfall: $4km = \mu^2$	136

Kapitel 1

Abbildungen, Ableitungen und Integrale

1.1 Abbildungen

Definition: Eine *Abbildung* [*map*, *mapping*] f von einer Menge [*set*] U in eine Menge V ist eine Zuordnungsvorschrift, die jedem Element $u \in U$ genau ein Element $v \in V$ zuordnet. Man schreibt in diesem Fall:

$$f : U \longrightarrow V \quad u \mapsto v = f(u). \quad (1.1)$$

v bezeichnet man als das *Bild* [*image*] von u . Umgekehrt nennt man die Menge aller $u \in U$ mit der Eigenschaft $f(u) = v$ die *Urbildmenge* [*inverse image*] von v . Die Menge aller v , sodass es (mindestens) ein $u \in U$ gibt mit $f(u) = v$ bezeichnet man als *Bildmenge* [*range*] der Abbildung f . U heißt Urbildmenge [*domain*] von f .

A: Eine Funktion kann auch als eine spezielle Form von *Relation* [*relation*] definiert werden. Eine Relation R zwischen den Elementen einer Menge U und den Elementen einer Menge V ist dabei eine Teilmenge des kartesischen Produkts von U und V , also $R \subset U \times V$. Eine Funktion f entspricht dann einer Relation R_f mit der Eigenschaft, dass es zu jedem Element $u \in U$ genau ein Element (und nur ein Element) $v \in V$ geben muss, sodass $(u, v) \in R_f$.

Definition: Eine Abbildung $f : U \rightarrow V$ heißt *injektiv* [*injective*], wenn aus $f(u) = f(u')$ folgt $u = u'$. (Die Urbildmenge zu jedem $v \in V$ enthält also maximal ein Element.)

Definition: Eine Abbildung $f : U \rightarrow V$ heißt *surjektiv* [surjective], wenn es zu jedem $v \in V$ (mindestens) ein $u \in U$ gibt, sodass $f(u) = v$. (Die Bildmenge der Abbildung f ist also gleich der gesamten Menge V .)

Definition: Eine Abbildung $f : U \rightarrow V$ heißt *bijektiv* [bijective], wenn sie sowohl surjektiv als auch injektiv ist. (Man bezeichnet die Mengen U und V in diesem Fall als *gleich mächtig*. Handelt es sich bei U und V um Mengen mit endlich vielen Elementen, so ist die Anzahl der Elemente von U und V in diesem Fall gleich.)

Definition: Die *Identitätsabbildung* [identity mapping] $\mathbf{1}_U$ von einer Menge U in sich selbst ist die Abbildung, die jedem Element $u \in U$ wieder u zuordnet: $\mathbf{1}_U(u) = u$. Diese Abbildung ist immer bijektiv.

Unter gewissen Voraussetzungen lassen sich zwei Abbildungen f und g hintereinander ausführen. Sei f eine Abbildung von U nach V und g eine Abbildung von V nach W , dann bezeichnet man mit $g \circ f : U \rightarrow W$ die Abbildung, die dem Element $u \in U$ das Element $g(f(u)) \in W$ zu ordnet. Die allgemeine Voraussetzung an die Hintereinanderschaltbarkeit zweier Abbildungen f und g ist, dass die Bildmenge von f in der Urbildmenge von g enthalten sein muss.

Definition: Eine Abbildung $f^{-1} : V \rightarrow U$ heißt *invers* [inverse] (genauer *links invers*) zu einer Abbildung $f : U \rightarrow V$ (oder auch *Umkehrabbildung* [inverse map] zu f), wenn die Hintereinanderschaltung f^{-1} nach f die Identitätsabbildung $\mathbf{1}_U : U \rightarrow U$ ist: $f^{-1}(f(u)) = u$. Notwendige und hinreichende Voraussetzung [necessary and sufficient condition] für die Existenz einer Umkehrabbildung zu einer Funktion f ist, dass f injektiv ist.

A: Ist f bijektiv, so ist die Umkehrabbildung f^{-1} eindeutig. Im Allgemeinen kann eine injektive Abbildung f mehrere Umkehrabbildungen besitzen.

1.2 Folgen und Grenzwerte

Definition: Das *kartesische Produkt* [cartesian product] zweier Mengen U und V ist die Menge

$$U \times V := \{(u, v) | u \in U, v \in V\}.$$

Diese Menge besteht also aus allen (geordneten) Punktepaaaren (u, v) , wobei $u \in U$ und $v \in V$.

Definition: Eine *Folge* [sequence] $\{a_i\}_{i \in \mathbf{N}}$ ist eine Abbildung von den natürlichen Zahlen $\mathbf{N}$ in eine Menge M . Die natürlichen Zahlen bezeichnet man in diesem Fall als die *Indexmenge* [index set].

Definition: Eine *Metrik* [*metric*] d auf einer Menge M ist eine Abbildung von $M \times M$ (dem kartesischen Produkt der Menge M mit sich selbst) in die nicht-negativen reellen Zahlen, die folgenden Bedingungen genügt:

1. Für alle $x, y \in M$ gilt: Aus $d(x, y) = 0$ folgt $x = y$.
2. Für alle $x, y \in M$ gilt: $d(x, y) = d(y, x)$.
3. Für alle $x, y, z \in M$ gilt (Dreiecksungleichung [*triangle inequality*]):

$$d(x, z) \leq d(x, y) + d(y, z) .$$

Eine Menge M , auf der eine Metrik d definiert ist, bezeichnet man als *metrischen Raum* [*metric space*].

Beispiel: Auf dem $\mathbf{R}^n$ definiert die Abbildung

$$d(\vec{x}, \vec{y}) = \sqrt{(x^1 - y^1)^2 + (x^2 - y^2)^2 + \dots + (x^n - y^n)^2}$$

eine Metrik.

Definition: Sei M ein metrischer Raum. Man sagt, eine Folge $\{a_i\}_{i \in \mathbf{N}}$ mit $a_i \in M$ *konvergiert* [*converges*] gegen ein Element $a \in M$, wenn es zu jedem $\epsilon > 0$ ein $n_0 \in \mathbf{N}$ gibt, sodass

$$d(a_i, a) < \epsilon \quad \text{für alle } i > n_0 .$$

Das Element a bezeichnet man als den *Grenzwert* [*limit value*] der Folge.

Diese Definition der Konvergenz setzt voraus, dass der Grenzwert a der Folge $\{a_i\}$ selber ein Element der Menge M ist. Es gibt aber auch Fälle, bei denen eine Folge von Elementen $\{a_i\} \in M$ konvergiert, der Grenzwert jedoch nicht Element der Menge M ist. Die folgende Definition der Cauchy-Konvergenz umgeht dieses Problem.

Definition: Eine Folge $\{a_i\}_{i \in \mathbf{N}}$ von Elementen eines metrischen Raumes M heißt *Cauchy konvergent* [*Cauchy convergent*], wenn es zu jedem $\epsilon > 0$ ein $n_0 \in \mathbf{N}$ gibt, sodass

$$d(a_i, a_j) < \epsilon \quad \text{für alle } i, j > n_0 .$$

1.3 Ableitungen

1.3.1 Definitionen

Im Folgenden betrachten wir nur Abbildungen von Teilmengen [subsets] U der reellen Zahlen [real numbers] $\mathbf{R}$ in die reellen Zahlen. Die Konzepte der Stetigkeit [continuity] und der Ableitung [derivative] lassen sich zwar für allgemeinere Räume definieren, werden aber in dieser Form vorläufig nicht benötigt.

Definition: Eine Abbildung $f : U \subset \mathbf{R} \rightarrow \mathbf{R}$ heißt *stetig* [continuous] im Punkte $x \in U$, wenn für jede Folge $x_n \rightarrow x$ gilt: $f(x_n) \rightarrow f(x)$.

A: Sofern auf allgemeinen Mengen U und V eine Topologie (Definition offener Teilmengen) gegeben ist, lässt sich die Stetigkeit einer Funktion $f : U \rightarrow V$ auch so definieren, dass das Urbild jeder offenen Menge in V wieder eine offene Menge in U sein muss. In der üblichen Topologie auf den reellen Zahlen $\mathbf{R}$ sind offene Mengen als Vereinigung von offenen Intervallen (a, b) definiert. Die beiden Definitionen von Stetigkeit sind dann äquivalent.

Definition: Eine stetige Abbildung $f : U \subset \mathbf{R} \rightarrow \mathbf{R}$ heißt *ableitbar* [differentiable] im Punkte $x \in U$, wenn der Grenzwert [limit]

$$f'(x) := \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} \quad (1.2)$$

existiert. Ist f für alle $x \in U$ ableitbar, so bezeichnet man f einfach als *ableitbar*. Die Funktion $f' : U \subset \mathbf{R} \rightarrow \mathbf{R}$ nennt man die *Ableitung* [derivative] der Funktion f .

A: In der Mathematik wird die Ableitung $Df(x) = f'(x)$ einer Funktion f im Punkte x als „bester linearer Fit“ der Funktion f in x definiert. Diese Linearität bezieht sich nicht auf das Argument x , sondern $Df(x)$ wird als lineare Abbildung auf das Argument h aufgefasst, d.h., es gilt:

$$Df(x)(h) = f'(x)h = f(x+h) - f(x) + o(h).$$

$o(h)$ bedeutet dabei, dass dieser Term im Grenzwert $h \rightarrow 0$ schneller verschwindet als h , d.h., $\lim_{h \rightarrow 0} o(h)/h = 0$. Später werden wir auch die Notation $O(h)$ verwenden, die besagt, dass der Ausdruck $O(h)/h$ im Grenzfall $h \rightarrow 0$ beschränkt bleibt.

1.3.2 Ableitungsregeln

Linearität: [*linearity*] Für zwei (ableitbare) Funktionen f und g gilt: Die Ableitung der Summe der beiden Funktionen ist gleich der Summe der Ableitungen,

$$(f + g)'(x) = f'(x) + g'(x),$$

und die Ableitung des Vielfachen einer Funktion ist gleich dem Vielfachen der Ableitung:

$$(\lambda f)'(x) = \lambda f'(x) \quad \lambda \in \mathbf{R}.$$

Produktregel: [*product rule*] Für die Ableitung der Produktfunktion $f \cdot g$ zweier (ableitbaren) Funktionen f und g gilt:

$$(f \cdot g)'(x) = f(x) \cdot g'(x) + f'(x) \cdot g(x).$$

Kettenregel: [*chain rule*] Seien f und g zwei (ableitbare) Funktionen von Teilmengen der reellen Zahlen in die reellen Zahlen, sodass das Bild von f im Urbild von g liegt. Für die Ableitung der Hintereinanderschaltung von g nach f — $(g \circ f)(x) = g(f(x))$ — gilt:

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

$(g'(f(x)))$ ist die Ableitung g' der Funktion g , ausgewertet an der Stelle $f(x)$.
Aus der Tatsache, dass die Ableitung der Funktion $f(x) = 1/x$ die Funktion $f'(x) = -1/x^2$ ist, folgt - zusammen mit der Produkt- und der Kettenregel - die *Quotientenregel*: Sei

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

dann gilt:

$$\left(\frac{f}{g}\right)'(x) = \frac{f'(x)g(x) - f(x)g'(x)}{g(x)^2}.$$

1.3.3 Notation

Statt $f'(x)$ schreibt man auch manchmal

$$f' = \frac{df}{dx}$$

und

$$f'(x) = \frac{df(x)}{dx} \equiv \left. \frac{df}{dx} \right|_x.$$

Damit soll zum Ausdruck gebracht werden, dass die Ableitung selber als eine (lineare) Abbildung auf dem Raum der Funktionen angesehen werden kann, d.h., $\frac{d}{dx}$ ist eine Abbildung, die der Funktion f die Funktion f' zuordnet. Die Linearität dieser Abbildung folgt aus der Linearität der Ableitung:

$$\frac{d}{dx}(\alpha f + \beta g) = \alpha \frac{df}{dx} + \beta \frac{dg}{dx}.$$

Ist die Ableitungsfunktion f' selber wieder ableitbar, so kann man auch die zweite Ableitung von f bilden und schreibt dafür:

$$f''(x) = \frac{d^2 f(x)}{dx^2}.$$

Während man für die dritte Ableitung einer Funktion manchmal noch f''' schreibt, kennzeichnet man die höheren Ableitung meist durch den in Klammern gesetzten Hochindex n :

$$f^{(n)}(x) = \frac{d^n f(x)}{dx^n}.$$

Ist die physikalische Interpretation des Arguments der Funktion f die Zeit t , so schreibt man für die Ableitung oft:

$$\frac{df(t)}{dt} = \dot{f}(t).$$

Entsprechend ist die zweite Ableitung $\ddot{f}(t)$.

Besitzt die Funktion f mehrere Argumente, handelt es sich also um eine Abbildung vom $\mathbf{R}^n$ (oder einer offenen Umgebung in $\mathbf{R}^n$) in die reellen Zahlen, so kann man die Ableitung von f bezüglich jedes dieser Argumente bilden. Die anderen Argumente werden dabei konstant gelassen. Für diese *partiellen Ableitungen* [*partial derivatives*] schreibt man meist:

$$\begin{aligned} \frac{\partial f(x_1, x_2, \dots, x_n)}{\partial x_i} &= \partial_i f(x_1, x_2, \dots, x_n) \equiv \partial_{x_i} f(x_1, x_2, \dots, x_n) \\ &= \lim_{h \rightarrow 0} \frac{f(x_1, \dots, x_i + h, \dots, x_n) - f(x_1, \dots, x_i, \dots, x_n)}{h}. \end{aligned}$$

Die Notation

$$f'(x) = \frac{df(x)}{dx}$$

soll auch andeuten, dass die Ableitung der Grenzwert des *Differenzenquotienten* [*difference quotient*] ist:

$$\frac{df(x)}{dx} = \lim_{\Delta x \rightarrow 0} \frac{\Delta f(x)}{\Delta x} \quad \text{mit} \quad \Delta f(x) = f(x + \Delta x) - f(x).$$

Unter geeigneten Voraussetzungen lassen sich die Symbole df und dx wie gewöhnliche Differenzen Δf und Δx manipulieren.

1.3.4 Taylor-Entwicklung

Ist eine Funktion f in der Umgebung eines Punktes x beliebig oft ableitbar, so gilt in dieser Umgebung die *Taylor-Entwicklung* [*Taylor expansion*]:

$$\begin{aligned} f(x+h) &= f(x) + f'(x)h + \frac{1}{2}f''(x)h^2 + \dots = \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(x)h^n + R(h) \\ &= \sum_{n=0}^{\infty} \frac{1}{n!} \frac{d^n f(x)}{dx^n} h^n + R(h). \end{aligned} \quad (1.3)$$

R ist hierbei ein Term, der für $h \rightarrow 0$ schneller gegen null geht als irgendeine Potenz von h , d.h.:

$$\lim_{h \rightarrow 0} \frac{R(h)}{h^n} = 0 \quad \text{für alle } n \in \mathbf{N}.$$

A: Die Funktion $R(h) = \exp(-1/h^2)$ hat die Eigenschaft, dass sie für $h \rightarrow 0$ schneller verschwindet als jede Potenz von h , und doch ist sie in einer beliebig kleinen Umgebung von $h = 0$ von null verschieden.

Insbesondere folgt für eine Funktion f , die in der Umgebung von 0 beliebig oft differenzierbar ist:

$$f(x) = \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(0) x^n + R(h).$$

Ist $R(h) = 0$ für eine nicht verschwindende Umgebung von $h = 0$, ist also die Funktion $f(x)$ eindeutig durch ihre Potenzreihenentwicklung gegeben, so bezeichnet man die Funktion $f(x)$ als *reell analytisch* in dieser Umgebung.

Ist eine Funktion k -mal stetig differenzierbar, so gilt:

$$f(x+h) = \sum_{n=0}^k \frac{1}{n!} f^{(n)}(x) h^n + o(h^k),$$

wobei

$$\lim_{h \rightarrow 0} \frac{o(h^k)}{h^k} = 0.$$

1.4 Integrale

Auch dieser Abschnitt soll nur die elementaren Begriffe im Zusammenhang mit Integralen zusammenfassen. Allgemeinere Definitionen von Integralen werden in der Mathematik im Rahmen der *Maßtheorie* [*measure theory*] behandelt. Wir betrachten im Folgenden meist stetige Funktionen über Bereichen der reellen Achse, allerdings lässt sich der Integrationsbegriff auch auf Funktionen mit Sprungstellen verallgemeinern, sofern diese Sprungstellen nirgendwo „dicht“ liegen.

1.4.1 Definitionen

Definition: Eine Funktion $F : U \subset \mathbf{R} \rightarrow \mathbf{R}$ heißt *Stammfunktion* zu einer Funktion $f : U \subset \mathbf{R} \rightarrow \mathbf{R}$, wenn gilt:

$$F'(x) \equiv \frac{dF(x)}{dx} = f(x). \quad (1.4)$$

Die Stammfunktion ist eindeutig bis auf eine von x unabhängige Konstante, d.h., mit $F(x)$ ist auch $F(x) + c$ Stammfunktion von $f(x)$.

Definition: Sei $[a, b] \subset U \subset \mathbf{R}$ ein Intervall und $f : U \rightarrow \mathbf{R}$ eine stetige Funktion mit Stammfunktion F , dann bezeichnen wir mit

$$I = \int_a^b dx f(x) = F(b) - F(a) \quad (1.5)$$

das *bestimmte Integral* [*definite integral*] der Funktion f über das Intervall $[a, b]$.

Anschaulich entspricht I der Fläche unter der Funktion f (zwischen dem Graph von f und der x -Achse) im Intervall $[a, b]$. Dabei werden Bereiche mit negativen Werte von f auch negativ berechnet. Wegen der Differenz auf der rechten Seite von Gl. 1.5 ist das bestimmte Integral unabhängig von der freien Konstanten der Stammfunktion und eindeutig.

Insbesondere folgt:

$$\int_a^b dx \frac{df(x)}{dx} = f(b) - f(a).$$

Wird das Integral ohne Grenzen angegeben, so spricht man vom *unbestimmten Integral* [*indefinite integral*] über die Funktion f und meint damit meist die Stammfunktion F :

$$F(x) = \int dx f(x).$$

1.4.2 Integrationsregeln

Auch das Integral ist eine lineare Abbildung für Funktionen, d.h., es gilt:

$$\int_a^b dx (\alpha f(x) + \beta g(x)) = \alpha \int_a^b dx f(x) + \beta \int_a^b dx g(x). \quad (1.6)$$

Allgemeiner gilt für stetige Funktionen f , g und h sogar:

$$\int_a^b dx h(x)(\alpha f(x) + \beta g(x)) = \alpha \int_a^b dx h(x)f(x) + \beta \int_a^b dx h(x)g(x). \quad (1.7)$$

A: Erweitert man den Funktionenraum für h geeignet und schränkt gleichzeitig die Möglichkeiten für f und g ein (auf den Raum so genannter *Testfunktionen*), so lässt sich jede lineare Abbildung auf den Testfunktionen als Integral über diese *verallgemeinerten Funktionen* bzw. *Distributionen* h schreiben.

Aus der Produktregel folgt:

$$f(b)g(b) - f(a)g(a) = \int_a^b dx \frac{d(f(x)g(x))}{dx} = \int_a^b dx f'(x)g(x) + \int_a^b dx f(x)g'(x),$$

bzw.

$$\int_a^b dx f'(x)g(x) = - \int_a^b dx f(x)g'(x) + f(b)g(b) - f(a)g(a). \quad (1.8)$$

Für $x = x(y)$ (lässt sich x also als Funktion eines anderen Arguments y schreiben) gilt:

$$\int_a^b f(x) dx = \int_{y_a}^{y_b} f(x(y)) \left| \frac{dx(y)}{dy} \right| dy, \quad (1.9)$$

wobei $a = x(y_a)$ und $b = x(y_b)$. Diese Formel beschreibt das Integral bei einer Variablentransformation.

Kapitel 2

Vektoren

Die Beschreibung von Bewegungen (die sogenannte *Kinematik*) bildet einen wesentlichen Bestandteil der theoretischen Physik. Demgegenüber beschäftigt sich die *Dynamik* mit den Ursachen der Bewegungen bzw. Bewegungsänderungen. Zur Beschreibung einer Bewegung gehört, den Ort eines Teilchens zu jedem Zeitpunkt angeben zu können. Die möglichen Orte eines Teilchens bilden unseren Anschauungsraum. Dieser Anschauungsraum kann durch den dreidimensionalen Raum $\mathbf{R}^3$ beschrieben werden. Dass es sich beim $\mathbf{R}^3$ um einen Vektorraum handelt, ist eher ein Zufall. Der Anschauungsraum $\mathbf{R}^3$ bildet eine Mannigfaltigkeit.

Betrachten wir jedoch die möglichen Geschwindigkeiten eines Teilchens, so handelt es sich ebenfalls um Elemente eines $\mathbf{R}^3$. In diesem Fall betrachtet man den $\mathbf{R}^3$ jedoch als einen *Tangentialraum* an die Mannigfaltigkeit der möglichen Orte. Dieser Tangentialraum ist immer ein Vektorraum.

Der $\mathbf{R}^3$ spielt also eine Doppelrolle: einmal als Mannigfaltigkeit der möglichen Orte eines Teilchens, und einmal als Vektorraum der möglichen Geschwindigkeiten. Diese Doppelrolle führt oftmals zur Identifikation von Konzepten, die eigentlich verschiedenen Ursprungs sind. Ein Ziel dieser Vorlesung wird sein, diese Doppelrolle hervorzuheben und zu betonen, an welchen Stellen die unterschiedlichen Konzepte auftreten. In diesem Kapitel behandeln wir die mathematische Struktur des Vektorraums, wobei der $\mathbf{R}^3$ oftmals als Beispiel herangezogen wird. In späteren Kapiteln werden wir uns dem $\mathbf{R}^3$ als Mannigfaltigkeit widmen.

2.1 Gruppen und Körper

Definition: Eine *Gruppe* [*group*] ist eine Menge G zusammen mit einer Abbildung $\circ : G \times G \rightarrow G$, die folgenden Bedingungen genügt:

- Assoziativität [*associativity*]: für alle $g_1, g_2, g_3 \in G$ gilt

$$(g_1 \circ g_2) \circ g_3 = g_1 \circ (g_2 \circ g_3).$$

- Existenz eines Einselements [*unit element*]: Es gibt ein Element e , sodass für alle $g \in G$ gilt

$$g \circ e = e \circ g = g.$$

- Existenz eines Inversen [*inverse element*]: Zu jedem $g \in G$ gibt es ein Element $g^{-1} \in G$, sodass

$$g \circ g^{-1} = g^{-1} \circ g = e.$$

Ist noch die folgende Bedingung erfüllt,

- Kommutativität [*commutativity*]: Für alle $g_1, g_2 \in G$ gilt

$$g_1 \circ g_2 = g_2 \circ g_1,$$

so bezeichnet man die Gruppe als *kommutative Gruppe* [*commutative group*] oder auch *Abel'sche Gruppe*. Bei einer kommutativen Gruppe schreibt man für $\circ$ auch häufig $+$.

Beispiel: die ganzen Zahlen bilden eine Gruppe bezüglich der Addition. Das Einselement ist die Null $e = 0$, und zu einer Zahl z ist $-z$ das inverse Element.

Definition: Ein *Körper* [*field*] ist eine Menge K (mit mindestens zwei Elementen) zusammen mit zwei Abbildungen $+: K \times K \rightarrow K$ und $\cdot: K \times K \rightarrow K$, sodass folgende Bedingungen erfüllt sind:

- Bezüglich $+$ bildet K eine kommutative Gruppe. Das Einselement wird als „0“ bezeichnet. Das zu einem Element $z \in K$ inverse Element bezeichnet man oft als $-z$.
- Bezüglich $\cdot$ bildet $K/\{0\}$ (die Menge K ohne das Element 0) eine kommutative Gruppe. Das Einselement bezeichnet man als „1“.
- Distributivgesetz [*distributive law*]: Für alle $a, b, c \in K$ gilt

$$a \cdot (b + c) = (a \cdot b) + (a \cdot c).$$

Beispiel: Die rationalen Zahlen [*rational numbers*] $\mathbf{Q}$ bilden einen Körper bezüglich der üblichen Addition und Multiplikation.

Beispiel: Die reellen Zahlen [*real numbers*] $\mathbf{R}$ bilden einen Körper bezüglich der üblichen Addition und Multiplikation.

Beispiel: Die Menge $F = \{0, 1\}$ bildet einen Körper bezüglich der Addition modulo 2 (d.h., $1 + 1 = 0$) und der üblichen Multiplikation.

A: Der wesentliche strukturelle Unterschied zwischen den rationalen und den reellen Zahlen besteht darin, dass die reellen Zahlen noch *vollständig* [*complete*] sind. Das bedeutet, dass der Grenzwert jeder Cauchy-Folge von reellen Zahlen auch wieder eine reelle Zahl ist. Dies gilt nicht für die rationalen Zahlen. Es gibt Cauchy-konvergente Folgen von rationalen Zahlen, deren Grenzwert eine reelle aber nicht rationale Zahl ist.

2.2 Vektorräume

Definition: Ein *Vektorraum* [*vector space*] über einem Körper K ist eine Menge V mit zwei Verknüpfungsregeln [*composition rules*],

$$+ : V \times V \rightarrow V \quad \text{und} \quad \cdot : K \times V \rightarrow V, \quad (2.1)$$

die folgenden Bedingungen genügen:

Bezüglich $+$ bildet V eine kommutative Gruppe: Für drei beliebige Elemente $\vec{a}, \vec{b}, \vec{c} \in V$ gilt:

$$\begin{aligned} \vec{a} + (\vec{b} + \vec{c}) &= (\vec{a} + \vec{b}) + \vec{c} && \text{Assoziativität,} \\ \vec{a} + \vec{b} &= \vec{b} + \vec{a} && \text{Kommutativität.} \end{aligned}$$

Es gibt ein Element $\vec{0}$, sodass für alle $\vec{a} \in V$ gilt:

$$\vec{a} + \vec{0} = \vec{0} + \vec{a} = \vec{a} \quad \text{Existenz eines Eins-Elements,}$$

und zu jedem Element $\vec{a} \in V$ gibt es ein Element $\vec{a}' = -\vec{a}$ mit der Eigenschaft:

$$\vec{a} + (-\vec{a}) = (-\vec{a}) + \vec{a} = \vec{0} \quad \text{Existenz eines Inversen.}$$

Für die Multiplikation gilt das Assoziativgesetz: Für alle $\alpha, \beta \in K$ und alle $\vec{a} \in V$ gilt:

$$\alpha \cdot (\beta \cdot \vec{a}) = (\alpha\beta) \cdot \vec{a}.$$

Außerdem gelten für alle $\vec{a}, \vec{b} \in V$ und alle $\alpha, \beta \in K$ die beiden Distributivgesetze:

$$\alpha \cdot (\vec{a} + \vec{b}) = \alpha \cdot \vec{a} + \alpha \cdot \vec{b} \quad \text{und} \quad (\alpha + \beta) \cdot \vec{a} = \alpha \cdot \vec{a} + \beta \cdot \vec{a}.$$

Der Vollständigkeit halber benötigt man noch:

$$1 \cdot \vec{a} = \vec{a}.$$

Wir werden im Folgenden ausschließlich Vektorräume über dem Körper der reellen Zahlen betrachten. In der Physik spielen aber auch Vektorräume über dem Körper der komplexen Zahlen eine wichtige Rolle. Ähnlich wie bei der Multiplikation reeller Zahlen lässt man den Punkt $\cdot$ für die Multiplikation eines Vektors mit einer reellen Zahl meist weg.

2.2.1 Beispiel 1: Verschiebungen in einer Ebene

Bekannt ist die Darstellung von Vektoren als „Äquivalenzklasse von Pfeilen“ in einer Ebene. Zwei Pfeile gelten dabei als äquivalent, wenn sie dieselbe Richtung und dieselbe Länge haben. Es kommt also nicht auf den „Angriffspunkt“ des Pfeiles an und wir können als Repräsentanten einer solchen Äquivalenzklasse immer einen Pfeil mit einem geeigneten Angriffspunkt wählen. Eine solche Äquivalenzklasse kann man auch als (globale) Verschiebung in einer Ebene interpretieren: Jedes Element der Ebene wird in Richtung des Pfeils um die Länge des Pfeils verschoben.

Die Addition zweier Vektoren wird üblicherweise über das so genannte *Vektorparallelogramm* (im Fall von Kräften auch als *Kräfteparallelogramm* bekannt) definiert. Der Null-Vektor $\vec{0}$ entspricht dem Pfeil der Länge 0, d.h. keiner Verschiebung. Die Multiplikation mit einer reellen Zahl entspricht einer Skalierung der Länge des Vektors um diese Zahl, die Richtung bleibt unverändert. Die Vektorraumaxiome lassen sich für dieses Beispiel leicht verifizieren.

2.2.2 Beispiel 2: n -Tupel von Zahlen

Ein weiteres wichtiges Beispiel für einen Vektorraum bilden die n -Tupel reeller Zahlen:

$$V = \{(x^1, x^2, \dots, x^n) | x^i \in \mathbf{R}\}.$$

Die einzelnen Zahlen x^i bezeichnet man auch als die *Komponenten* [*components*] des Vektors. Im Rahmen einer Konvention, die sich als sehr sinnvoll erwiesen

hat, schreibt man die Indizes der Komponenten eines Vektors meist nach oben. (Lediglich im Falle des so genannten kartesischen Skalarprodukts, das wir später verwenden werden, spielt die Stellung der Indizes keine Rolle. Die Gründe und die Zusammenhänge werden später noch erläutert.) Eine Verwechslung von der Komponente x^i mit der Potenzierung der Zahl x zur i -ten Potenz sollte in der Praxis ausgeschlossen sein. Im Zweifelsfall schreibt man beispielsweise $(x^i)^2$ um das Quadrat der i -ten Komponente auszudrücken.

Die Addition von Vektoren sowie die Multiplikation mit einer reellen Zahl sind komponentenweise definiert:

$$\begin{aligned}(x^1, x^2, \dots, x^n) + (y^1, y^2, \dots, y^n) &= (x^1 + y^1, x^2 + y^2, \dots, x^n + y^n) \\ \alpha \cdot (x^1, x^2, \dots, x^n) &= (\alpha x^1, \alpha x^2, \dots, \alpha x^n).\end{aligned}$$

Auch hier lassen sich die Vektorraumaxiome leicht überprüfen.

2.3 Dimension, Basis, lineare Unabhängigkeit

Definition: Ein Satz von Vektoren $\{\vec{v}_i\}_{i=1, \dots, n}$ ($\vec{v}_i \neq \vec{0}$) heißt *linear unabhängig* [*linearly independent*], wenn aus

$$\sum_{i=1}^n \alpha^i \vec{v}_i = \vec{0} \tag{2.2}$$

folgt: $\alpha^i \equiv 0$ für alle $i = 1, \dots, n$. Andernfalls heißen die Vektoren *linear abhängig* [*linearly dependent*].

Definition: Die maximale Zahl n , sodass es n linear unabhängige (nicht-verschwindende) Vektoren gibt, bezeichnet man als die *Dimension* [*dimension*] des Vektorraums.

A: Wir werden im Folgenden nur Vektorräume betrachten, deren Dimension endlich ist. Es gibt auch Vektorräume mit unendlicher Dimension, wobei man noch zwischen abzählbarer und überabzählbarer Dimension unterscheiden muss. Auf die Besonderheiten unendlich dimensionaler Vektorräume gehen wir hier nicht näher ein, obwohl einige dieser Konzepte im Zusammenhang mit der Quantenmechanik von großer Bedeutung sind.

Für einen Vektorraum der Dimension n bilden je n linear unabhängige Vektoren $\{\vec{e}_i\}$ eine *Basis* [*basis*]. (Bei der Durchnummerierung der Basisvektoren setzt

man die Indizes nach unten.) Jeder Vektor $\vec{x} \in V$ lässt sich als Linearkombination dieser Basis schreiben:

$$\vec{x} = \sum_{i=1}^n x^i \vec{e}_i. \quad (2.3)$$

Die $n + 1$ Vektoren $\{\vec{e}_i\}$ plus $\vec{x}$ müssen nämlich linear abhängig sein. Daher gibt es einen nicht identisch verschwindenden Satz reeller Zahlen $\alpha^0, \{\alpha^i\}$ mit

$$\alpha^0 \vec{x} + \sum_i \alpha^i \vec{e}_i = 0.$$

α^0 kann nicht 0 sein, da die Vektoren $\{\vec{e}_i\}$ eine Basis bilden sollen und somit linear unabhängig sind. Also kann man obige Gleichung durch α^0 dividieren und definiert: $x^i = -\alpha^i/\alpha^0$. Wiederum bezeichnet man die $\{x^i\}$ als die Komponenten des Vektors $\vec{x}$ bezüglich der Basis $\{\vec{e}_i\}$.

Hat man zwei Vektoren $\vec{x}$ und $\vec{y}$ in der Basis $\{\vec{e}_i\}$ ausgedrückt,

$$\vec{x} = \sum_i x^i \vec{e}_i \quad \text{und} \quad \vec{y} = \sum_i y^i \vec{e}_i,$$

so folgt für die Summe der beiden Vektoren:

$$\vec{x} + \vec{y} = \sum_i (x^i + y^i) \vec{e}_i.$$

Hat man sich einmal auf eine feste Basis $\{\vec{e}_i\}$ geeinigt, lässt sich somit jeder Vektor als n -Tupel seiner Komponenten ausdrücken und wir erhalten den Vektorraum aus Beispiel 2.

2.4 Der duale Vektorraum

Eine Abbildung $\omega : V \rightarrow \mathbf{R}$ bezeichnet man als *linear*, wenn für alle $\alpha, \beta \in \mathbf{R}$ und alle $\vec{x}, \vec{y} \in V$ gilt:

$$\omega(\alpha \vec{x} + \beta \vec{y}) = \alpha \omega(\vec{x}) + \beta \omega(\vec{y}). \quad (2.4)$$

Für zwei lineare Abbildungen ω_1 und ω_2 können wir durch folgende Vorschrift eine Addition sowie eine Multiplikation mit reellen Zahlen definieren:

$$(\alpha \omega_1 + \beta \omega_2)(\vec{x}) = \alpha \omega_1(\vec{x}) + \beta \omega_2(\vec{x}) \quad \text{für alle } \vec{x} \in V.$$

Man kann sich leicht davon überzeugen, dass diese Definition die Menge der linearen Abbildungen von V in $\mathbf{R}$ zu einem Vektorraum macht. Diesen Vektorraum nennt man den *Dualraum* [*dual space*] zu V und bezeichnet ihn mit V^* .

Die Anwendung eines Elements $\omega \in V^*$ auf ein Element $\vec{x} \in V$, also $\omega(\vec{x})$, bezeichnet man auch als *inneres Produkt* [*inner product*] von ω mit $\vec{x}$.

Ein einfaches Beispiel einer linearen Abbildung von einem Vektorraum in den zugrundeliegenden Körper (und gleichzeitig ein Beispiel, aus dem sich der allgemeine Fall zusammensetzen lässt) ist die Abbildung, die jedem Vektor bezüglich einer vorgegebenen Basis eine bestimmte Komponente - beispielsweise die i -te Komponente - zuordnet.

Wegen der Bedingung (2.4) ist ein Element $\omega \in V^*$ eindeutig gegeben, wenn seine Anwendung auf die Basisvektoren von V bekannt ist:

$$\omega(\vec{x}) = \omega\left(\sum_i x^i \vec{e}_i\right) = \sum_i x^i \omega(\vec{e}_i). \quad (2.5)$$

Die n Zahlen $\omega_i = \omega(\vec{e}_i)$ legen das Element ω somit eindeutig fest. Man kann sich nun leicht davon überzeugen, dass V^* dieselbe Dimension wie V hat.

Ist für V eine Basis $\{\vec{e}_i\}$ gegeben, so kann man für V^* eine ausgezeichnete Basis $\{\epsilon^i\}$ durch die folgende Vorschrift definieren:

$$\epsilon^i \vec{e}_j = \delta_j^i := \begin{cases} 1 & \text{für } i = j \\ 0 & \text{sonst} \end{cases} \quad (2.6)$$

Das Symbol δ_j^i bezeichnet man als *Kronecker- δ* . Die so definierte Basis im Dualraum bezeichnet man als *duale Basis* [*dual basis*].

2.4.1 Index-Schreibweise und Summenkonvention

Wir haben die Indizes zur Nummerierung der Basisvektoren eines Vektorraums V nach unten gesetzt - $\{\vec{e}_i\}$ - und die Indizes der Komponenten eines Vektors in dieser Basis nach oben: $\vec{x} = \sum_i x^i \vec{e}_i$. Für den dualen Vektorraum V^* gilt die umgekehrte Konvention: Die Indizes zur Nummerierung der dualen Basis $\{\epsilon^i\}$ wurden nach oben gesetzt, während wir die Indizes der Komponenten einer linearen Abbildung ω in dieser Basis nach unten schreiben: $\omega = \sum_i \omega_i \epsilon^i$. Wenden wir die Abbildung ω auf einen Vektor $\vec{x}$ an, so folgt:

$$\omega(\vec{x}) = \left(\sum_i \omega_i \epsilon^i\right) \left(\sum_j x^j \vec{e}_j\right) = \sum_{i,j} \omega_i x^j \epsilon^i(\vec{e}_j) = \sum_{i,j} \omega_i x^j \delta_j^i = \sum_i \omega_i x^i.$$

(Zur Überprüfung schreibe man sich diese Beziehung für den Fall $n = 3$ explizit hin und verwende die bisher angeführten Regeln.)

Man erkennt, dass ein Index, über den summiert wird, immer zweimal auftritt - einmal oben und einmal unten. Bei einer konsequenten Einhaltung der Schreibweise muss das auch immer so sein, daher verwendet man oft die folgende *Einstein'sche Summenkonvention*: Tritt ein Index einmal oben und einmal unten auf, so ist über ihn zu summieren. Mit dieser Summenkonvention würde obige Gleichung folgendermaßen aussehen:

$$\omega(\vec{x}) = (\omega_i \epsilon^i)(x^j \vec{e}_j) = \omega_i x^j \epsilon^i(\vec{e}_j) = \omega_i x^j \delta_j^i = \omega_i x^i.$$

Wenn ein Index auf einer Seite einer Gleichung nur einmal auftritt (und es wird nicht über ihn summiert), so muss er auch auf der anderen Seite der Gleichung auftreten, und zwar an derselben Stelle (oben bzw. unten). Als Beispiel bestimmen wir die lineare Abbildung, die jedem Vektor die i -te Komponente zuordnet. Diese Abbildung ist durch das entsprechende Element der dualen Basis gegeben:

$$\epsilon^i(\vec{x}) = \epsilon^i(x^j \vec{e}_j) = x^j \epsilon^i(\vec{e}_j) = x^j \delta_j^i = x^i.$$

Die Summenkonvention ist in der allgemeinen Relativitätstheorie von großer Bedeutung. In Kap. 6 werden wir im Zusammenhang mit den Ableitungen von Feldern einige Identitäten verwenden, deren Beweis ohne die Summenkonvention oft recht aufwendig ist.

Abschließend soll noch vermerkt werden, dass man die Elemente von V in Komponentenschreibweise oft als Spaltenvektor [*column vector*] schreibt, während die Elemente von V^* in Komponentenschreibweise als Zeilenvektor [*row vector*] geschrieben werden:

$$\vec{x} \simeq \begin{pmatrix} x^1 \\ x^2 \\ \vdots \\ x^n \end{pmatrix} \quad \omega = (\omega_1, \omega_2, \dots, \omega_n).$$

Für die Anwendung von ω auf $\vec{x}$ gilt dann „Linke Zeile multipliziert rechte Spalte“:

$$\omega(\vec{x}) = (\omega_1, \omega_2, \dots, \omega_n) \begin{pmatrix} x^1 \\ x^2 \\ \vdots \\ x^n \end{pmatrix} = \omega_1 x^1 + \omega_2 x^2 + \dots + \omega_n x^n.$$

2.5 Das Skalarprodukt

Definition: Eine Abbildung $g(\cdot, \cdot) : V \times V \rightarrow \mathbf{R}$ heißt *Skalarprodukt* [*scalar product*] (oder auch *inneres Produkt* [*inner product*]) auf einem Vektorraum V , wenn für alle Vektoren $\vec{x}, \vec{y}, \vec{z}$ gilt:

$$\begin{aligned} g(\vec{x}, \vec{y}) &= g(\vec{y}, \vec{x}) && \text{Symmetrie} \\ g(\vec{x}, \alpha\vec{y} + \beta\vec{z}) &= \alpha g(\vec{x}, \vec{y}) + \beta g(\vec{x}, \vec{z}) && \text{Bilinearität.} \end{aligned}$$

Gilt außerdem $g(\vec{x}, \vec{x}) > 0$ für alle $\vec{x} \neq \vec{0}$, so bezeichnet man das Skalarprodukt als *positiv definit* und *nicht entartet* [*non-degenerate*]. Ein positives, nicht entartetes Skalarprodukt nennt man auch manchmal eine *Metrik* [*metric*]. Der Grund ist, dass die Definition

$$d(\vec{x}, \vec{y}) = \sqrt{g(\vec{x} - \vec{y}, \vec{x} - \vec{y})}$$

die Bedingungen einer Metrik erfüllt, also:

$$\begin{aligned} d(\vec{x}, \vec{y}) &\geq 0 && \text{Positivität} \\ d(\vec{x}, \vec{y}) &= 0 \Rightarrow \vec{x} = \vec{y} \\ d(\vec{x}, \vec{y}) &= d(\vec{y}, \vec{x}) && \text{Symmetrie} \\ d(\vec{x}, \vec{y}) + d(\vec{y}, \vec{z}) &\geq d(\vec{x}, \vec{z}) && \text{Dreiecksungleichung.} \end{aligned}$$

A: Allgemeiner heißt ein Skalarprodukt *nicht entartet*, wenn aus $g(\vec{x}, \vec{y}) = 0$ für alle $\vec{y} \in V$ folgt, dass $\vec{x} = \vec{0}$. Ist $g(\vec{x}, \vec{x}) \geq 0$ für alle $\vec{x} \in V$, so ist das Skalarprodukt nicht entartet, wenn das Gleichheitszeichen nur für $\vec{x} = \vec{0}$ angenommen wird.

Wegen der Symmetrie des Skalarprodukts gilt die Linearität auch im ersten Argument (daher auch die Bezeichnung *Bi*-linearität: das Skalarprodukt ist bezüglich beider Argumente linear):

$$g(\alpha\vec{x} + \beta\vec{y}, \vec{z}) = \alpha g(\vec{x}, \vec{z}) + \beta g(\vec{y}, \vec{z}). \quad (2.7)$$

Ganz allgemein kann man daher schreiben (Summenkonvention!):

$$g(\vec{x}, \vec{y}) = g(x^i \vec{e}_i, y^j \vec{e}_j) = x^i y^j g(\vec{e}_i, \vec{e}_j) = x^i y^j g_{ij}. \quad (2.8)$$

Hierbei wurde definiert:

$$g_{ij} := g(\vec{e}_i, \vec{e}_j). \quad (2.9)$$

Die Symmetriebedingung bedeutet:

$$g_{ij} = g_{ji}. \quad (2.10)$$

Aus der Positivität,

$$\sum_{ij} g_{ij} x^i x^j > 0 \quad \text{für beliebige } \{x^i\} \ (\vec{x} \neq 0), \quad (2.11)$$

folgt, dass die Diagonalelemente positiv sein müssen:

$$g_{ii} > 0 \quad \text{für alle } i.$$

Die Einschränkungen an die nicht diagonalen Elemente sind (abgesehen von der Symmetrie) weniger einfach.

Als „Länge“ eines Vektors definiert man:

$$|\vec{x}| = \sqrt{g(\vec{x}, \vec{x})}. \quad (2.12)$$

Wir haben das Skalarprodukt als eine Abbildung von $g : V \times V \rightarrow \mathbf{R}$ definiert. Ein solches Skalarprodukt erlaubt aber auch die Definition einer Abbildung $g : V \rightarrow V^*$ nach folgender Vorschrift: Das Element $\vec{x} \in V$ wird abgebildet auf das Element $\mathbf{g}(\vec{x}) \in V^*$, sodass für alle $\vec{y} \in V$ gilt:

$$\mathbf{g}(\vec{x})(\vec{y}) = g(\vec{x}, \vec{y}). \quad (2.13)$$

Auf der linken Seite dieser Definitionsgleichung steht ein Element des Dualraums — $\mathbf{g}(\vec{x}) \in V^*$ — angewandt auf einen Vektor $\vec{y} \in V$. Auf der rechten Seite steht das Skalarprodukt angewandt auf zwei Vektoren in V . Ausgedrückt in der kanonischen Basis von V^* (also $\epsilon^j(\vec{e}_i) = \delta_i^j$) folgt für die Komponenten von $\mathbf{g}(\vec{x})$:

$$\mathbf{g}(\vec{x})_i = \sum_j g_{ij} x^j. \quad (2.14)$$

Mithilfe des Skalarprodukts kann man also Indizes „von oben nach unten ziehen“. Oftmals schreibt man vereinfacht:

$$g(\vec{x})_i = x_i. \quad (2.15)$$

Bei dieser Schreibweise wird nochmals deutlich, weshalb die Unterscheidung zwischen oben und unten stehenden Indizes im Allgemeinen so wichtig ist.

Die Nicht-Entartetheit des Skalarprodukts bedeutet, dass die Abbildung $g : V \rightarrow V^*$ bijektiv und damit eindeutig umkehrbar ist. Es gibt also eine Abbildung $g^{-1} : V^* \rightarrow V$, sodass $g^{-1} \circ g = \mathbf{1}$, oder ausgedrückt in der Indexschreibweise:

$$\sum_k (g^{-1})^{ik} g_{kj} = \delta_j^i. \quad (2.16)$$

Vereinfacht schreibt man meist:

$$(g^{-1})^{ij} = g^{ij}.$$

Mithilfe dieser inversen Abbildung lassen sich „Indizes von unten nach oben ziehen“:

$$g^{-1}(\omega)^i = \sum_j g^{ij} \omega_j. \quad (2.17)$$

A: Gerade diese letzten Bemerkungen gelten in dieser Form nicht mehr für Vektorräume mit unendlich vielen Dimensionen. Der duale Vektorraum kann dort „mehr“ Elemente enthalten als der Vektorraum selber.

2.6 Tensoren

Wir haben gesehen, dass die Elemente des Dualraums zu einem Vektorraum V aus den linearen Abbildungen von V in die reellen Zahlen bestehen, $\omega : V \rightarrow \mathbf{R}$. Das Skalarprodukt lässt sich auffassen als bilineare Abbildung von $V \times V$ in die reellen Zahlen, $g : V \times V \rightarrow \mathbf{R}$, oder auch als lineare Abbildung von V in seinen Dualraum, $g : V \rightarrow V^*$. Umgekehrt lässt sich g^{-1} als lineare Abbildung von V^* nach V oder auch als bilineare Abbildung von $V^* \times V^* \rightarrow \mathbf{R}$ auffassen.

Ganz allgemein bezeichnet man multilineare Abbildungen von kartesischen Produkten von V und/oder V^* in die reellen Zahlen oder auch in kartesische Produkte von V und/oder V^* als *Tensoren* [*tensors*]. Betrachten wir als Beispiel eine Abbildung $\tilde{T} : V \times V \times V \rightarrow \mathbf{R}$. Die Abbildung $\tilde{T}$ ordnet also drei beliebigen Vektoren aus V eine reelle Zahl zu. Die Multilinearität bedeutet dabei, dass $\tilde{T}$ in allen drei Argumenten linear ist, also für alle $\vec{x}_i, \vec{y}_i, \vec{z}_i \in V$ und alle $\alpha, \beta \in \mathbf{R}$ gilt:

$$\begin{aligned} \tilde{T}(\alpha \vec{x}_1 + \beta \vec{x}_2, \vec{y}, \vec{z}) &= \alpha \tilde{T}(\vec{x}_1, \vec{y}, \vec{z}) + \beta \tilde{T}(\vec{x}_2, \vec{y}, \vec{z}) \\ \tilde{T}(\vec{x}, \alpha \vec{y}_1 + \beta \vec{y}_2, \vec{z}) &= \alpha \tilde{T}(\vec{x}, \vec{y}_1, \vec{z}) + \beta \tilde{T}(\vec{x}, \vec{y}_2, \vec{z}) \\ \tilde{T}(\vec{x}, \vec{y}, \alpha \vec{z}_1 + \beta \vec{z}_2) &= \alpha \tilde{T}(\vec{x}, \vec{y}, \vec{z}_1) + \beta \tilde{T}(\vec{x}, \vec{y}, \vec{z}_2). \end{aligned}$$

Wegen der Multilinearität liegt die Abbildung $\tilde{T}$ bereits eindeutig fest, wenn ihre Wirkung auf die Basisvektoren gegeben ist, wenn also die Werte

$$\tilde{T}_{ijk} = \tilde{T}(\vec{e}_i, \vec{e}_j, \vec{e}_k) \quad (2.18)$$

bekannt sind. In Komponentenschreibweise kann man daher auch schreiben:

$$\tilde{T}(\vec{x}, \vec{y}, \vec{z}) = \sum_{ijk} \tilde{T}_{ijk} x^i y^j z^k. \quad (2.19)$$

Hätte man andererseits eine bilineare Abbildung $T : V \times V \rightarrow V$ gegeben, So könnte man das Bild zweier Vektoren $T(\vec{x}, \vec{y})$ wieder nach den Basisvektoren entwickeln und erhielte:

$$T(\vec{x}, \vec{y}) = \sum_i T(\vec{x}, \vec{y})^i \vec{e}_i = \sum_{ijk} x^j y^k T(\vec{e}_j, \vec{e}_k)^i \vec{e}_i = \sum_{ijk} T_{jk}^i x^j y^k \vec{e}_i. \quad (2.20)$$

Die Abbildung T lässt sich bezüglich einer Basis in V also durch die Komponenten T_{jk}^i charakterisieren.

Andererseits kann man T auch als Abbildung von $V^* \times V \times V \rightarrow \mathbf{R}$ auffassen. Dazu definiert man für $\omega \in V^*$ und $\vec{x}, \vec{y} \in V$:

$$T(\omega, \vec{x}, \vec{y}) = \omega(T(\vec{x}, \vec{y})).$$

Sowohl T als auch $\tilde{T}$ bezeichnet man als Tensoren dritter Stufe. Allgemeiner heißt eine multilineare Abbildung T von einem m -fachen kartesischen Produkt von V und/oder V^* in ein $n - m$ -faches kartesisches Produkt von V und/oder V^* als Tensor n -ter Stufe.

2.7 Der euklidische Vektorraum

Einen Vektorraum mit einem positiv definiten Skalarprodukt bezeichnet man auch als *euklidischen Vektorraum* [*Euclidean vector space*]. Wir werden im Folgenden ein spezielles Skalarprodukt wählen, das man manchmal als kartesisches Skalarprodukt bezeichnet, und einige bekannte Tatsachen aus der Geometrie der Ebene oder des dreidimensionalen Raumes herleiten und verallgemeinern. Außerdem werden wir speziell für den dreidimensionalen Vektorraum das so genannte Vektorprodukt (oder auch Kreuzprodukt) sowie das Spatprodukt betrachten.

2.7.1 Das kartesische Skalarprodukt

Das kartesische Skalarprodukt hat folgende Form:

$$g(\vec{x}, \vec{y}) = \delta_{ij} x^i y^j = x_i y^i = x^1 y^1 + x^2 y^2 + x^3 y^3 + \dots \quad (2.21)$$

Da, etwas lax gesprochen,

$$x_i = x^i \quad \left(\text{genauer } x_i = \begin{cases} x^j & \text{für } i = j \\ 0 & \text{sonst} \end{cases} \right),$$

spielt es für die konkreten Zahlenwerte der Komponenten keine Rolle, ob die Indizes unten oder oben geschrieben werden. Trotzdem empfiehlt es sich, bei allgemeinen Rechnungen die Indexkonvention beizubehalten.

Statt der umständlichen Notation $g(\vec{x}, \vec{y})$ kennzeichnen wir das karthetische Skalarprodukt in Zukunft einfach durch einen „Multiplikationspunkt“:

$$g(\vec{x}, \vec{y}) \equiv \vec{x} \cdot \vec{y}.$$

Setzen wir $x = y$ so folgt:

$$\vec{x} \cdot \vec{x} \equiv x^2 = (x^1)^2 + (x^2)^2 + (x^3)^2 + \dots, \quad (2.22)$$

In diesem Ausdruck erkennen wir den euklidischen Abstand des Punktes $(x^1, x^2, x^3, \dots)$ vom Ursprung wieder bzw. die Länge der Verbindungslinie $\vec{x}$ vom Ursprung zum Punkt $(x^1, x^2, x^3, \dots)$. Wir schreiben auch oft:

$$|\vec{x}| = x = \sqrt{\vec{x} \cdot \vec{x}}.$$

Wir berechnen nun

$$(\vec{x} - \vec{y})^2 \equiv (\vec{x} - \vec{y}) \cdot (\vec{x} - \vec{y}) = x^2 + y^2 - 2\vec{x} \cdot \vec{y}. \quad (2.23)$$

Die Vektoren $\vec{x}$, $\vec{y}$ und $\vec{z} = \vec{x} - \vec{y}$ bilden die Seiten eines Dreiecks, und nach obiger Formel gilt:

$$z^2 = x^2 + y^2 - 2\vec{x} \cdot \vec{y}. \quad (2.24)$$

Durch Vergleich mit dem allgemeinen Kosinus-Satz der Trigonometrie können wir daher folgern:

$$\vec{x} \cdot \vec{y} = |\vec{x}| |\vec{y}| \cos \gamma, \quad (2.25)$$

wobei γ der von den Vektoren $\vec{x}$ und $\vec{y}$ eingeschlossene Winkel ist.

Über das Skalarprodukt können wir also sowohl die Länge eines Vektors bestimmen als auch den Winkel, der von zwei Vektoren eingeschlossen wird. Insbesondere gilt für zwei Vektoren, die beide vom Nullvektor verschieden sind (für die also gilt $|\vec{x}| \neq 0 \neq |\vec{y}|$):

$$\vec{x} \cdot \vec{y} = 0 \quad \Leftrightarrow \quad \gamma = 90^\circ. \quad (2.26)$$

Zwei Vektoren, für die das Skalarprodukt verschwindet, stehen somit senkrecht aufeinander.

Ist das Skalarprodukt $\cdot$ gegeben, so kann man eine Basis konstruieren, die folgende Bedingungen erfüllt:

$$\vec{e}_i \cdot \vec{e}_j = \delta_{ij}.$$

Die Basisvektoren haben alle die Länge 1 (es handelt sich um so genannte *Einheitsvektoren* [unit vectors]), und stehen jeweils senkrecht aufeinander. Eine solche Basis bezeichnet man auch als *Orthonormalbasis*. Für den $\mathbf{R}^3$ bietet sich daher folgende Basis an:

$$\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \vec{e}_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}, \quad \vec{e}_3 = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

Diese Basis bezeichnet man auch manchmal als die *kanonische Basis* [canonical basis].

2.7.2 Das Kreuzprodukt

Während die bisherigen Überlegungen unabhängig von der Anzahl der Dimensionen waren, betrachten wir nun eine Struktur speziell für einen dreidimensionalen Vektorraum. Es gibt zwar Verallgemeinerungen für beliebige Dimensionen, doch der in der Physik häufigste Fall bezieht sich auf den dreidimensionalen euklidischen Vektorraum.

Wir betrachten den dreidimensionalen Vektorraum $\mathbf{R}^3$, dargestellt durch Tripel reeller Zahlen $\mathbf{R}^3 \simeq \{(x^1, x^2, x^3) | x^i \in \mathbf{R}\}$. Auf diesem Raum definieren wir eine Abbildung, die zwei Vektoren wieder einen Vektor zuordnet,

$$\times : \mathbf{R}^3 \times \mathbf{R}^3 \longrightarrow \mathbf{R}^3,$$

nach folgender Vorschrift:

$$\begin{pmatrix} x^1 \\ x^2 \\ x^3 \end{pmatrix} \times \begin{pmatrix} y^1 \\ y^2 \\ y^3 \end{pmatrix} = \begin{pmatrix} x^2 y^3 - x^3 y^2 \\ -x^1 y^3 + x^3 y^1 \\ x^1 y^2 - x^2 y^1 \end{pmatrix}. \quad (2.27)$$

Drei Eigenschaften dieses *Kreuzprodukts* (manchmal auch *Vektorprodukt* oder *äußeres Produkt* genannt) zweier Vektoren lassen sich leicht überprüfen:

$$\begin{aligned} \vec{x} \times \vec{y} &= -(\vec{y} \times \vec{x}) && \text{Antisymmetry} \\ \vec{x} \times (\vec{y} + \vec{z}) &= (\vec{x} \times \vec{y}) + (\vec{x} \times \vec{z}) && \text{Bilinearität} \\ \vec{x} \times (\alpha \vec{y}) &= \alpha(\vec{x} \times \vec{y}) \end{aligned}$$

Aus der Antisymmetrie folgt unmittelbar:

$$\vec{x} \times \vec{x} = 0.$$

Das Kreuzprodukt eines Vektors mit sich selber verschwindet somit. Wegen der Linearität bedeutet dies, dass auch das Kreuzprodukt zweier kollinear Vektoren (d.h., $\vec{y} = \alpha \vec{x}$) verschwindet. Außerdem impliziert die Antisymmetrie, dass das Kreuzprodukt auch im ersten Argument linear ist.

Um ein Gespür für das Kreuzprodukt zu erhalten, betrachten wir zwei Vektoren $\vec{x}$ und $\vec{y}$, die in der 1-2-Ebene des $\mathbf{R}^3$ liegen, deren 3-Komponenten also verschwinden. Dann folgt:

$$\begin{pmatrix} x^1 \\ x^2 \\ 0 \end{pmatrix} \times \begin{pmatrix} y^1 \\ y^2 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ x^1 y^2 - x^2 y^1 \end{pmatrix}. \quad (2.28)$$

Der Vektor $\vec{x} \times \vec{y}$ besitzt nur eine 3-Komponente, steht also senkrecht sowohl auf $\vec{x}$ als auch auf $\vec{y}$. Bezeichnen wir mit α den Winkel, den $\vec{x}$ mit der x -Achse einschließt und mit β den Winkel, den $\vec{y}$ mit der x -Achse einschließt, so gilt:

$$x^1 = |\vec{x}| \cos \alpha, \quad x^2 = |\vec{x}| \sin \alpha, \quad y^1 = |\vec{y}| \cos \beta, \quad y^2 = |\vec{y}| \sin \beta.$$

Damit erhalten wir:

$$x^1 y^2 - x^2 y^1 = |\vec{x}| |\vec{y}| (\cos \alpha \sin \beta - \cos \beta \sin \alpha) = |\vec{x}| |\vec{y}| \sin(\beta - \alpha).$$

Der Winkel $\gamma = \beta - \alpha$ ist wieder der von $\vec{x}$ und $\vec{y}$ eingeschlossene Winkel. Damit folgt:

$$\begin{pmatrix} x^1 \\ x^2 \\ 0 \end{pmatrix} \times \begin{pmatrix} y^1 \\ y^2 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ |\vec{x}| |\vec{y}| \sin \gamma \end{pmatrix}. \quad (2.29)$$

$|\vec{x}| |\vec{y}| \sin \gamma$ ist gleich der Fläche, die von den beiden Vektoren aufgespannt wird.

A: Dieses Beispiel zeigt auch, wie sich das Kreuzprodukt für zweidimensionale Vektoren definieren lässt. Das Ergebnis ist in diesem Fall allerdings nicht wieder ein Vektor, sondern einfach eine Zahl — die Fläche des von den beiden Vektoren aufgespannten Parallelogramms.

Dieses Ergebnis lässt sich zu folgender Aussage verallgemeinern: Das Kreuzprodukt zweier Vektoren $\vec{x}$ und $\vec{y}$ ist ein Vektor, der senkrecht auf der von $\vec{x}$ und

$\vec{y}$ aufgespannten Ebene steht, und dessen Betrag gleich der Fläche des von $\vec{x}$ und $\vec{y}$ aufgespannten Parallelogramms ist.

Aus dieser Darstellung folgt auch nochmals für zwei nichtverschwindende Vektoren $\vec{x}$ und $\vec{y}$:

$$\vec{x} \times \vec{y} = 0 \quad \Leftrightarrow \quad \gamma = 0^\circ \text{ oder } \gamma = 180^\circ,$$

d.h., die beiden Vektoren sind kollinear.

Die Richtung des Vektors $\vec{x} \times \vec{y}$, für die nach obiger Aussage immer noch zwei Möglichkeiten in Frage kommen, lässt sich nach der „Rechte-Hand-Regel“ bestimmen: Spreizt man Daumen, Zeigefinger und Mittelfinger der rechten Hand, sodass sie (mehr oder weniger) senkrecht aufeinander stehen, und zeigt der Daumen in die Richtung von $\vec{x}$ und der Zeigefinger in die Richtung von $\vec{y}$, so zeigt der Mittelfinger in die Richtung von $\vec{x} \times \vec{y}$.

Das Kreuzprodukt erlaubt es somit, die von zwei Vektoren aufgespannte Fläche zu berechnen (gleich dem Betrag des Kreuzprodukts der beiden Vektoren), oder auch zu zwei linear unabhängigen Vektoren einen dritten Vektor zu konstruieren, der senkrecht auf den beiden Vektoren steht.

Für die kanonische Basis gilt:

$$\vec{e}_1 \times \vec{e}_2 = \vec{e}_3, \quad \vec{e}_2 \times \vec{e}_3 = \vec{e}_1, \quad \vec{e}_3 \times \vec{e}_1 = \vec{e}_2. \quad (2.30)$$

2.7.3 Der ϵ -Tensor

Wie wir gesehen haben, ist das Kreuzprodukt im $\mathbf{R}^3$ eine bilineare Abbildung von $\mathbf{R}^3 \times \mathbf{R}^3$ in den $\mathbf{R}^3$. Damit ist das Kreuzprodukt ein Tensor dritter Stufe. Diesen Tensor bezeichnet man auch als ϵ -Tensor.

Ausgedrückt in einer Basis des $\mathbf{R}^3$ lässt sich der ϵ -Tensor durch seine Komponenten ϵ^i_{jk} ausdrücken. Jeder dieser Indizes kann die Werte 1, 2 oder 3 annehmen. (Wir schreiben die Indizes zwar nach oben bzw. unten, wie es der korrekten Schreibweise entspricht, allerdings unterscheiden sich für das kartesische Skalarprodukt, das wir im Folgenden verwenden werden, die Komponenten der ϵ -Tensoren für unterschiedliche Indexstellungen nicht.) Der ϵ -Tensor ist durch folgende Vorschrift definiert:

$$\epsilon^i_{jk} = \begin{cases} 0 & \text{wenn zwei der Indizes gleich sind} \\ 1 & \text{wenn die drei Indizes zyklisch angeordnet sind} \\ -1 & \text{wenn die drei Indizes anti-zyklisch angeordnet sind} \end{cases} \quad (2.31)$$

Es gilt also:

$$\begin{aligned}\epsilon^1_{23} &= \epsilon^2_{31} = \epsilon^3_{12} &= 1 \\ \epsilon^2_{13} &= \epsilon^1_{32} = \epsilon^3_{21} &= -1 \\ \text{sonst } \epsilon^i_{jk} &= 0.\end{aligned}$$

Mit dieser Definition lässt sich das Kreuzprodukt in Komponenten schreiben:

$$(\vec{x} \times \vec{y})^i = \sum_{j,k} \epsilon^i_{jk} x^j y^k. \quad (2.32)$$

Als Beispiel betrachten wir die erste Komponente, also $i = 1$:

$$(\vec{x} \times \vec{y})^1 = \sum_{j,k} \epsilon^1_{jk} x^j y^k.$$

Bezüglich der Indizes j und k gilt die Summenkonvention, über diese Indizes ist also jeweils von 1 bis 3 zu summieren. Allerdings verschwindet der ϵ -Tensor, wenn zwei der Indizes gleich sind, insbesondere also wenn j oder k gleich 1 ist, oder wenn j gleich k ist. Damit bleiben nur noch zwei Möglichkeiten übrig: $i = 2$, $j = 3$ sowie $i = 3$, $j = 2$:

$$(\vec{x} \times \vec{y})^1 = \epsilon^1_{23} x^2 y^3 + \epsilon^1_{32} x^3 y^2.$$

Beim ersten Term auf der rechten Seite sind die Indizes in zyklischer Reihenfolge, also $\epsilon^1_{23} = 1$, beim zweiten Term sind sie antizyklisch, also $\epsilon^1_{32} = -1$. Damit erhalten wir das bekannte Ergebnis:

$$(\vec{x} \times \vec{y})^1 = x^2 y^3 - x^3 y^2. \quad (2.33)$$

Die Beziehungen aus Gl. 2.30 lassen sich nun in einer geschlossenen Form ausdrücken:

$$\vec{e}_j \times \vec{e}_k = \epsilon^i_{jk} \vec{e}_i$$

(wobei im Prinzip über i zu summieren ist, allerdings gibt es nur einen Wert für i , der zu dieser Summe beiträgt, wenn j und k festgelegt wurden.)

A: Es gibt auch einen ϵ -Tensor für zweidimensionale Vektoren. Dieser Tensor hat zwei Indizes, die dabei die Werte 1 und 2 annehmen:

$$\epsilon_{ij} = \begin{cases} 0 & \text{für } i = j \\ 1 & \text{für } i = 1 \text{ und } j = 2 \\ -1 & \text{für } i = 2 \text{ und } j = 1 \end{cases}$$

Für zweidimensionale Vektoren ist das Kreuzprodukt eine Zahl und es gilt:

$$\vec{x} \times \vec{y} = \sum_{i,j} \epsilon_{ij} x^i y^j = x^1 y^2 - x^2 y^1.$$

Wir können uns nun auch nochmals mithilfe des ϵ -Tensors überlegen, weshalb das Kreuzprodukt eines Vektors mit sich selber verschwindet. Es gilt:

$$(\vec{x} \times \vec{x})^i = \sum_{j,k} \epsilon^i_{jk} x^j x^k.$$

Der ϵ -Tensor ist „total antisymmetrisch“, d.h., die Vertauschung zweier Indizes erwirkt ein Minuszeichen:

$$\epsilon^i_{jk} = -\epsilon^i_{kj}.$$

Andererseits ist der Ausdruck $x^j x^k$ symmetrisch in den beiden Indizes:

$$x^j x^k = x^k x^j.$$

Die Verkürzung zweier Indizes (Summation über j und k) bei einem Produkt von einem symmetrischen und einem total antisymmetrischen Objekt ist aber immer null, da $x^j x^k - x^k x^j = 0$. Natürlich hätten wir dieses Ergebnis wesentlich einfacher ableiten können, da offensichtlich die rechte Seite von Gl. 2.33 verschwindet, wenn $\vec{x} = \vec{y}$. Das obige Argument tritt aber sehr häufig bei Berechnungen im „Indexkalkül“ auf, unabhängig vom speziellen ϵ -Tensor: Wird über zwei Indizes summiert und ist der eine Term in diesen beiden Indizes symmetrisch und der andere Term antisymmetrisch, so verschwindet das Produkt.

2.7.4 Das Spatprodukt

Das *Spatprodukt* ist eine Kombination aus Skalarprodukt und Vektorprodukt, bei dem drei Vektoren miteinander multipliziert werden. Für drei Vektoren $\vec{x}$, $\vec{y}$ und $\vec{z}$ gilt:

$$\text{Vol}(\vec{x}, \vec{y}, \vec{z}) = \vec{x} \cdot (\vec{y} \times \vec{z}), \quad (2.34)$$

wobei $\text{Vol}(\vec{x}, \vec{y}, \vec{z})$ das (orientierte) Volumen des von den drei Vektoren $\vec{x}$, $\vec{y}$ und $\vec{z}$ aufgespannten Parallelepipeds ist. Das orientierte Volumen ist positiv (gleich dem gewöhnlichen Volumen), wenn die drei Vektoren der „Rechte-Hand-Regel“ genügen, andernfalls ist das Volumen negativ (gleicher Betrag wie das gewöhnliche Volumen).

Diese Beziehung kann man sich aus folgender Überlegung verdeutlichen:

$$\vec{x} \cdot (\vec{y} \times \vec{z}) = |\vec{x}| |\vec{y} \times \vec{z}| \cos \gamma,$$

wobei γ der Winkel zwischen einem Vektor senkrecht zu der von $\vec{y}$ und $\vec{z}$ aufgespannten Ebene und dem Vektor $\vec{x}$ ist. Daher folgt:

- $|\vec{x}| \cos \gamma$ ist gleich der Komponente von $\vec{x}$ senkrecht zu der von $\vec{y}$ und $\vec{z}$ aufgespannten Fläche, also gleich der „Höhe“ von $\vec{x}$ über dieser Fläche,
- $|\vec{x} \times \vec{y}|$ ist gleich dem Betrag der von $\vec{x}$ und $\vec{y}$ aufgespannten Fläche.

Der Betrag der Grundfläche multipliziert mit der Höhe ergibt die Fläche des Parallelepipedes.

Aus der expliziten Darstellung,

$$\vec{x} \cdot (\vec{y} \times \vec{z}) = x^1 y^2 z^3 - x^1 y^3 z^2 + x^2 y^3 z^1 - x^2 y^1 z^3 + x^3 y^1 z^2 - x^3 y^2 z^1$$

kann man auch leicht die zyklische Symmetrie des Spatprodukts ableiten:

$$\vec{x} \cdot (\vec{y} \times \vec{z}) = \vec{z} \cdot (\vec{x} \times \vec{y}) = \vec{y} \cdot (\vec{z} \times \vec{x}).$$

In Indeschreibweise hat das Spatprodukt die Form:

$$\text{Vol}(\vec{x}, \vec{y}, \vec{z}) = \vec{x} \cdot (\vec{y} \times \vec{z}) = \delta_{li} x^l (\vec{y} \times \vec{z})^i = x_i \epsilon^i_{jk} y^j z^k = \epsilon_{ijk} x^i y^j z^k. \quad \text{Summenkonvention!}$$

A: Diese Formel kann man auch zum Ausgangspunkt für eine Definition des ϵ -Tensors machen, die nicht mehr vom kartesischen Skalarprodukt abhängt: Für einen d -dimensionalen Vektorraum V ist der ϵ -Tensor als eine total antisymmetrische multilineare Abbildung $\epsilon : V \times V \times \dots \times V \rightarrow \mathbf{R}$ (d -faches kartesisches Produkt) definiert, die den d Vektoren das von ihnen aufgespannte Volumen zuordnet.

2.7.5 Determinanten

In Komponentenschreibweise kann man d linear unabhängige Vektoren des $\mathbf{R}^d$ auch in Form einer $d \times d$ Matrix M schreiben:

$$M = (\vec{x}_1, \dots, \vec{x}_d) = \begin{pmatrix} x^1_1 & x^1_2 & \dots & x^1_d \\ x^2_1 & x^2_2 & \dots & x^2_d \\ \vdots & \vdots & \ddots & \vdots \\ x^n_1 & x^n_2 & \dots & x^n_d \end{pmatrix}.$$

(Entsprechend unserer Konvention nummeriert der untere Index die Vektoren, der obere Index die Komponenten der Vektoren.) Die Matrix M beschreibt eine lineare Abbildung, welche die Basiselemente $\vec{e}_i$ auf die Vektoren $\vec{x}_i$ abbildet. Das von den d Vektoren aufgespannte Volumen bezeichnet man auch als die *Determinante* [*determinant*] der Matrix M :

$$\text{Vol}(\vec{x}_1, \dots, \vec{x}_d) = \det M = \begin{vmatrix} x_1^1 & x_1^2 & \cdots & x_1^d \\ x_2^1 & x_2^2 & \cdots & x_2^d \\ \vdots & \vdots & \ddots & \vdots \\ x_d^1 & x_d^2 & \cdots & x_d^d \end{vmatrix}. \quad (2.35)$$

Wie wir gesehen haben, gilt in zwei bzw. drei Dimensionen:

$$\begin{vmatrix} x_1^1 & x_1^2 \\ x_2^1 & x_2^2 \end{vmatrix} = x_1^1 x_2^2 - x_2^1 x_1^2$$

$$\begin{vmatrix} x_1^1 & x_1^2 & x_1^3 \\ x_2^1 & x_2^2 & x_2^3 \\ x_3^1 & x_3^2 & x_3^3 \end{vmatrix} = x_1^1 x_2^2 x_3^3 - x_1^1 x_3^2 x_2^3 + x_1^2 x_3^2 x_1^3 - x_1^2 x_2^1 x_3^3 + x_1^3 x_2^1 x_1^2 - x_1^3 x_2^2 x_1^1.$$

A: Die allgemeine Regel zur Auswertung der Determinante kann man in folgender Formel zusammenfassen:

$$\det M = \sum_p (-1)^{|p|} x_{p_1}^1 x_{p_2}^2 x_{p_3}^3 \cdots x_{p_d}^d,$$

wobei über alle Permutationen p mit $p(1, 2, 3, \dots, d) = (p_1, p_2, p_3, \dots, p_d)$ der Indizes zu summieren ist. Eine Permutation, die durch eine gerade Anzahl von paarweisen Vertauschungen aus der normalen Reihenfolge $(1, 2, 3, \dots, d)$ hervorgeht, bezeichnet man als gerade (für sie gilt $(-1)^{|p|} = 1$), anderenfalls als ungerade ($(-1)^{|p|} = -1$). Der d -dimensionale ϵ -Tensor ist gerade durch obige Summe definiert:

$$\det M = \sum_{p_1, p_2, \dots, p_d} \epsilon_{p_1 p_2 \dots p_d} x_{p_1}^1 x_{p_2}^2 \cdots x_{p_d}^d,$$

bzw.

$$\epsilon_{p_1 p_2 \dots p_d} = \begin{cases} 0 & \text{wenn zwei Indizes gleich sind} \\ 1 & \text{wenn die Sequenz } p_1, p_2, \dots, p_d \text{ eine gerade Permutation ist} \\ -1 & \text{wenn die Permutation ungerade ist} \end{cases}.$$

Die Determinante gibt unter anderem an, ob die d Vektoren $\{\vec{x}_i\}$ linear unabhängig sind oder nicht. Das von linear unabhängigen Vektoren aufgespannte Volumen ist von null verschieden, wohingegen das Volumen zu d linear abhängigen Vektoren verschwindet. Die Determinante einer $d \times d$ Matrix verschwindet also genau dann nicht, wenn der Bildraum der Matrix wieder der gesamte Vektorraum ist (die Abbildung also bijektiv ist), andernfalls ist die Determinante null.

Kapitel 3

Lineare Abbildungen

In diesem Kapitel beschäftigen wir uns mit Abbildungen von einem Vektorraum V in einen Vektorraum W , wobei wir auch oft $W = V$ wählen. Diese Abbildungen sollen die besonderen Eigenschaften von Vektorräumen unverändert lassen, was uns zu so genannten *linearen Abbildungen* [*linear mappings*] führen wird. Diese besitzen eine Darstellung in Form von Matrizen [*matrices*]. Spezielle lineare Abbildungen wie das Skalarprodukt haben wir schon im letzten Kapitel kennengelernt.

3.1 Lineare Abbildungen

Definition: Eine Abbildung $A : V \rightarrow W$ heißt *linear*, wenn gilt:

$$A(\alpha \vec{x} + \beta \vec{y}) = \alpha A(\vec{x}) + \beta A(\vec{y}). \quad (3.1)$$

Für $\vec{x} \in V$ ist also $A(\vec{x}) \in W$. Die beiden Vektorräume V und W müssen dabei nicht dieselbe Dimension haben. Im Folgenden sei d_v die Dimension von V und d_w die Dimension von W .

Da sich jeder Vektor $\vec{x} \in V$ nach einer in V gewählten Basis $\{\vec{e}_i\}$ entwickeln lässt, und wegen der Linearität von A gilt

$$A(x^1 \vec{e}_1 + x^2 \vec{e}_2 + \dots) = x^1 A(\vec{e}_1) + x^2 A(\vec{e}_2) + \dots,$$

ist die Abbildung A bereits eindeutig gegeben, wenn die Bilder der Basisvektoren bekannt sind, also wenn $\{A(\vec{e}_i)\}$ bekannt ist. Ist andererseits $\{\vec{f}_i\}$ eine Basis von

W , so können wir jeden Vektor in W , insbesondere auch die Vektoren $\{A(\vec{e}_i)\}$, nach dieser Basis entwickeln:

$$\begin{aligned} A(\vec{e}_1) &= A^1_1 \vec{f}_1 + A^2_1 \vec{f}_2 + A^3_1 \vec{f}_3 + \dots \\ A(\vec{e}_2) &= A^1_2 \vec{f}_1 + A^2_2 \vec{f}_2 + A^3_2 \vec{f}_3 + \dots \\ A(\vec{e}_3) &= A^1_3 \vec{f}_1 + A^2_3 \vec{f}_2 + A^3_3 \vec{f}_3 + \dots \\ &\vdots \quad \quad \quad \vdots \end{aligned}$$

oder in einer Gleichung zusammengefasst:

$$A(\vec{e}_i) = \sum_j A^j_i \vec{f}_j. \quad (3.2)$$

(Auch hier hätten wir natürlich wieder die Summenkonvention verwenden können, doch in diesem Kapitel werden wir in den meisten Fällen die Summationen explizit schreiben.) Die Stellung der Indizes ergibt sich aus folgender Regel:

- Bezüglich eines *oberen* Index transformiert sich die Matrix wie die Komponenten eines Vektors (kontravariant). Ein oberer Index wird beispielsweise kontrahiert mit den Komponenten eines dualen Vektors oder mit den Basisvektoren eines Vektorraums.
- Entsprechend transformiert sich eine Matrix bezüglich eines *unten* stehenden Index wie die Basisvektoren eines Vektorraums (kovariant). Ein unten stehender Index wird beispielsweise mit den Komponenten eines Vektors kontrahiert.
- Der *rechts* stehende Index bezieht sich auf den Urbildraum der Abbildung.
- Der *links* stehende Index bezieht sich auf den Bildraum der Abbildung.

Die Zahlen A^j_i bezeichnet man als die *Matrixelemente* [*matrix elements*] der Abbildung A . Man beachte, dass der Index j dabei von 1 bis d_w (der Dimension von W) läuft, während sich der Index i von 1 bis d_v (der Dimension von V) erstreckt.

A: Wenn eine Unterscheidung zwischen oben und unten stehenden Indizes nicht notwendig ist (oder diese Unterscheidung aus welchen Gründen auch immer nicht getroffen wird), schreibt man meist A_{ji} für diese Matrixelemente.

Ist $\epsilon^j \in W^*$ die kanonische Basis des Dualraums von W , definiert durch die Bedingung

$$\epsilon^j(\vec{f}_i) = \delta_i^j ,$$

so können wir die Matrixelemente A^j_i auch in der folgenden Form schreiben:

$$A^j_i = \epsilon^j(A(\vec{e}_i)) .$$

Wir können uns nun überlegen, wie die Komponenten eines Vektors unter der Abbildung A transformiert werden. Sei

$$\vec{x} = \sum_i x^i \vec{e}_i ,$$

dann gilt:

$$A(\vec{x}) = \sum_i x^i A(\vec{e}_i) .$$

Umgekehrt ist

$$\epsilon^j(A(\vec{x})) = y^j$$

die j -te Komponente des Bildvektors. Also folgt:

$$y^j = \epsilon^j(A(\vec{x})) = \sum_i x^i \epsilon^j(A(\vec{e}_i)) = \sum_i A^j_i x^i . \quad (3.3)$$

Nun wird auch die Bedeutung des ko- bzw. kontravarianten Transformationsverhaltens deutlich. Die Komponenten eines Vektors transformieren sich wie in Gl. 3.3, die Basisvektoren wie in Gl. 3.2.

Oft schreibt man die Elemente A^j_i (oder vereinfacht auch A_{ji}) in Form einer Matrix auf:

$$A \simeq \begin{pmatrix} A^1_1 & A^1_2 & A^1_3 & \dots & A^1_n \\ A^2_1 & A^2_2 & A^2_3 & \dots & A^2_n \\ A^3_1 & A^3_2 & A^3_3 & \dots & A^3_n \\ \vdots & \vdots & \vdots & \dots & \vdots \\ A^m_1 & A^m_2 & A^m_3 & \dots & A^m_n \end{pmatrix} . \quad (3.4)$$

Die Anwendung dieser Matrix auf einen Spaltenvektor $\vec{x} \in V$ ergibt den Spalten-

vektor $\vec{y} \in W$:

$$\begin{pmatrix} y^1 \\ y^2 \\ y^3 \\ \vdots \\ y^m \end{pmatrix} = \begin{pmatrix} A^1_1 & A^1_2 & A^1_3 & \dots & A^1_n \\ A^2_1 & A^2_2 & A^2_3 & \dots & A^2_n \\ A^3_1 & A^3_2 & A^3_3 & \dots & A^3_n \\ \vdots & \vdots & \vdots & \dots & \vdots \\ A^m_1 & A^m_2 & A^m_3 & \dots & A^m_n \end{pmatrix} \begin{pmatrix} x^1 \\ x^2 \\ x^3 \\ \vdots \\ x^n \end{pmatrix}$$

$$= \begin{pmatrix} A^1_1 x^1 + A^1_2 x^2 + \dots + A^1_n x^n \\ A^2_1 x^1 + A^2_2 x^2 + \dots + A^2_n x^n \\ A^3_1 x^1 + A^3_2 x^2 + \dots + A^3_n x^n \\ \vdots \\ A^m_1 x^1 + A^m_2 x^2 + \dots + A^m_n x^n \end{pmatrix}$$

Die j -te Komponente y^j des Bildvektors erhält man, indem man die Elemente in der j -ten Zeile der Matrix mit den Elementen des Vektors multipliziert (Elemente in der 1. Spalte in der j -ten Zeile mit der ersten Komponente des Vektors, Element in der 2. Spalte der j -ten Zeile mit der zweiten Komponente des Vektors, etc.) und die Ergebnisse addiert.

Definition: Der *Kern* [*kernel*] einer linearen Abbildung $A : V \rightarrow W$ ist der lineare Unterraum von V für den gilt:

$$\text{Kern}_A = \{ \vec{x} \in V \mid A(\vec{x}) = 0 \}.$$

(Wegen der Linearität von A kann man sich leicht überzeugen, dass es sich bei Kern_A tatsächlich um einen linearen Unterraum (Untervektorraum) handelt. Ebenso kann man sich leicht überzeugen, dass auch das *Bild* [*image*] unter der Abbildung A ein linearer Unterraum von W ist. Die Dimension des Bildraumes von A bezeichnet man als den *Rang* [*rank*] der Abbildung A .

3.2 Basistransformationen

Wir haben bisher immer eine feste Basis in einem Vektorraum betrachtet und haben gesehen, dass sich jeder Vektor als lineare Kombination dieser Basis schreiben lässt:

$$\vec{x} = x^1 \vec{e}_1 + x^2 \vec{e}_2 + \dots = \sum_i x^i \vec{e}_i$$

Was passiert jedoch, wenn wir die Basis ändern, wenn wir also statt der Basis $\{\vec{e}_i\}$ eine andere Basis linear unabhängiger Vektoren $\{\vec{f}_i\}$ wählen?

Wir betrachten dazu eine lineare Abbildung $F : V \rightarrow V$. Diese Abbildung soll die bisherige Basis $\{\vec{e}_i\}$ in V auf eine neue Basis $\{\vec{f}_i\}$ in V abbilden:

$$\vec{f}_i = F(\vec{e}_i).$$

Zunächst ist klar, dass sich jeder Basisvektor der neuen Basis als Linearkombination der alten Basisvektoren schreiben lässt (an dieser Stelle schreiben wir die Summen nochmals explizit, obwohl sich natürlich auch hier die Summenkonvention anwenden lässt):

$$\begin{aligned}\vec{f}_1 &= F_1^1 \vec{e}_1 + F_1^2 \vec{e}_2 + F_1^3 \vec{e}_3 + \dots = \sum_i F_1^i \vec{e}_i \\ \vec{f}_2 &= F_2^1 \vec{e}_1 + F_2^2 \vec{e}_2 + F_2^3 \vec{e}_3 + \dots = \sum_i F_2^i \vec{e}_i \\ \vec{f}_3 &= F_3^1 \vec{e}_1 + F_3^2 \vec{e}_2 + F_3^3 \vec{e}_3 + \dots = \sum_i F_3^i \vec{e}_i \\ &\vdots \\ &\vdots\end{aligned}$$

Diese Beziehungen lassen sich wieder in einer Gleichung zusammenfassen:

$$\vec{f}_j = \sum_i F_j^i \vec{e}_i. \quad (3.5)$$

Diese Beziehung gibt an, wie sich die neuen Basisvektoren $\{\vec{f}_i\}$ als Linearkombinationen der alten Basisvektoren $\{\vec{e}_i\}$ ausdrücken lassen. Die Koeffizienten dieser Linearkombinationen sind die Koeffizienten $\{F_j^i\}$ der Abbildungsmatrix von F .

Wir untersuchen nun, wie sich die Komponenten eines Vektors bei einem Wechsel der Basis verhalten. Man beachte, dass durch eine Basistransformation die Vektoren $\{\vec{x}\}$ nicht verändert werden sondern nur dessen Komponenten. Wir bezeichnen mit x_f^i die Komponenten des Vektors $\vec{x}$ in der Basis $\{\vec{f}_j\}$ und mit x_e^i die Komponenten *desselben* Vektors $\vec{x}$ in der Basis $\{\vec{e}_i\}$. Es gilt somit

$$\vec{x} = \sum_i x_e^i \vec{e}_i = \sum_j x_f^j \vec{f}_j. \quad (3.6)$$

(Dass wir einmal den Summationsindex mit i und einmal mit j bezeichnet haben, ist rein willkürlich, erleichtert aber die folgenden Rechnungen.) Wir können

nun mithilfe von Gl. 3.5 die Basisvektoren $\{\vec{f}_j\}$ durch die Basisvektoren $\{\vec{e}_i\}$ ausdrücken und erhalten:

$$\sum_i x_e^i \vec{e}_i = \sum_j x_f^j \left(\sum_i F_j^i \vec{e}_i \right) = \sum_i \left(\sum_j x_f^j F_j^i \right) \vec{e}_i. \quad (3.7)$$

Da die Basisvektoren $\{\vec{e}_i\}$ linear unabhängig sind und die Zerlegung eines Vektors nach einer solchen Basis eindeutig ist, müssen die Komponenten vor gleichen Basisvektoren auch gleich sein. Das bedeutet:

$$x_e^i = \sum_j F_j^i x_f^j. \quad (3.8)$$

Äquivalent können wir natürlich die Indizes auch umbenennen und schreiben:

$$x_e^j = \sum_i F_i^j x_f^i. \quad (3.9)$$

Diese Gleichung entspricht vollkommen Gl. 3.3, d.h., die Transformation der Komponenten bei einer Basistransformation entspricht einer Abbildung des Vektorraums auf sich selber. Der Unterschied zwischen den beiden Situationen besteht jedoch in einer Interpretation der Veränderungen: Im ersten Fall (Gl. 3.3) wird ein Vektor von V im Allgemeinen auf einen anderen Vektor in V abgebildet. Die Basis bleibt unverändert und man drückt die Komponenten des neuen Vektors durch die Komponenten des alten Vektors aus. In diesem Fall spricht man von einer *aktiven Transformation* [*active transformation*]. Im zweiten Fall (Gl. 3.9) bleiben die Vektoren des Vektorraums unverändert, werden aber durch unterschiedliche Basisvektoren beschrieben. Man spricht hier von einer *passiven Transformation* [*passive transformation*]. Die Unterscheidung zwischen aktiven und passiven Transformationen spielt eine große Rolle im Zusammenhang mit Symmetrien.

Interessant ist auch der Unterschied zwischen Gl. 3.9 und Gl. 3.5. Im einen Fall (Gl. 3.9) wird über den unteren Index der Elemente $\{F_i^j\}$ summiert, also über die Spalten der zugehörigen Matrix. Bei der Transformation der Basisvektoren (Gl. 3.5) wird jedoch über den oberen Index von $\{F_i^j\}$ summiert, also über die Zeilen der zugehörigen Matrix. Man kann nun alle Objekte mit einem Index danach klassifizieren, ob sie sich wie die Basisvektoren transformieren (nach Gl. 3.5), dann spricht man von *kovariantem* [*covariant*] Transformationsverhalten, oder wie die Komponenten der Vektoren (Gl. 3.9), was man als *kontravariantes* [*contravariant*] Transformationsverhalten bezeichnet. Man kann sich leicht überlegen, dass sich die Basisvektoren des dualen Vektorraums kontravariant, die Komponenten von dualen Vektoren aber kovariant transformieren.

3.3 Spezielle lineare Abbildungen

3.3.1 Multiplikation von Matrizen

Ist das Bild einer linearen Abbildung $A : V_1 \rightarrow V_2$ enthalten im Urbild einer zweiten linearen Abbildung $B : V_2 \rightarrow V_3$, so lassen sich die beiden Abbildungen zu einer linearen Abbildung $BA : V_1 \rightarrow V_3$ hintereinander schalten. Sei $\vec{y} \in V_2$ das Bild eines Vektors $\vec{x} \in V_1$, so gilt für die Komponenten:

$$y^k = \sum_j A^k_j x^j,$$

und wenn $\vec{z} \in V_3$ das Bild von $\vec{y}$ unter der Abbildung B ist so gilt entsprechend:

$$z^i = \sum_k B^i_k y^k.$$

Für die Komponenten von $\vec{z}$ ausgedrückt durch die Komponenten von $\vec{x}$ folgt daher:

$$z^i = \sum_k B^i_k \left(\sum_j A^k_j x^j \right) = \sum_j \left(\sum_k B^i_k A^k_j \right) x^j.$$

Durch Vergleich mit der allgemeinen Formel

$$z^i = \sum_j (BA)^i_j x^j$$

finden wir für die Matrixelemente der kombinierten Matrix BA :

$$(BA)^i_j = \sum_k B^i_k A^k_j, \quad (3.10)$$

bzw. in der bekannteren Notation, bei der beide Indizes unten geschrieben werden — erst der Zeilenindex, dann der Spaltenindex:

$$(BA)_{ij} = \sum_k B_{ik} A_{kj}. \quad (3.11)$$

Es werden also jeweils die Zeilen von B mit den Spalten von A verkürzt.

3.3.2 Die inverse Matrix

Zu einer bijektiven linearen Abbildung $A : V \rightarrow W$ gibt es entsprechend auch eine inverse lineare Abbildung A^{-1} , sodass gilt:

$$A^{-1}A = \mathbf{1}.$$

In Indexschreibweise erfüllen die Matrixelemente von A^{-1} nach Gl. 3.10 die Bedingung:

$$\sum_k (A^{-1})^i_k A^k_j = \delta^i_j. \quad (3.12)$$

Die Formeln, nach denen sich die Matrixelemente von A^{-1} direkt durch die Matrixelemente von A ausdrücken lassen, sind recht aufwendig, sofern man nicht vorab zusätzliche Strukturen definiert. Für den Fall von 2-dimensionalen Vektorräumen wollen wir die Lösung explizit angeben:

$$\begin{aligned} (A^{-1})^1_1 &= \frac{1}{\det A} A^2_2 \\ (A^{-1})^1_2 &= -\frac{1}{\det A} A^1_2 \\ (A^{-1})^2_1 &= -\frac{1}{\det A} A^2_1 \\ (A^{-1})^2_2 &= \frac{1}{\det A} A^1_1 \end{aligned} \quad (3.13)$$

mit der Determinante von A :

$$\det A = A^1_1 A^2_2 - A^1_2 A^2_1. \quad (3.14)$$

Durch explizites Nachrechnen kann man sich leicht von der Richtigkeit obiger Ausdrücke überzeugen.

3.3.3 Die transponierte Matrix

Definition: Sei $A : V \rightarrow V$ eine lineare Abbildung auf einem Vektorraum V , auf dem ein Skalarprodukt $g : V \times V \rightarrow \mathbf{R}$ gegeben ist. Dann ist die zu A *adjungierte* [*adjoint*] Abbildung $\hat{A}$ durch folgende Bedingung gegeben, die für alle $\vec{x}, \vec{y} \in V$ erfüllt sein soll:

$$g(\hat{A}(\vec{x}), \vec{y}) = g(\vec{x}, A(\vec{y})). \quad (3.15)$$

In Indexnotation geschrieben lautet die linke Seite dieser Gleichung:

$$\sum_{jk} g_{kj} (\hat{A}(\vec{x}))^k y^j = \sum_{ijk} g_{kj} \hat{A}_i^k x^i y^j.$$

Entsprechend folgt für die rechte Seite:

$$\sum_{ik} g_{ik} x^i (A(\vec{y}))^k = \sum_{ijk} g_{ik} A_j^k x^i y^j.$$

Da die beiden Seiten der Gleichung für beliebige $\vec{x}$ und $\vec{y}$ gleich sein sollen, folgt für die Matrixelemente der adjungierten Matrix:

$$\sum_k \hat{A}_i^k g_{kj} = \sum_k g_{ik} A_j^k,$$

oder, nachdem beide Seiten von rechts mit der inversen Matrix zum Skalarprodukt g^{jl} multipliziert werden:

$$\hat{A}_i^l = \sum_{kj} g_{ik} A_j^k g^{jl}.$$

Insbesondere erhält man für das kartesische Skalarprodukt $g_{ij} = \delta_{ij}$:

$$\hat{A}_i^l = A_i^l.$$

Ist ein kartesisches Skalarprodukt gegeben, so erhält man die Matrixelemente der adjungierten Matrix $\hat{A}$, indem man die Zeilen und Spalten der Matrix von A vertauscht. Man spricht in diesem Fall auch von der *transponierten* Matrix und schreibt $\hat{A} = A^T$.

Im Zusammenhang mit ko- und kontravariantem Transformationsverhalten hatten wir gesehen, dass sich die kovarianten Transformationsgesetze (Gl. 3.5) ebenso wie die kontravarianten Transformationsgesetze (Gl. 3.9) schreiben lassen, vorausgesetzt man vertauscht die Zeilen und Spalten der Transformationsmatrix. Man könnte also sagen, dass sich die Komponenten der dualen Vektoren bzw. die Basiselemente nach der transponierten Transformationsmatrix transformieren.

3.3.4 Orthogonale Matrizen

Eine wichtige Klasse von Matrizen sind die so genannten *orthogonalen* Matrizen [*orthogonal matrices*], die anschaulich Rotationen bzw. Kombinationen von Rotationen und Spiegelungen entsprechen. Unter einer Rotation verändern sich weder

die Längen von Vektoren noch die Winkel zwischen ihnen. Dies ist unabhängig von der Dimension n des Vektorraums. Daher gilt für Rotation, dass das kartesische Skalarprodukt $\vec{x} \cdot \vec{y}$ (das sich ja durch die Längen der beiden Vektoren und den Winkel zwischen ihnen ausdrücken lässt) nicht verändert. Die Menge aller orthogonalen Matrizen bildet eine Gruppe, die so genannte *orthogonale Gruppe* [*orthogonal group*] $O(n)$. Neben den reinen Rotationen erfüllen noch Punktspiegelungen sowie Spiegelungen an Ebenen diese Bedingung. Die Gruppe der reinen Rotationen bezeichnet man auch als *spezielle orthogonale Gruppe* [*special orthogonal group*] $SO(n)$. Während für orthogonale Matrizen immer $\det R = \pm 1$ gelten muss (wie wir gleich am Beispiel der 2-dimensionalen Drehungen zeigen werden), sind die Elemente der $SO(n)$ durch die Bedingung $\det R = +1$ definiert.

Seien $\vec{x}$ und $\vec{y}$ zwei Vektoren mit Skalarprodukt $\vec{x} \cdot \vec{y}$ und sei R eine Rotation. Die Aussage, dass sich das Skalarprodukt nicht ändern soll, bedeutet:

$$\vec{x} \cdot \vec{y} = R(\vec{x}) \cdot R(\vec{y}) . \quad (3.16)$$

Nach unseren Überlegungen des letzten Abschnitts folgt daraus:

$$\vec{x} \cdot \vec{y} = \vec{x} \cdot R^T R(\vec{y}) ,$$

wobei $R^T R$ die Matrix ist, die man durch Hintereinanderschaltung der Rotationsmatrix R und der zugehörigen transponierten Matrix R^T erhält.

Damit diese Gleichung für beliebige Vektoren $\vec{x}$ und $\vec{y}$ gelten kann, darf $R^T R$ den Vektor $\vec{y}$ nicht verändern, d.h., es muss gelten:

$$R^T R = \mathbf{1} , \quad (3.17)$$

oder gleichbedeutend:

$$R^T = R^{-1} . \quad (3.18)$$

Orthogonale Matrizen sind gerade durch diese Bedingung gekennzeichnet: Die transponierte Matrix ist gleich der inversen Matrix. Aus diesen Bedingungen kann man auch allgemein zeigen, dass $(\det R)^2 = 1$ gelten muss.

Als Beispiel betrachten wir Rotationen im $\mathbf{R}^2$. Gleichung 3.18 führt in zwei Dimensionen zunächst explizit auf die Bedingung (vgl. Gl. 3.13):

$$(\det R)^2 = (R^1_1 R^2_2 - R^1_2 R^2_1)^2 = 1 , \quad (3.19)$$

oder

$$\det R = \pm 1 .$$

Wir wählen die positive Lösung ($\det R = -1$ beschreibt neben Rotationen noch Spiegelungen). Damit gelangt man zu folgenden Bedingungen:

$$R^1_1 = R^2_2 \quad \text{und} \quad R^1_2 = -R^2_1,$$

die schließlich auf folgende Gleichung führen:

$$(R^1_1)^2 + (R^1_2)^2 = 1.$$

Die allgemeine Lösung dieser Gleichung lässt sich durch die Winkelfunktionen Sinus und Kosinus ausdrücken und durch einen Winkel φ parametrisieren. Man erhält:

$$R = \begin{pmatrix} \cos \varphi & \sin \varphi \\ -\sin \varphi & \cos \varphi \end{pmatrix}. \quad (3.20)$$

Wie man sich auch anhand einer einfachen Zeichnung und Rechnung verdeutlichen kann, beschreibt diese Matrix in zwei Dimensionen eine Rotation um den Nullpunkt mit Winkel φ .

3.3.5 Eigenwerte und Eigenvektoren

Wenn es zu einer linearen Abbildung $A : V \rightarrow V$ einen Vektor $\vec{x}$ gibt und eine Zahl λ , sodass gilt

$$A\vec{x} = \lambda\vec{x}, \quad (3.21)$$

so bezeichnet man $\vec{x}$ als einen *Eigenvektor* [*eigenvector*] von A und λ als den zugehörigen *Eigenwert* [*eigenvalue*]. Wir können obige Gleichung auch in folgender Form schreiben:

$$(A - \lambda \cdot \mathbf{1})\vec{x} = 0.$$

Zu der Matrix $A - \lambda \cdot \mathbf{1}$ gibt es also einen linearen Unterraum (Vielfache von $\vec{x}$), der auf die Null abgebildet wird. Damit kann das Bild der Abbildung $A - \lambda \cdot \mathbf{1}$ nicht mehr der gesamte Vektorraum sein, und das von dem Bild der Basisvektoren aufgespannte Volumen verschwindet. Mit anderen Worten:

$$\det(A - \lambda \cdot \mathbf{1}) = 0. \quad (3.22)$$

Diese Gleichung (eine polynomiale Gleichung n -ter Ordnung — n gleich der Dimension von V — für die Eigenwerte λ) bezeichnet man auch als *charakteristische Gleichung* [*characteristic equation*] für die Eigenwerte. Die linke Seite der Gleichung bezeichnet man als *charakteristisches Polynom* [*characteristic polynomial*].

Im Bereich komplexer Zahlen besitzt die charakteristische Gleichung immer n Lösungen, d.h., das charakteristische Polynom zu Gl. 3.22 lässt sich in folgende Form bringen:

$$\det(A - \lambda \cdot \mathbf{1}) = \prod_{i=1}^n (\lambda - \lambda_i).$$

Die Werte $\{\lambda_i\}$ bezeichnet man als die Lösungen der charakteristischen Gleichung und gleichzeitig als die Eigenwerte der Matrix A . Sind einige der λ_i gleich, so spricht man von *entarteten Eigenwerten* [*degenerate eigenvalues*].

Im Bereich reeller Zahlen kann es sein, dass man gar keine Lösungen zur charakteristischen Gleichung findet oder nur einen eingeschränkten Satz von Lösungen (weniger als n). Beispielsweise lautet die charakteristische Gleichung zur 2-dimensionalen Rotationsmatrix Gl. 3.20:

$$\lambda^2 - 2 \cos \varphi \lambda + 1 = 0.$$

Die Lösungen dieser Gleichung sind (zu komplexen Zahlen vgl. Kap. 4):

$$\lambda_{1/2} = \cos \varphi \pm i \sin \varphi = e^{\pm i \varphi}.$$

Außer für $\varphi = 0$ und $\varphi = 180^\circ$ gibt es keine reellen Lösungen.

3.3.6 Selbstadjungierte Matrizen

Ist auf einem Vektorraum V ein Skalarprodukt g definiert, so bezeichnet man eine lineare Abbildung $A : V \rightarrow V$ als *selbstadjungiert* [*self-adjoint*], wenn

$$g(\vec{x}, A\vec{y}) = g(A\vec{x}, \vec{y}),$$

bzw.

$$\hat{A} = A.$$

Speziell für das kartesische Skalarprodukt ist eine (reelle) Matrix also selbstadjungiert, wenn sie symmetrisch ist: $A^i_j = A_j^i$.

Ein wichtiger Satz aus der linearen Algebra lautet: Die Eigenwerte selbstadjungierter Matrizen ($A = A^T$ für reelle Vektorräume und kartesischem Skalarprodukt) sind reell und die Eigenvektoren zu verschiedenen Eigenwerten sind jeweils orthogonal. Der Beweis der Realität der Eigenwerte ist (fast) trivial, wenn

man komplexe Vektorräume mit hermiteschem Skalarprodukt (d.h., die Symmetriebedingung wird durch $g(\vec{x}, \vec{y}) = g(\vec{y}, \vec{x})^*$ ersetzt) betrachtet. Die Orthogonalität der Eigenvektoren lässt sich leicht zeigen. Seien $\vec{x}$ und $\vec{y}$ Eigenvektoren von A zu verschiedenen Eigenwerten λ_1 bzw. λ_2 . Dann gilt:

$$\vec{y} \cdot A(\vec{x}) = \lambda_1 \vec{y} \cdot \vec{x},$$

andererseits gilt wegen der Symmetrie von A auch:

$$\vec{y} \cdot A(\vec{x}) = A(\vec{y}) \cdot \vec{x} = \lambda_2 \vec{y} \cdot \vec{x}.$$

Es folgt also:

$$\lambda_1 \vec{y} \cdot \vec{x} = \lambda_2 \vec{y} \cdot \vec{x}.$$

Da wir $\lambda_1 \neq \lambda_2$ vorausgesetzt haben, muss $\vec{y} \cdot \vec{x} = 0$ gelten. Die beiden Eigenvektoren sind also orthogonal.

3.4 Skalare - Pseudoskalare - Vektoren - Pseudovektoren

Physikalische Größen lassen sich danach unterteilen, wie sie sich bei bestimmten Basistransformationen verhalten. Im Folgenden betrachten wir ganz konkret Rotationen der Basis sowie Punktspiegelungen der Basisvektoren ($\vec{e}_i \rightarrow -\vec{e}_i$). (Wir wählen grundsätzlich oben stehende Vektorindizes. Natürlich ließe sich noch unterscheiden, ob die Indizes oben oder unten stehen, d.h., ob sich die entsprechenden Komponenten bei Rotationen nach R oder nach $R^T = R^{-1}$ transformieren. Wir wollen die Sache an dieser Stelle jedoch nicht zu kompliziert machen.)

Ein *Skalar* [*scalar*] ist eine Größe, die sich unter Rotationen der Basis nicht verändert. Das einfachste Beispiel ist das Skalarprodukt zweier Vektoren, insbesondere auch der Betrag eines Vektors. Während sich die einzelnen Komponenten der Vektoren unter Rotationen und Punktspiegelung verändern,

$$\begin{aligned} x^i &\longrightarrow \tilde{x}^i = \sum_j R^i_j x^j && \text{Rotation} \\ x^i &\longrightarrow \tilde{x}^i = -x^i && \text{Punktspiegelung,} \end{aligned}$$

gilt dies nicht für das Skalarprodukt:

$$S = \vec{x} \cdot \vec{y} = \sum_i x_i y^i = \sum_i \tilde{x}_i \tilde{y}^i = S'.$$

Ein *Pseudoskalar* [*pseudo-scalar*] ist eine Größe, die sich unter Rotationen nicht verändert, unter einer Punktspiegelung aber das Vorzeichen wechselt. Ein Beispiel für einen Pseudoskalar ist das Spatprodukt dreier Vektoren:

$$P = (\vec{x} \times \vec{y}) \cdot \vec{z}.$$

Das Spatprodukt beschreibt das (orientierte) Volumen des Parallelepipeds, das von den drei Vektoren aufgespannt wird. Dieses Volumen ändert sich natürlich nicht, wenn die drei Vektoren gleichermaßen gedreht werden, aber unter einer Punktspiegelung wechseln alle drei Vektoren das Vorzeichen und somit auch das orientierte Volumen. Bei einer Punktspiegelung „geht die rechte Hand in die linke Hand“ über.

Ein *Vektor* V ist ein Triplet von Größen (V^1, V^2, V^3) , die sich bei einer Rotation bzw. Punktspiegelung ebenso transformieren wie die Komponenten eines Vektors:

$$V^i \longrightarrow \tilde{V}^i = \sum_j R^i_j V^j \qquad V^i \longrightarrow \tilde{V}^i = -V^i. \quad (3.23)$$

Das triviale Beispiel für einen Vektor sind natürlich die Komponenten eines Vektors selber.

Ein *Pseudovektor* V^P [*pseudo-vector*] ist ein Triplet von Größen, die sich bei einer Rotation genauso verhalten wie die Komponenten eines Vektors, die unter einer Punktspiegelung jedoch nicht das Vorzeichen wechseln:

$$(V^P)^i \longrightarrow (\tilde{V}^P)^i = \sum_j R^i_j (V^P)^j \qquad (V^P)^i \longrightarrow (\tilde{V}^P)^i = (V^P)^i. \quad (3.24)$$

Ein Beispiel für einen Pseudovektor ist das (3-dimensionale) Kreuzprodukt zweier Vektoren:

$$V^P = \vec{x} \times \vec{y}.$$

Das Ergebnis ist ein Vektor, der senkrecht auf $\vec{x}$ und $\vec{y}$ steht, und diese Eigenschaft bleibt unter einer Rotation erhalten. Der Betrag des Vektors entspricht der Fläche des von den beiden Vektoren aufgespannten Parallelogramms und auch diese Fläche bleibt unter Rotationen erhalten. Andererseits wechseln sowohl die Komponenten von $\vec{x}$ als auch die Komponenten von $\vec{y}$ das Vorzeichen unter einer Punktspiegelung, daher bleiben die Komponenten des Kreuzprodukts unter einer Punktspiegelung unverändert.

Auch Tensoren lassen sich danach klassifizieren, wie sich die Komponenten unter Rotationen und Punktspiegelungen verhalten.

Kapitel 4

Komplexe Zahlen

Komplexe Zahlen [*complex numbers*] spielen in der Physik eine wichtige Rolle. Viele Gleichungen lassen sich mithilfe komplexer Zahlen wesentlich einfacher lösen. Algebraische Gleichungen (Polynomgleichungen in den unbekannten Variablen) besitzen immer einen vollständigen Satz komplexer Lösungen, während es oftmals weniger oder gar keine reellen Lösungen gibt. Das folgende Kapitel kann natürlich nur einen elementaren Einstieg in die Theorie der komplexen Zahlen geben, insbesondere der wichtige Bereich der Theorie komplexer Funktionen und der Wegintegration über solche Funktionen kann kaum mehr als angedeutet werden.

4.1 Definition und einfache Relationen

Im Rahmen der reellen Zahlen besitzt die Gleichung

$$x^2 = -1$$

keine Lösungen. Wir definieren nun ein Element i als Lösung dieser Gleichung, erweitern die reellen Zahlen um dieses Element i und vervollständigen sie unter den Operationen der Addition und der Multiplikation.

Oft wird gefragt, wieso man eine Lösung obiger Gleichung einfach „definieren“ kann. Weshalb kann man nicht einfach akzeptieren, dass obige Gleichung keine Lösungen besitzt? Grundsätzlich ist dieser Einwand richtig, doch dieses Verfahren, eine Zahlenmenge um bestimmte Lösungen von Gleichungen zu erweitern, ist uns eigentlich schon vertraut. Auch Zahlen wie $\sqrt{2}$ sind streng genommen zunächst als Lösungen von Gleichungen definiert (hier $x^2 = 2$). Außerdem hat

sich die Erweiterung der reellen Zahlen zum komplexen Zahlenkörper als außerordentlich fruchtbares Konzept erwiesen. Das wichtigste Argument ist aber, dass eine solche Definition und die Erweiterung der reellen Zahlen um das Element i widerspruchsfrei möglich sind.

Fügt man das Element i zu den reellen Zahlen hinzu und bildet alle möglichen Kombinationen von Additionen und Multiplikationen, so erhält man die komplexen Zahlen. Sie lassen sich allgemein in der Form

$$\mathbf{C} = \{z = a + ib \mid a, b \in \mathbf{R}\}$$

schreiben. In der Darstellung $z = a + ib$ bezeichnet man a als den *Realteil* [*real part*] und b als den *Imaginärteil* [*imaginary part*] der komplexen Zahl.

Zusammen mit der Forderung $i^2 = -1$ erhält man damit folgende Rechenregeln für komplexe Zahlen:

$$z_1 + z_2 = (a_1 + ib_1) + (a_2 + ib_2) = (a_1 + a_2) + i(b_1 + b_2),$$

und

$$z_1 \cdot z_2 = (a_1 + ib_1) \cdot (a_2 + ib_2) = (a_1a_2 - b_1b_2) + i(a_1b_2 + a_2b_1).$$

Außerdem definieren wir noch die *komplexe Konjugation* [*complex conjugation*] als zusätzliche Struktur auf den komplexen Zahlen:

$$z^* \equiv (a + ib)^* = a - ib.$$

Der Betrag [*absolute value*] einer komplexen Zahl ist definiert als:

$$|z| = \sqrt{zz^*} = \sqrt{a^2 + b^2}.$$

Für die Division durch eine komplexe Zahl erhält man:

$$\frac{z_1}{z_2} = \frac{z_1 z_2^*}{z_2 z_2^*} = \frac{1}{a_2^2 + b_2^2} (a_1 a_2 + b_1 b_2 + i(a_2 b_1 - a_1 b_2)).$$

Insbesondere ist

$$\frac{1}{z} \equiv z^{-1} = \frac{1}{a + ib} = \frac{a - ib}{a^2 + b^2}.$$

Real- bzw. Imaginärteil einer komplexen Zahl erhält man nach folgenden Regeln:

$$\operatorname{Re}(z) = \frac{1}{2}(z + z^*) \quad \operatorname{Im}(z) = \frac{1}{2i}(z - z^*). \quad (4.1)$$

Für die komplexen Zahlen gelten dieselben Rechenregeln wie für reelle Zahlen, insbesondere auch

$$z_1 \cdot z_2 = 0 \implies z_1 = 0 \text{ oder } z_2 = 0$$

und

$$|z_1 z_2| = (a_1 a_2 - b_1 b_2)^2 + (a_1 b_2 + a_2 b_1)^2 = a_1^2 a_2^2 + b_1^2 b_2^2 + a_1^2 b_2^2 + a_2^2 b_1^2 = |z_1| |z_2|.$$

Lediglich eine ausgezeichnete Ordnungsrelation, $x > y$, wie sie für die reellen Zahlen existiert, gibt es bei komplexen Zahlen nicht mehr. (Man kann natürlich Ordnungsrelation auf den komplexen Zahlen definieren, allerdings sind diese immer mehr oder weniger willkürlich.)

Ein interessantes Ergebnis erhält man, wenn man das Produkt zweier komplexer Zahlen bildet, von denen eine jedoch noch komplex konjugiert wird:

$$z_1^* \cdot z_2 = (a_1 a_2 + b_1 b_2) + i(a_1 b_2 - a_2 b_1).$$

Der Realteil dieses Produkts entspricht dem Skalarprodukt von zweidimensionalen Vektoren mit Komponenten (a_i, b_i) , der Imaginärteil entspricht dem Kreuzprodukt zweier solcher Vektoren.

Funktionen von komplexen Zahlen sind meist ähnlich definiert wie Funktionen von reellen Zahlen. Ist beispielsweise eine Funktion durch ihre Potenzreihenentwicklung definierbar,

$$f(x) = \sum_{n=0} a_n x^n,$$

so ist $f(z)$ durch dieselbe Reihe definiert:

$$f(z) = \sum_{n=0} a_n z^n.$$

Der Konvergenzradius einer solchen Funktion ist nun durch die Werte von z gegeben, für die

$$|f(z)| = \sum_{n=0} |a_n z|^n$$

konvergiert. Natürlich muss es sich nicht unbedingt um eine Entwicklung um den Punkt $z = 0$ handeln. Auch folgende Reihe definiert eine komplexe Funktion:

$$f(z) = \sum_{n=0} a_n (z - z_0)^n.$$

Eine Funktion kann auch durch eine Relation definiert sein. Beispielsweise ist die Funktion

$$f(z) = \sqrt{z}$$

durch die Relation

$$f(z)^2 = z$$

gegeben. Zu einer gegebenen komplexen Zahl z sind also komplexe Zahlen $\sqrt{z}$ gesucht, deren Quadrat gleich z ist. Diese Funktion ist, ähnlich wie schon für reelle Zahlen, nicht eindeutig und die Theorie mehrdeutiger komplexer Funktionen wird rasch sehr aufwendig, sodass wir dieses Thema an dieser Stelle nicht weiter vertiefen. Auf die komplexe Exponentialfunktion $f(z) = \exp(z)$ kommen wir in einem späteren Abschnitt zu sprechen.

4.2 Analytische Funktionen

Definition: Eine komplexwertige Funktion $f(z)$ auf einem Bereich der komplexen Ebene heißt *homeomorph* [*homeomorphic*] an einem Punkt z , wenn die Ableitung $f'(z)$ der Funktion am Punkt z , definiert durch

$$f(z+h) = f(z) + f'(z) \cdot h + O(h^2),$$

eindeutig ist. Ist eine Funktion $f(z)$ in allen Punkten einer offenen Umgebung D in der komplexen Ebene homeomorph, so bezeichnet man die Funktion als homeomorph auf D .

Da sich die komplexen Zahlen $z = x + iy$ als Elemente der Zahlenebene $\mathbf{R}^2$ darstellen lassen, könnte man zunächst den Eindruck haben, als ob jede ableitbare zweikomponentige Funktion auf dem $\mathbf{R}^2$ eine homeomorphe Funktion definiert. Dies ist aber nicht der Fall. Die Eindeutigkeit der Ableitung, d.h. die Unabhängigkeit von der Richtung, aus der man sich in der komplexen Ebene dem Punkt z nähert, ist sehr weitreichend und schränkt die möglichen Funktionen auf der komplexen Ebene sehr ein. Um das zu sehen, betrachten wir eine beliebige zweikomponentige Funktion auf dem $\mathbf{R}^2$, also $(u(x, y), v(x, y)) \in \mathbf{R}^2$. Wir definieren die Ableitung nach der allgemeinen Vorschrift:

$$\begin{aligned} (u(x+h_1, y+h_2), v(x+h_1, y+h_2)) &= (u(x, y), v(x, y)) + \\ &\quad \left(\frac{\partial u}{\partial x} h_1 + \frac{\partial u}{\partial y} h_2, \frac{\partial v}{\partial x} h_1 + \frac{\partial v}{\partial y} h_2 \right) + O(\|h\|^2) \end{aligned}$$

Nun kommt die entscheidende Einschränkung: Der in h_i lineare Ausdruck soll sich als (komplexes) Produkt aus der komplexen Zahl (der Ableitung) $f'(z) = (a(x, y), b(x, y))$ mit der komplexen Zahl $h = (h_1, h_2)$ schreiben lassen. Damit muss gelten:

$$f'(z) \cdot h = (ah_1 - bh_2, ah_2 + bh_1).$$

Durch Vergleich mit Gl. 4.2 erkennen wir, dass dies nur möglich ist, wenn folgende Bedingungen erfüllt sind:

$$\frac{\partial u}{\partial x} = \frac{\partial v}{\partial y} \quad \text{und} \quad \frac{\partial u}{\partial y} = -\frac{\partial v}{\partial x}. \quad (4.3)$$

Man bezeichnet diese beiden Gleichungen als *Cauchy-Gleichungen*. Es handelt sich um Bedingungen an die Ableitungen einer 2-komponentigen Funktion auf dem $\mathbf{R}^2$. Wenn diese Bedingungen erfüllt sind, handelt es sich um eine homeomorphe Funktion.

Definition: Eine Funktion, die sich in einer Umgebung von jedem Punkt z_0 in einer offenen Umgebung D der komplexen Ebene als Potenzreihe schreiben lässt,

$$f(z) = \sum_{n=0}^{\infty} a_n (z - z_0)^n,$$

(wobei a_n beliebige komplexe Zahlen sein können, sodass obige Reihe für nichtverschwindende Werte von z konvergiert) bezeichnet man als *analytische Funktion* [*analytic function*] in D .

Es zeigt sich, dass jede in einem Gebiet D homeomorphe Funktion auch in D analytisch ist.

4.3 Die komplexe Exponentialfunktion

Die Exponentialfunktion ist für reelle Zahlen durch ihre Potenzreihenentwicklung definiert:

$$e^x = \sum_{n=0}^{\infty} \frac{1}{n!} x^n.$$

Dieselbe Potenzreihe definiert auch die Exponentialfunktion komplexer Argumente:

$$e^z = \sum_{n=0}^{\infty} \frac{1}{n!} z^n.$$

Für eine komplexe Zahl z ist $\exp(z)$ im allgemeinen wieder eine komplexe Zahl. Um ein Gespür für die Exponentialfunktion zu erhalten, betrachten wir sie zunächst für rein imaginäre Argumente:

$$e^{ix} = \sum_{n=0}^{\infty} \frac{1}{n!} i^n x^n.$$

Wegen

$$i^{4n} = 1, \quad i^{4n+1} = i, \quad i^{4n+2} = -1, \quad i^{4n+3} = -i$$

erhalten wir:

$$e^{ix} = \sum_{n=0}^{\infty} \frac{1}{(2n)!} (-1)^n x^{2n} + i \sum_{n=0}^{\infty} \frac{1}{(2n+1)!} (-1)^n x^{2n+1}$$

was wir in der bekannten *Euler'schen Gleichung* zusammenfassen können:

$$e^{ix} = \cos x + i \sin x. \quad (4.4)$$

Als Beispiel für eine Anwendung der Euler'schen Gleichung sowie den Umgang mit komplexen Zahlen wollen wir aus dieser Beziehung die Additionstheoreme für Winkelfunktionen ableiten. Es gilt:

$$e^{i(x+y)} = \cos(x+y) + i \sin(x+y).$$

Andererseits folgt:

$$\begin{aligned} e^{i(x+y)} &= e^{ix} e^{iy} = (\cos x + i \sin x) (\cos y + i \sin y) \\ &= (\cos x \cos y - \sin x \sin y) + i(\cos x \sin y + \sin x \cos y). \end{aligned}$$

Durch Vergleich der beiden Gleichungen (und unter Ausnutzung der Eindeutigkeit der Zerlegung einer komplexen Zahl in Real- und Imaginärteil) erhält man:

$$\begin{aligned} \cos(x+y) &= \cos x \cos y - \sin x \sin y \\ \sin(x+y) &= \cos x \sin y + \sin x \cos y. \end{aligned}$$

Mithilfe der Euler-Formel können wir für die Exponentialfunktion eines komplexen Arguments noch eine andere Darstellung angeben:

$$e^{x+iy} = e^x e^{iy} = e^x (\cos y + i \sin y).$$

Der Realteil des komplexen Arguments liefert somit einen reellen Faktor, welcher der üblichen Exponentialfunktion entspricht. Der Imaginärteil des Arguments kann mit der Euler-Formel als Summe des Kosinus- und Sinus-Anteils geschrieben werden.

Eine ähnliche Zerlegung gibt es für eine beliebige komplexe Zahl. Wenn wir uns $z = a + ib$ als Punkt in einer Ebene vorstellen, wobei die x -Achse die reellen Zahlen und die y -Achse die imaginären Zahlen darstellt, so kann man für den Punkt (a, b) auch schreiben:

$$(a, b) = (r \cos \varphi, r \sin \varphi) \quad \text{mit} \quad r = \sqrt{a^2 + b^2} \quad \text{und} \quad \tan \varphi = \frac{b}{a}.$$

Damit erhalten wir die so genannte *Polarzerlegung* [*polar decomposition*] einer komplexen Zahl:

$$z = a + ib = r \cos \varphi + ir \sin \varphi = r e^{i\varphi}.$$

(Man beachte, dass φ in diesem Fall als Radiant angegeben werden sollte, d.h., der Vollwinkel entspricht 2π .) Eine komplexe Zahl z wird in diesem Fall durch ihren Betrag (r) und den Winkel φ (als Radiant) zur reellen Achse ausgedrückt.

In dieser Form lässt sich aus einer komplexen Zahl auch leicht die Wurzel ziehen:

$$\sqrt{z} = \sqrt{r} \exp\left(\frac{1}{2}i\varphi\right).$$

Die mehrdeutigkeit der Wurzelfunktion kommt dadurch herein, dass man die Zahl z auch durch

$$z = r e^{i(\varphi+2\pi)}$$

hätte ausdrücken können und bei der Halbierung des Winkelarguments ein Winkel von π addiert würde. Da

$$e^{i\pi} = -1$$

entspricht dies einem Vorzeichen, wie bei der gewöhnlichen Wurzelfunktion.

4.4 Die Exponentialfunktion und 2-dimensionale Drehungen

Wir haben schon mehrfach angedeutet, dass man die komplexen Zahlen auch als Punkte in der so genannten komplexen Zahlenebene auffassen kann. Die x -Koordinate entspricht dem Realteil der komplexen Zahl und die y -Koordinate dem Imaginärteil.

Dreht man den Punkt (a, b) in der Ebene um einen Winkel φ , so erhält man den neuen Punkt (a', b') , indem man die Rotationsmatrix (Gl. 3.20) auf (a, b) anwendet:

$$\begin{pmatrix} a' \\ b' \end{pmatrix} = \begin{pmatrix} \cos \varphi & \sin \varphi \\ -\sin \varphi & \cos \varphi \end{pmatrix} \begin{pmatrix} a \\ b \end{pmatrix} = \begin{pmatrix} a \cos \varphi + b \sin \varphi \\ -a \sin \varphi + b \cos \varphi \end{pmatrix}.$$

Wie man sich leicht überzeugen kann, erhält man die um den Winkel φ „gedrehte“ komplexe Zahl z einfach durch Multiplikation mit

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

Der Zusammenhang zwischen diesen beiden Darstellungen wird deutlicher, wenn man die Drehmatrix in zwei Anteile zerlegt:

$$\begin{pmatrix} \cos \varphi & \sin \varphi \\ -\sin \varphi & \cos \varphi \end{pmatrix} = \cos \varphi \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} + \sin \varphi \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} = \cos \varphi \mathbf{1} + \sin \varphi \mathbf{I}.$$

Die Matrix

$$\mathbf{I} = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$$

erfüllt dieselben Multiplikationsregeln wie die „Zahl“ i :

$$\mathbf{I}^2 = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix} = \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix} = -\mathbf{1}. \quad (4.5)$$

Tatsächlich kann man die Menge der komplexen Zahlen auch durch die Menge der 2×2 -Matrizen

$$z = a + ib \sim \begin{pmatrix} a & b \\ -b & a \end{pmatrix}$$

ausdrücken. Die Multiplikation zweier solcher Matrizen führt wieder auf eine solche Matrix und entspricht genau den Multiplikationsregeln komplexer Zahlen:

$$\begin{pmatrix} a_1 & b_1 \\ -b_1 & a_1 \end{pmatrix} \begin{pmatrix} a_2 & b_2 \\ -b_2 & a_2 \end{pmatrix} = \begin{pmatrix} a_1 a_2 - b_1 b_2 & a_1 b_2 + a_2 b_1 \\ -(a_1 b_2 + a_2 b_1) & a_1 a_2 - b_1 b_2 \end{pmatrix}.$$

Der Zusammenhang zwischen der Exponentialfunktion und den Rotationsmatrizen wird ebenfalls deutlich, wenn wir die Exponentialfunktion der Matrix $\mathbf{I}$ betrachten. (Die Funktion einer Matrix ist wieder eine Matrix und definiert durch die Potenzreihenentwicklung der Funktion.)

Wegen der Relationen (4.5) gilt:

$$\begin{aligned}\exp(\varphi \mathbf{I}) &= \sum_{n=0}^{\infty} \frac{1}{n!} \varphi^n \mathbf{I}^n \\ &= \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} + \begin{pmatrix} 0 & \varphi \\ -\varphi & 0 \end{pmatrix} + \frac{1}{2} \begin{pmatrix} -\varphi^2 & 0 \\ 0 & -\varphi^2 \end{pmatrix} + \dots \\ &= \begin{pmatrix} \cos \varphi & \sin \varphi \\ -\sin \varphi & \cos \varphi \end{pmatrix}\end{aligned}$$

Die Exponentialfunktion der Matrix $\mathbf{I}$ (multipliziert mit einem Winkel φ) führt also direkt auf die 2-dimensionale Rotationsmatrix.

Übrigens entspricht die Matrix $\mathbf{I}$ der Matrixdarstellung des 2-dimensionalen ϵ -Tensors:

$$\begin{pmatrix} \epsilon_{11} & \epsilon_{12} \\ \epsilon_{21} & \epsilon_{22} \end{pmatrix} = \mathbf{I} = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}.$$

Damit wurde eigentlich nur die Spitze des Eisbergs an Beziehungen angesprochen, die zwischen den komplexen Zahlen, dem 2-dimensionalen ϵ -Tensor und den Rotationen in einer Ebene bestehen.

Kapitel 5

Kurven und Flächen

Im Folgenden betrachten wir den $\mathbf{R}^n$ als den Raum möglicher Lagen von Körpern und wir wollen beispielsweise die Bahnkurven bzw. Trajektorien von Körpern im $\mathbf{R}^n$ beschreiben. Obwohl wir zunächst mit dem kartesischen Skalarprodukt rechnen, werden wir die Stellung der Indizes entsprechend ihrer allgemeinen Bedeutung schreiben.

5.1 Darstellung von Wegen im $\mathbf{R}^n$

Einen *Weg* [*path*] oder eine *Kurve* können wir mathematisch dadurch beschreiben, dass wir diesen Weg durch eine reelle Zahl parametrisieren. Dabei ist es durchaus möglich, dass demselben Raumpunkt zwei (oder mehrere) Parameterwerte zugeordnet werden, wenn nämlich der Weg denselben Raumpunkt mehrmals durchläuft. Ein Weg γ ist somit eine Abbildung von einem Intervall der reellen Zahlen (das den möglichen Parameterwerten entspricht) in den Raum $\mathbf{R}^n$:

$$\gamma : I \in \mathbf{R} \longrightarrow \mathbf{R}^n \quad t \in I \mapsto \gamma(t) \simeq \vec{x}(t) \in \mathbf{R}^n .$$

Eine physikalische oder mathematische Interpretation des Parameters t besteht zunächst nicht, obwohl dieser Parameter in der Physik häufig die Rolle der Zeit spielt (d.h., wir geben an, an welchem Raumpunkt $\vec{x}(t)$ sich ein Körper zu einem bestimmten Zeitpunkt t befindet). In der Mathematik wählt man t oft als die so genannte Bogenlänge. Auf diese spezielle Parametrisierung kommen wir noch zu sprechen.

Da wir einen Punkt im $\mathbf{R}^n$ durch seine n Koordinaten (beispielsweise in einem kartesischen Koordinatensystem) charakterisieren können, erhalten wir folgende

Beschreibung für einen Weg (oder eine Bahnkurve) im $\mathbf{R}^n$:

$$t \in I \mapsto \vec{x}(t) = (x^1(t), x^2(t), \dots, x^n(t)). \quad (5.1)$$

Eine solche Darstellung einer Kurve bezeichnet man auch als *Parameterdarstellung*. Ist $I = [t_0, t_1]$, so beginnt der Weg bei $\vec{x}(t_0) = (x^1(t_0), x^2(t_0), \dots, x^n(t_0))$ und er endet bei $\vec{x}(t_1) = (x^1(t_1), x^2(t_1), \dots, x^n(t_1))$. Wir fordern, dass die Abbildungen

$$\gamma : t \mapsto x^i(t) \quad i = 1, \dots, n$$

stetig und, sofern notwendig, ausreichend oft differenzierbar sind.

Eine andere Parametrisierung desselben Weges entspricht anderen Funktionen $x^i(t')$. Dabei sollte sich t' als bijektive Funktion von t ausdrücken lassen. Genau genommen sollte man zwischen dem unparametrisierten Weg (d.h., der Punktmenge, aus denen der Weg besteht) und dem parametrisierten Weg (also der genauen Zuordnung zwischen Parameter und Wegpunkt) unterscheiden. Insbesondere spricht man in der Physik oft von der *Bahnkurve*, wenn man den durch die Zeit parametrisierten Weg eines Punktteilchens meint, und von der *Trajektorie*, wenn die unparametrisierte Punktmenge bezeichnet werden soll. Beispielsweise entspricht die Trajektorie der Erde um die Sonne einer Ellipse, unabhängig davon, wie diese Ellipse beschrieben wird, während man zur Angabe der Bahnkurve der Erde spezifizieren muss, wo sich die Erde zu einem bestimmten Zeitpunkt befindet.

Beispiel: Die Bahnkurve einer Helix im $\mathbf{R}^3$ entlang der z -Achse lässt sich durch folgende Abbildung beschreiben:

$$t \rightarrow (x^1(t), x^2(t), x^3(t)) = (r \cos \omega t, r \sin \omega t, \alpha t). \quad (5.2)$$

r , ω und α sind feste Parameter zur Charakterisierung der Helix. So ist r der Radius der Helix (in der Projektion auf die xy -Ebene), ω legt den Parameterbereich T fest, der einem Umlauf der Helix entspricht ($T = 2\pi/\omega$), und α charakterisiert (zusammen mit ω) die so genannte *Ganghöhe* h , um welche die Helix bei einem Umlauf zunimmt ($h = 2\pi\alpha/\omega$).

Zur Notation: Im Folgenden bezeichnen wir mit t meist den Parameter „Zeit“ und mit s die so genannte Bogenlänge (auf deren genaue Bedeutung wir noch zu sprechen kommen).

5.2 Der Tangentialvektor - Die Geschwindigkeit

Die Ableitung eines Weges $\vec{x}(t)$ nach dem Parameter t liefert einen Tangentialvektor an die Kurve am Punkte $\vec{x}(t)$. Handelt es sich bei dem Parameter t um die Zeit

und bezeichnet $\vec{x}(t)$ die Bahnkurve eines Teilchens, so ist dieser Tangentialvektor die Geschwindigkeit des Teilchens:

$$\vec{v}(t) = \frac{d\vec{x}(t)}{dt} = \dot{\vec{x}}(t) \quad (5.3)$$

bzw.

$$(v^1(t), v^2(t), \dots, v^n(t)) = \left(\frac{dx^1(t)}{dt}, \frac{dx^2(t)}{dt}, \dots, \frac{dx^n(t)}{dt} \right) = (\dot{x}^1(t), \dot{x}^2(t), \dots, \dot{x}^n(t)). \quad (5.4)$$

Die Geschwindigkeit ist zu jedem Zeitpunkt tangential zur räumlichen Trajektorie. Da die Geschwindigkeit im Laufe der Zeit verschiedene Werte annehmen kann, entspricht sie im mathematischen Sinne ebenfalls einer Bahnkurve bzw. Trajektorie, allerdings nicht im gewöhnlichen Ortsraum, sondern im Raum der Geschwindigkeiten.

Die Geschwindigkeit ist eine gerichtete Größe und daher ein Vektor. Es besteht allerdings ein grundlegender Unterschied in der Interpretation des Vektors $\vec{x}$ zur Beschreibung des Orts eines Teilchens und $\vec{v}$ zur Beschreibung seiner Geschwindigkeit. Es gibt zunächst keinen Grund, weshalb der physikalische Raum, in dem sich ein Teilchen bewegt, ein Vektorraum sein muss. Beispielsweise könnte sich das Teilchen auch auf der Oberfläche einer Kugel oder einer allgemeinen Mannigfaltigkeit bewegen (wie es beispielsweise in der allgemeinen Relativitätstheorie der Fall ist). Auf einer solchen Mannigfaltigkeit muss (zumindest lokal, d.h. in der Umgebung von jedem Punkt) ein Koordinatensystem definiert sein, das ganz allgemein einem Ausschnitt (einer offenen Menge) eines $\mathbf{R}^n$ entspricht. Doch das macht diese Mannigfaltigkeit noch nicht zu einem Vektorraum. So ist meist weder eine Addition von Raumpunkten noch die Multiplikation eines Raumpunkts mit einer reellen Zahl definiert. Der $\mathbf{R}^n$ dient in diesem Fall lediglich als Koordinatensystem zur Beschreibung der Raumpunkte.

Bei der Geschwindigkeit (bzw. bei einem Tangentialvektor) ist es jedoch anders. Geschwindigkeiten kann man addieren und auch mit reellen Zahlen multiplizieren. Geschwindigkeiten sind also Vektoren im eigentlichen Sinne. Bei einer allgemeinen Mannigfaltigkeit drückt sich das dadurch aus, dass die Geschwindigkeit nicht ein Element der Mannigfaltigkeit selber ist, sondern zum Tangentialraum an einem bestimmten Punkt der Mannigfaltigkeit gehört, der ein Vektorraum ist.

Der Betrag des Tangentialvektors ist

$$|\vec{v}(t)| = \sqrt{\vec{v}(t) \cdot \vec{v}(t)}.$$

Beispiel: Die Geschwindigkeit der Helix erhalten wir durch Ableitung von Gl. 5.2 nach der Zeit:

$$\vec{v}(t) = (-r\omega \sin \omega t, r\omega \cos \omega t, \alpha). \quad (5.5)$$

Für den Betrag der Geschwindigkeit finden wir:

$$|\vec{v}(t)| = \sqrt{r^2\omega^2 \sin^2 \omega t + r^2\omega^2 \cos^2 \omega t + \alpha^2} = \sqrt{r^2\omega^2 + \alpha^2}. \quad (5.6)$$

Der Betrag der Geschwindigkeit ist für die angegebene Parameterdarstellung (t entspricht der Zeit) der Helix zeitlich konstant. Die Trajektorie wird also mit konstanter Geschwindigkeit durchlaufen, allerdings ändert die Geschwindigkeit ständig ihre Richtung.

Bei einer Parametrisierungsänderung

$$t \longrightarrow t' = t'(t)$$

ändert sich der Betrag des Tangentialvektors, nicht aber seine Richtung:

$$\frac{d\vec{x}(t)}{dt} \longrightarrow \frac{d\vec{x}(t')}{dt'} = \frac{d\vec{x}(t)}{dt} \frac{dt}{dt'}.$$

Durch eine geeignete Parametrisierung kann man daher erreichen, dass der Tangentialvektor immer den Betrag 1 hat. Eine Parametrisierung $\vec{x}(s)$ einer Kurve, sodass

$$\left| \frac{d\vec{x}(s)}{ds} \right| = 1,$$

bezeichnet man als *Bogenlängenparametrisierung*. Der Grund für diese Bezeichnung wird deutlich, wenn wir Kurvenintegrale betrachten (vgl. Kap. 6.3).

5.3 Die Beschleunigung

Die Beschleunigung eines Teilchens entspricht der Ableitung der Geschwindigkeit nach der Zeit:

$$\vec{a}(t) = \frac{d\vec{v}(t)}{dt} = \dot{\vec{v}}(t) = \frac{d^2\vec{x}(t)}{dt^2} = \ddot{\vec{x}}(t) \quad (5.7)$$

bzw.

$$(a^1(t), a^2(t), \dots, a^n(t)) = \left(\frac{dv^1(t)}{dt}, \frac{dv^2(t)}{dt}, \dots, \frac{dv^n(t)}{dt} \right). \quad (5.8)$$

Auch die Beschleunigung ist Element eines Vektorraums.

Beispiel: Für die Beschleunigung der Helix-Trajektorie erhalten wir:

$$\vec{a}(t) = (-r\omega^2 \cos \omega t, -r\omega^2 \sin \omega t, 0). \quad (5.9)$$

Die Beschleunigung hat in diesem Fall keine 3-Komponente mehr, was auch verständlich ist, da die Bewegung des Teilchens in 3-Richtung einer konstanten Bewegung entspricht. Der Betrag der Beschleunigung,

$$|\vec{a}(t)| = \sqrt{r^2\omega^4 \cos^2 \omega t + r^2\omega^4 \sin^2 \omega t} = r\omega^2, \quad (5.10)$$

ist für die Helix zeitlich ebenfalls konstant.

Das Skalarprodukt von Geschwindigkeit und Beschleunigung ergibt für die Helix:

$$\vec{v}(t) \cdot \vec{a}(t) = r^2\omega^3 \sin \omega t \cos \omega t - r^2\omega^3 \cos \omega t \sin \omega t = 0. \quad (5.11)$$

Da weder die Geschwindigkeit noch die Beschleunigung dem Null-Vektor entsprechen bedeutet dieses Ergebnis, dass die Beschleunigung bei der Helix-Trajektorie zu jedem Zeitpunkt senkrecht zur Geschwindigkeit gerichtet ist. Wie das folgende Zwischenkapitel zeigt, ist dies eine unmittelbar Konsequenz aus der Tatsache, dass der Betrag der Geschwindigkeit konstant ist. In diesem Fall beschreibt die Beschleunigung nämlich nur eine Richtungsänderung der Geschwindigkeit, und diese ist immer senkrecht zur Geschwindigkeit selber gerichtet.

Wählt man zur Beschreibung der Kurve die Bogenlängenparametrisierung ($|\vec{x}(s)'| = 1$), so steht die Beschleunigung immer senkrecht auf dem Tangentenvektor. In diesem Fall bezeichnet man die Beschleunigung auch als *Krümmungsvektor* der Kurve im Punkte $\vec{x}(s)$. Der Betrag des Krümmungsvektors ist die *Krümmung*,

$$\kappa(s) = \left| \frac{d^2 \vec{x}(s)}{ds^2} \right|,$$

und der reziproke Wert

$$\rho(s) = \frac{1}{\kappa(s)}$$

heißt *Krümmungsradius* der Kurve im Punkte $\vec{x}(s)$. Für eine allgemeine Parametrisierung $\vec{x}(t)$ der Kurve gilt

$$\kappa(t) = \frac{\sqrt{(\vec{v}(t) \cdot \vec{v}(t))(\vec{a}(t) \cdot \vec{a}(t)) - (\vec{v}(t) \cdot \vec{a}(t))^2}}{(\vec{v}(t) \cdot \vec{v}(t))^{3/2}}.$$

5.4 Zwischenabschnitt: der Gradient

Wie wir gesehen haben, wird die Ableitung der Trajektorie eines Teilchens bzw. seiner Geschwindigkeit nach der Zeit komponentenweise durchgeführt und entspricht daher gewöhnlichen Ableitungen. Wie sieht es jedoch aus, wenn wir die Ableitung einer skalaren Funktion auf dem $\mathbf{R}^n$, nach der Zeit (bzw. nach dem Kurvenparameter) bestimmen wollen?

In einem konkreten Fall, wo die Parametrisierung der Bahnkurve explizit gegeben ist, können wir natürlich zunächst die entsprechende Funktion berechnen und dann nach der Zeit ableiten. Für viele Überlegungen ist es jedoch von Interesse, einen allgemeineren Ausdruck für die Ableitung einer Funktion auf dem $\mathbf{R}^n$ nach dem Parameter eines Weges zu haben.

Sei $U : \mathbf{R}^n \rightarrow \mathbf{R}$ eine Funktion, die jedem Vektor $\vec{x} \in \mathbf{R}^n$ eine reelle Zahl $U(\vec{x}) \in \mathbf{R}$ zuordnet. Eine solche Funktion bezeichnet man in der Physik auch häufig als *Feld*. Ein Beispiel ist der Betrag des Vektors zu einem Punkt:

$$U(\vec{x}) = |\vec{x}| = \sqrt{\vec{x} \cdot \vec{x}}. \quad (5.12)$$

Da wir den Vektor durch seine Komponenten in einer bestimmten Basis ausdrücken können, $\vec{x} = (x^1, x^2, \dots, x^n)$, wird auch die Funktion $U(\vec{x})$ zu einer Funktion der Komponenten $U(x^1, x^2, \dots, x^n)$, in obigem Beispiel:

$$U(x^1, x^2, \dots, x^n) = \sqrt{(x^1)^2 + (x^2)^2 + \dots + (x^n)^2}. \quad (5.13)$$

A: An dieser Stelle nehmen wir eine Identifikation vor, die in der Physik so selbstverständlich ist, dass sie meist nie erwähnt wird, und die doch häufig zu Komplikationen und Verständnisproblem führt. Streng genommen haben wir es in den beiden Gleichungen (5.12) und (5.13) mit zwei verschiedenen Funktionen zu tun. Gleichung 5.12 bezieht sich allgemein auf eine Funktion $U : M \rightarrow \mathbf{R}$, die jedem Punkt einer Mannigfaltigkeit M eine reelle Zahl zuordnet. Diese Funktion ist unabhängig von irgendeiner Koordinatenwahl oder Darstellung des Punktes. In dem konkreten Beispiel in Gl. 5.13 betrachten wir zunächst einen Punkt des Ortsraumes $\mathbf{R}^n$, der allerdings nicht als Vektorraum sondern als so genannter *affiner Raum* aufzufassen ist. Ein affiner Raum ist im Wesentlichen ein Vektorraum, bei dem allerdings der Nullpunkt nicht ausgezeichnet ist. Nicht die Punkte eines affinen Raums sind Vektoren, sondern die Differenzen zwischen zwei Punkten. Wir haben im $\mathbf{R}^n$ einen Nullpunkt ausgezeichnet und eine kartesische Basis gewählt, bezüglich der wir den Vektor Punkt $\vec{x}$ durch die Basis-komponenten $(x^1, x^2, \dots, x^n)$ ausdrücken. Dadurch wird $U : \mathbf{R}^n \rightarrow \mathbf{R}$ zu einer konkreten

Funktion von n Argumenten. Wie wir später noch sehen werden, können wir einen Vektor beispielsweise im $\mathbf{R}^3$ auch durch seinen Betrag und zwei Winkelvariable darstellen (in so genannten Polarkoordinaten) oder durch einen Radialvektor, eine Winkelvariable und die z -Komponente (in Zylinderkoordinaten). Die Funktion U soll immer noch den Betrag des Vektors angeben, aber die konkrete funktionale Abhängigkeit von den drei Variablen, durch die der Vektor beschrieben wird, ist je nach Koordinatensystem eine andere.

Bisher beschreibt U eine Funktion auf dem gesamten Vektorraum. Nun soll U jedoch als eine Funktion auf dem durch $\vec{x}(t)$ beschriebenen Weg in diesem Vektorraum aufgefaßt werden und wir erhalten:

$$U(\vec{x}(t)) = U(x^1(t), x^2(t), \dots, x^n(t)) \quad (5.14)$$

in obigem Beispiel also:

$$U(t) = \sqrt{x^1(t)^2 + x^2(t)^2 + \dots + x^n(t)^2}. \quad (5.15)$$

A: Auch hier findet wieder eine Identifikation statt. $U(t)$ ist eine Funktion des Parameters t , während $U(\vec{x})$ eine Funktion auf $\mathbf{R}^n$ ist. Strenggenommen entspricht $U(t)$ der Hintereinanderschaltung $U \circ \gamma$ der Abbildung $U : \mathbf{R}^n \rightarrow \mathbf{R}$ hinter der Abbildung für den Weg $\gamma : I \subset \mathbf{R} \rightarrow \mathbf{R}^n$.

Da das Argument t in allen Komponenten $x_i(t)$ steht, müssen wir für die Ableitung dieser Funktion nach dem Kurvenparameter t die Kettenregel bezüglich dieser drei Komponenten anwenden:

$$\frac{dU(\vec{x}(t))}{dt} = \frac{\partial U(\vec{x})}{\partial x^1} \frac{dx^1(t)}{dt} + \frac{\partial U(\vec{x})}{\partial x^2} \frac{dx^2(t)}{dt} + \dots + \frac{\partial U(\vec{x})}{\partial x^n} \frac{dx^n(t)}{dt}, \quad (5.16)$$

was für den Betrag des Vektor auf folgenden Ausdruck führt:

$$\frac{d|\vec{x}(t)|}{dt} = \frac{x_1(t)}{|\vec{x}(t)|} \frac{dx^1(t)}{dt} + \frac{x_2(t)}{|\vec{x}(t)|} \frac{dx^2(t)}{dt} + \dots + \frac{x_n(t)}{|\vec{x}(t)|} \frac{dx^n(t)}{dt}. \quad (5.17)$$

Die Notation $\partial/\partial x^i$ in Gl. 5.16 soll andeuten, dass U als Funktion von mehreren Argumenten aufgefasst wird, in dem jeweiligen Ausdruck aber nur nach dem Argument x^i abgeleitet wird. Außerdem erkennt man in Gl. 5.16 eine Verallgemeinerung der Kettenregel wieder, wenn man nämlich $U(t)$ als Hintereinanderschaltung der Abbildungen U und γ auffasst.

Wir sollten an dieser Stelle eine Bemerkung einfügen, weshalb an den Ableitungen von $|\vec{x}|$ nach den Komponenten von $\vec{x}$ die Indizes unten stehen:

$$\frac{\partial |\vec{x}|}{\partial x^i} = \frac{x_i}{|\vec{x}|}.$$

Der spezielle Grund im vorliegenden Fall ist, dass der Betrag eines Vektors über das Skalarprodukt des Vektors mit sich selber definiert ist:

$$|\vec{x}| = \sqrt{\sum_{ij} \delta_{ij} x^i x^j}.$$

Die Ableitung nach der Komponente x^i ergibt somit:

$$\frac{\partial |\vec{x}|}{\partial x^i} = \frac{\sum_j \delta_{ij} x^j}{|\vec{x}|} = \frac{x_i}{|\vec{x}|}.$$

Ganz allgemein lässt sich zeigen, dass sich die Ableitung einer Funktion nach den Komponenten eines Vektors wie die Komponenten eines dualen Vektors transformieren muss. Daher schreibt man oft auch:

$$\partial_i U = \frac{\partial U}{\partial x^i}.$$

Auf der rechten Seite von Gleichung 5.16 wird somit ein dualer Vektor auf einen gewöhnlichen Vektor angewandt. Der gewöhnliche Vektor ist der Geschwindigkeitsvektor,

$$\vec{v}(t) = \left(\frac{dx^1(t)}{dt}, \frac{dx^2(t)}{dt}, \dots, \frac{dx^n(t)}{dt} \right),$$

der duale Vektor, dessen Komponenten aus den partiellen Ableitungen der Funktion U nach den jeweiligen Ableitungen besteht, bezeichnet man als den *Gradienten* der Funktion U :

$$\text{grad } U(\vec{x}) \equiv \vec{\nabla} U(\vec{x}) = \left(\frac{\partial U(x^1, x^2, \dots, x^n)}{\partial x^1}, \frac{\partial U(x^1, x^2, \dots, x^n)}{\partial x^2}, \dots, \frac{\partial U(x^1, x^2, \dots, x^n)}{\partial x^n} \right). \quad (5.18)$$

Beide Notationen ($\text{grad } U$ und $\vec{\nabla} U$) sind in Gebrauch, insbesondere in neuen Texten setzt sich aber $\vec{\nabla} U$ immer mehr durch. In Komponentenschreibweise gilt also:

$$(\vec{\nabla} U)_i = \partial_i U = \frac{\partial U(\vec{x})}{\partial x^i}. \quad (5.19)$$

Der Gradient — die Ableitungen einer Funktion nach den Koordinaten eines Raumes — ist die erste einer ganzen Reihe von Ableitungsoperatoren, die auf Felder angewandt werden können. Von seiner Konstruktion her wird der Gradient in erster Linie auf skalare Felder angewandt, das Ergebnis ist dann ein (dualer) Vektor.

Für die Ableitung einer skalaren Funktion auf dem $\mathbf{R}^n$ nach dem Kurvenparameter eines Weges in diesem Raum erhalten wir also (Gl. 5.17):

$$\frac{dU(\vec{x}(t))}{dt} = \vec{\nabla}U(\vec{x}(t)) \cdot \frac{d\vec{x}(t)}{dt} = \vec{\nabla}U(\vec{x}(t)) \cdot \vec{v}(t). \quad (5.20)$$

Wir schreiben das Ergebnis hier als Skalarprodukt von einem Vektor ($\vec{\nabla}U$) mit dem Geschwindigkeitsvektor. Dabei haben wir den Gradienten (der eigentlich ein dualer Vektor ist und auf $\vec{v}$ angewandt wird) über das kartesische Skalarprodukt mit einem gewöhnlichen Vektor identifiziert, der nun mit dem Geschwindigkeitsvektor skalar *multipliziert* wird. Mit diesem Ergebnis gelangen wir zu einer anschaulichen Interpretation des Gradienten. Dazu müssen wir jedoch etwas ausholen.

Unter den *Äquipotenzialflächen* (bzw. *Äquipotenziallinien* in zwei Dimensionen) einer (skalaren) Funktion $U(\vec{x})$ versteht man Flächen (bzw. Linien), auf denen die Funktion $U(\vec{x})$ konstant ist, also ihren Wert nicht ändert. Allgemein handelt es sich dabei um eine Verallgemeinerung einer Fläche mit $n - 1$ Dimensionen. Sei $\vec{x}(t)$ nun eine Kurve auf einer solchen Äquipotenzialfläche, dann ist die Funktion $U(\vec{x}(t))$ als Funktion von t eine konstante Funktion, und die Ableitung dieser Funktion nach t ist daher Null. Andererseits ist nach Gl. (5.20) die Ableitung von $U(\vec{x}(t))$ nach t gleich dem Gradienten von $U(\vec{x})$ angewandt auf den Tangentialvektor $\vec{v}(t)$ an die Kurve. Unter Zuhilfenahme des Skalarprodukts schreiben wir dies als das Produkt von zwei Vektoren. Unter der Annahme, dass beide Vektoren von Null verschieden sind, verschwindet das Skalarprodukt dann und nur dann, wenn die beiden Vektoren senkrecht aufeinander stehen. Da dies für jede beliebige Kurve auf der Äquipotenzialfläche gilt, steht der Gradient von $U(\vec{x})$ also immer senkrecht auf allen Vektoren, die tangential zu den Äquipotenzialflächen sind: Der Gradient einer (skalaren) Funktion steht bei einem gegebenen Punkt $\vec{x}$ senkrecht auf der Tangentialebene an die Äquipotenzialfläche in diesem Punkt und zeigt in die Richtung, in welche die Funktion $U(\vec{x})$ (in linearer Näherung) am stärksten zunimmt. Der Betrag des Gradienten ist proportional zur Steigung dieser Zunahme.

Betrachten wir als Beispiel den Betrag eines Vektors. Der Gradient ist in

diesem Fall (vgl. 5.17):

$$\vec{\nabla}|\vec{x}| = \frac{\vec{x}}{|\vec{x}|} \quad (5.21)$$

bzw. in Komponentenschreibweise:

$$\partial_i|\vec{x}| = \frac{x_i}{|\vec{x}|}.$$

Auf der rechten Seite von Gl. 5.21 steht der Vektor $\vec{x}$ selber, wobei jedoch jede Komponente durch den Betrag des Vektors dividiert wird. Die Richtung des Gradienten ist daher parallel zum Vektor $\vec{x}$ (zeigt also, anschaulich gesprochen, radial nach außen), sein Betrag ist 1:

$$\left| \frac{\vec{x}}{|\vec{x}|} \right| = 1.$$

Die „Zunahme“ der Länge eines Vektors ist in radialer Richtung also an jedem Punkt dieselbe. Senkrecht zur radialen Richtung, also auf den Kugelflächen, die nun die Äquipotenzialflächen bilden, bleibt die Länge der Vektoren konstant.

Einen Vektor vom Betrag 1 bezeichnet man als *Einheitsvektor*. Zu jedem beliebigen Vektor (der nicht der Null-Vektor ist) erhält man den Einheitsvektor, dessen Betrag 1 ist und der in dieselbe Richtung wie der gegebene Vektor zeigt, indem man die Komponenten des Vektors durch den Betrag des Vektors dividiert.

Wir können uns nun auch überlegen, weshalb der Beschleunigungsvektor bei einer Kurve, die mit konstantem Geschwindigkeitsbetrag durchlaufen wird, immer senkrecht auf dem Geschwindigkeitsvektor steht. Es gilt nämlich:

$$\frac{d[\vec{v}(t)]^2}{dt} = \nabla(\vec{v}^2) \cdot \vec{a}(t) = 2\vec{v}(t) \cdot \vec{a}(t).$$

Da der Betrag der Geschwindigkeit konstant sein soll verschwindet die linke Seite und damit auch der Ausdruck auf der rechten Seite dieser Gleichung. $\vec{v}(t)$ steht somit senkrecht auf $\vec{a}(t)$.

5.5 Das mitgeführte Dreibein

Im Folgenden betrachten wir speziell Kurven im $\mathbf{R}^3$. Für eine solche Kurve, wie die Helix, kann man in jedem ihrer Punkte ein Koordinatensystem definieren, das beispielsweise ein Beobachter, der sich entlang der Kurve bewegt, als Koordinatensystem verwenden kann. Man denke als Beispiel an einen Beobachter auf der

Erde, der den Erdmittelpunkt als Ursprung seines Koordinatensystems verwendet (so genannte *geozentrische Koordinaten*). Als ausgezeichnete Richtungen bieten sich beispielsweise die Bewegungsrichtung der Erde (momentaner Geschwindigkeitsvektor), die Verbindungslinie zur Sonne, also der Radialvektor (bei einer Kreisbahn würde dieser senkrecht zum Geschwindigkeitsvektor stehen) und ein Vektor senkrecht auf diesen beiden Vektoren an.

Im Folgenden betrachten wir ein ausgezeichnetes Koordinatensystem, das für nahezu jede Kurve aus rein geometrischen Bedingungen konstruiert werden kann. Dieses Koordinatensystem bezeichnet man auch als *mitgeführtes Dreibein*.

Zwei besondere Vektoren haben wir für eine Kurve schon kennengelernt: den Geschwindigkeitsvektor (allgemeiner Tangentialvektor) $\vec{v}(t)$ und den Beschleunigungsvektor $\vec{a}(t)$. Beide Vektoren hängen von der gewählten Parametrisierung der Kurve ab. Der Geschwindigkeitsvektor ist jedoch immer tangential zu einer Kurve, daher ist seine Richtung parametrisierungsunabhängig. Dividieren wir den Geschwindigkeitsvektor $\vec{v}$ durch seinen Betrag, so erhalten wir einen Einheitsvektor

$$\vec{e}_v(t) = \frac{\vec{v}(t)}{|\vec{v}(t)|},$$

der an jedem Punkt tangential zur Kurve ist. Dieser Vektor ist der erste Vektor unseres Dreibeins. Notwendig (und hinreichend) für seine Existenz ist, dass der Geschwindigkeitsvektor nicht null ist. Dies lässt sich durch eine geeignete Parametrisierung der Kurve immer erreichen. In der Bogenlängenparametrisierung hat der Tangentialvektor bereits den Betrag 1 und entspricht somit dem ersten Vektor des Dreibeins.

Für den Beschleunigungsvektor $\vec{a}(t)$ hängen sowohl der Betrag als auch die Richtung von der Parametrisierung der Kurve ab. Wir können jedoch den Beschleunigungsvektor immer in einen Anteil parallel zur Tangente - $\vec{a}_{\parallel}$ - und einen Anteil senkrecht zur Tangente - $\vec{a}_{\perp}$ - aufspalten. Die Komponente parallel zur Tangente erhalten wir durch das Skalarprodukt mit $\vec{e}_v$ und es gilt:

$$\vec{a}_{\parallel} = (\vec{e}_v \cdot \vec{a}) \vec{e}_v = \frac{(\vec{v} \cdot \vec{a})}{|\vec{v}|^2} \vec{v}. \quad (5.22)$$

Diese Komponente beschreibt die Änderung des Betrags der Geschwindigkeit. Wäre die Geschwindigkeit konstant, so steht $\vec{a}$ senkrecht auf $\vec{v}$ und dieser Anteil verschwindet.

Den senkrechten Anteil der Beschleunigung erhalten wir, indem wir den par-

alleen Anteil von $\vec{a}$ subtrahieren:

$$\vec{a}_\perp = \vec{a} - \vec{a}_\parallel = \vec{a} - \frac{(\vec{v} \cdot \vec{a})}{|\vec{v}|^2} \vec{v} = \frac{1}{|\vec{v}|^2} \left((\vec{v} \cdot \vec{v}) \vec{a} - (\vec{v} \cdot \vec{a}) \vec{v} \right).$$

Dieser Anteil beschreibt das Maß der Richtungsänderung der Kurve. Als zweiten Vektor unseres Dreiecks wählen wir den Einheitsvektor, der in die Richtung von $\vec{a}_\perp$ zeigt:

$$\vec{e}_a(t) = \frac{\vec{a}(t)_\perp}{|\vec{a}(t)_\perp|}. \quad (5.23)$$

Diesen Vektor bezeichnet man auch als *Hauptnormalenvektor*. In der Bogenlängenparametrisierung ist

$$\vec{e}_a(s) = \frac{\vec{a}(s)}{|\vec{a}(s)|}.$$

Damit $\vec{e}_a(t)$ eindeutig gegeben ist, muss die Kurve in dem Punkt $\vec{x}(t)$ ihre Richtung ändern. Für eine Gerade (bzw. Ausschnitte einer Geraden) ist dieser Vektor nicht mehr eindeutig definierbar.

Die von den beiden Vektoren $\vec{e}_v(t)$ und $\vec{e}_a(t)$ (bzw. äquivalent $\vec{v}(t)$ und $\vec{a}(t)$) aufgespannte Ebene bezeichnet man als *Schmiegeebene* an die Kurve im Punkte $\vec{x}(t)$.

Als dritten Vektor des ausgezeichneten Dreiecks wählen wir einen Einheitsvektor, der senkrecht auf $\vec{v}(t)$ und senkrecht auf $\vec{a}(t)$ steht. Diesen erhalten wir aus dem Kreuzprodukt:

$$\vec{e}_n(t) = \vec{e}_v(t) \times \vec{e}_a(t) = \frac{\vec{v}(t) \times \vec{a}(t)}{|\vec{v}(t) \times \vec{a}(t)|}. \quad (5.24)$$

Man bezeichnet $\vec{e}_n(t)$ auch als *Binormalenvektor* der Kurve im Punkte $\vec{x}(t)$.

Die drei Vektoren $(\vec{e}_v(t), \vec{e}_a(t), \vec{e}_n(t))$ (Tangentenvektor, Hauptnormalenvektor, Binormalenvektor) bilden zu jedem Zeitpunkt t ein kartesisches Koordinatensystem (es handelt sich um drei Einheitsvektoren, die jeweils senkrecht aufeinander stehen), das durch die Richtung und die Krümmung der Kurve ausgezeichnet ist.

5.6 Flächen

5.6.1 Parameterdarstellung von Flächen

Eine Abbildung von einem Ausschnitt des $\mathbf{R}^2$ (beispielsweise einem Rechteck oder einer Kreisfläche) in den $\mathbf{R}^n$ beschreibt eine Fläche (ähnlich wie eine Kurve einer

Abbildung von einem Ausschnitt des $\mathbf{R}$ in den $\mathbf{R}^n$ entspricht). Sei $U \subset \mathbf{R}^2$ ein Parametergebiet, beschrieben durch Koordinaten (u, v) , so ist eine Fläche im $\mathbf{R}^n$ gegeben durch:

$$(u, v) \in U \subset \mathbf{R} \longrightarrow \vec{x}(u, v) \in \mathbf{R}^n, \quad (5.25)$$

bzw. ausgedrückt in Koordinaten:

$$(u, v) \longrightarrow (x^1(u, v), x^2(u, v), \dots, x^n(u, v)). \quad (5.26)$$

Diese Beschreibung einer Fläche bezeichnet man wiederum als die *Parameterdarstellung* der Fläche.

Als Beispiel betrachten wir die Kugeloberfläche im $\mathbf{R}^3$. Eine Kugel lässt sich durch die Bedingung

$$(x^1)^2 + (x^2)^2 + (x^3)^2 = R^2$$

charakterisieren. Diese Bedingung können wir nach jeder der Koordinaten auflösen, beispielsweise nach x^3 :

$$x^3 = \pm \sqrt{R^2 - (x^1)^2 - (x^2)^2}.$$

Wir wählen nun x^1 und x^2 als unabhängige Parameter, setzen also $x^1 = u$ und $x^2 = v$, und erhalten:

$$x^1(u, v) = u \quad x^2(u, v) = v \quad x^3(u, v) = \pm \sqrt{R^2 - u^2 - v^2}.$$

Das positive Vorzeichen beschreibt die „nördliche“ Halbkugel, das negative Vorzeichen die südliche. Die Parameter (u, v) sind durch die Bedingung $u^2 + v^2 = R^2$ eingeschränkt. Das Parametergebiet im $\mathbf{R}^2$ ist also eine Kreisscheibe.

Wir können eine Kugeloberfläche aber auch durch zwei Winkel charakterisieren. Dazu überzeugen wir uns zunächst, dass sich jeder Vektor $\vec{x} = (x^1, x^2, x^3)$ im $\mathbf{R}^3$ durch seinen Abstand r vom Nullpunkt, dem Winkel θ zur z -Achse und den Winkel φ zwischen der Projektion auf die xy -Ebene und der x -Achse beschreiben lässt:

$$(x^1, x^2, x^3) = (r \cos \varphi \sin \theta, r \sin \varphi \sin \theta, r \cos \theta).$$

Es handelt sich hierbei um eine andere Beschreibung der Punkte des $\mathbf{R}^3$. Solche Koordinatentransformationen werden uns in einem der nächsten Kapitel noch beschäftigen. Im vorliegenden Fall genügt es, dass die Kugeloberfläche durch konstantes r gekennzeichnet ist. Wir erhalten daher für die Beschreibung der Kugeloberfläche:

$$(\varphi, \theta) \rightarrow (x^1(\varphi, \theta), x^2(\varphi, \theta), x^3(\varphi, \theta)) = (r \cos \varphi \sin \theta, r \sin \varphi \sin \theta, r \cos \theta).$$

Für die Parameter gilt $\varphi \in [0, 2\pi]$ und $\theta \in [0, \pi]$. Das Parametergebiet ist in diesem Fall ein Rechteck.

Damit es sich allgemein um eine reguläre Fläche handelt, verlangen wir nicht nur, dass die Funktionen $x^i(u, v)$ stetig und hinreichend oft differenzierbar sind, sondern wir verlangen auch, dass die Parametrisierung der Fläche in folgendem Sinne nicht entartet ist: An jedem beliebigen Punkt $\vec{x}(u_0, v_0)$ auf der Fläche sollen sich die Kurven $\vec{x}(u, v_0)$ (aufgefasst als Kurve mit Kurvenparameter u) und $\vec{x}(u_0, v)$ (Kurvenparameter v) unter einem nicht-verschwindenden Winkel schneiden. Konkret bedeutet dies: Die beiden Tangentialvektoren

$$\vec{f}_u = \left. \frac{\partial \vec{x}(u, v_0)}{\partial u} \right|_{u=u_0} \quad \text{und} \quad \vec{f}_v = \left. \frac{\partial \vec{x}(u_0, v)}{\partial v} \right|_{v=v_0} \quad (5.27)$$

sollen jeweils von null verschieden und linear unabhängig sein.

Im $\mathbf{R}^3$ sind diese beiden Bedingungen genau dann erfüllt, wenn

$$\vec{f}_n(u, v) := \vec{f}_u \times \vec{f}_v = \frac{\partial \vec{x}(u, v)}{\partial u} \times \frac{\partial \vec{x}(u, v)}{\partial v} \quad (5.28)$$

für alle Werte $(u, v) \in U$ von null verschieden ist. $\vec{f}_n(u, v)$ ist ein Vektor, der senkrecht auf der Fläche am Punkte $\vec{x}(u, v)$ steht. Diesen Vektor bezeichnet man auch als *Normalvektor*.

Am Beispiel der Kugeloberfläche erkennt man, dass diese Bedingungen für alle (u, v) erfüllt sind, für die $u^2 + v^2 < R^2$. Wenn wir Punkte auf dem Äquator der Kugeloberfläche beschreiben wollen, sollten wir nicht nach x^3 auflösen, sondern eher nach x^1 oder x^2 . Auch bei der zweiten Parametrisierung der Kugeloberfläche gibt es „singuläre“ Stellen: $\theta = 0$ und $\theta = \pi$. Man spricht in solchen Fällen auch von *Koordinatensingularitäten*.

Da die Parametrisierung der Fläche nicht entartet sein soll, können wir den Normalvektor auch durch seinen Betrag dividieren und erhalten den auf 1 normierten Normalvektor:

$$\vec{n}(u, v) = \frac{\left(\frac{\partial \vec{x}(u, v)}{\partial u} \times \frac{\partial \vec{x}(u, v)}{\partial v} \right)}{\left| \frac{\partial \vec{x}(u, v)}{\partial u} \times \frac{\partial \vec{x}(u, v)}{\partial v} \right|}. \quad (5.29)$$

5.6.2 Charakterisierung durch Bedingungen

Oftmals ist es nicht leicht, eine geeignete Parameterdarstellung einer Fläche zu finden. Insbesondere ist es für viele Flächen überhaupt nicht möglich, sie durch

eine Abbildung der im letzten Abschnitt genannten Art überall und frei von Singularitäten oder Entartungen der Parametrisierung zu definieren. Bekannte Beispiele sind, wie wir gesehen haben, die Kugeloberfläche oder auch die Oberfläche eines Torus.

Statt einer Parametrisierung ist es manchmal einfacher, eine Fläche durch eine geeignete Bedingung zu beschreiben. Im $\mathbf{R}^3$ lässt sich eine solche Bedingung häufig in der Form

$$f(x^1, x^2, x^3) = 0$$

schreiben. Für die Kugeloberfläche im $\mathbf{R}^3$ gilt beispielsweise

$$f(x^1, x^2, x^3) = (x^1)^2 + (x^2)^2 + (x^3)^2 - R^2,$$

was auf die Gleichung

$$(x^1)^2 + (x^2)^2 + (x^3)^2 = R^2. \quad (5.30)$$

führt.

Die Bedingung $f(x^1, x^2, x^3) = 0$ lässt sich (zumindest für bestimmte Werte der Koordinaten) nach einer der Koordinaten auflösen - beispielsweise $x^3 = \hat{f}(x^1, x^2)$ - und nun lässt sich die Fläche in der Form $(u, v, \hat{f}(u, v))$ beschreiben. Häufig ist diese Beschreibung jedoch nicht für die ganze Fläche gültig, sondern nur für bestimmte Gebiete der Fläche, die so genannten *Karten*. Wenn es für jeden Punkt der Fläche eine geeignete Karte gibt, sodass die Fläche in der Umgebung dieses Punktes frei von Koordinatensingularitäten beschrieben werden kann, so bezeichnet man die Menge dieser Karten auch als *Atlas*. Aus der Beschreibung der Kugeloberfläche durch Gl. (5.30) lassen sich beispielsweise sechs besondere Karten finden, und jeder Punkt der Kugeloberfläche liegt in mindestens einer dieser Karten. Diese sechs Karten erhält man, indem man Gl. (5.30) nach jeweils einer der Koordinaten x^i ($i = 1, 2, 3$) auflöst und dann noch die beiden Vorzeichen der Wurzel berücksichtigt. Jede der Karten beschreibt eine Halbkugel.

Auch Kurven bzw. Wege lassen sich durch Bedingungen beschreiben. Beispielsweise erhält man im Allgemeinen eindimensionale Linien, wenn man zwei Flächen schneidet, also zwei Bedingungen an die Koordinaten stellt:

$$f_1(x^1, x^2, x^3) = 0 \quad f_2(x^1, x^2, x^3) = 0.$$

Zur Parameterdarstellung gelangt man in diesem Fall, indem man beide Bedingungen nach einem Parameter, beispielsweise nach $x^1 = t$ auflöst:

$$\gamma(t) = (t, x^2(t), x^3(t)).$$

Ein anderes Beispiel für Kurven, die durch Bedingungen charakterisiert und zunächst nicht durch eine Parametrisierung gegeben sind, sind die Kegelschmitte, die in Abschnitt 5.8 kurz behandelt werden.

5.7 Parameterdarstellung d -dimensionaler Räume im $\mathbf{R}^n$

Die Konzepte der vergangenen Abschnitte können wir folgendermaßen verallgemeinern: ein d -dimensionaler Raum, eingebettet im $\mathbf{R}^n$ (zunächst verlangen wir noch $d < n$, später betrachten wir auch den Fall $d = n$), besitzt folgende Parameterdarstellung:

$$(u^1, u^2, \dots, u^d) \rightarrow \vec{x}(u^1, u^2, \dots, u^d), \quad (5.31)$$

wobei jede Komponente x^i ($i = 1, \dots, n$) eine (stetige und genügend oft differenzierbare) Funktion der Parameter $\{u^\alpha\}$ ($\alpha = 1, \dots, d$) ist. Jeder Punkt p auf diesem d -dimensionalen Raum soll eindeutig durch seine Koordinaten $\{u^\alpha\}$ mit $\vec{x}(u^1, \dots, u^d) \simeq p$ bestimmt sein. Statt von einem d -dimensionalen Raum spricht man auch manchmal von einer d -dimensionalen Mannigfaltigkeit. (Mannigfaltigkeiten lassen sich auch allgemeiner definieren, ohne dass man auf die Einbettung der Mannigfaltigkeit in einen $\mathbf{R}^n$ Bezug nehmen muss. Das soll hier aber nicht geschehen.)

Zu jeder der Komponenten u^α können wir die Kurve $\gamma^\alpha \simeq \vec{x}(u^\alpha)$ betrachten, die man erhält, wenn man alle anderen Komponenten $\{u^\beta\}$, $\beta \neq \alpha$ festhält. Durch jeden Punkt des d -dimensionalen Raumes (festgelegt durch die Koordinaten $\{u^\alpha\}$, wobei dieser Punkt nicht zu einer Koordinatensingularität gehören soll) finden wir daher d ausgezeichnete Kurven, die sich in diesem Punkt schneiden. Zu jeder dieser Kurven können wir den Tangentialvektor berechnen:

$$\vec{f}_\alpha(u^1, \dots, u^d) = \frac{\partial \vec{x}(u^1, \dots, u^d)}{\partial u^\alpha}.$$

Dieser Tangentialvektor ist natürlich eine Funktion der Koordinaten des Punktes, bei dem der Tangentialvektor berechnet wird. Wenn der Punkt $p \simeq \vec{x}(u^1, \dots, u^d)$ nicht zu einer Koordinatensingularität gehört, sind die d Tangentialvektoren $\vec{f}_\alpha(p)$ alle von null verschieden und linear unabhängig. Sie spannen somit einen d -dimensionalen Vektorraum auf, den so genannten *Tangentialraum* an die Mannigfaltigkeit im Punkte p .

Man sollte an dieser Stelle beachten, dass der Tangentialraum im Allgemeinen kein linearer Teilraum des $\mathbf{R}^n$ im üblichen Sinne ist. Das wäre nur der Fall,

wenn der Punkt p gerade dem Nullpunkt des $\mathbf{R}^n$ entspricht. Man kann einen Tangentialraum im obigen Sinne als einen *affinen Teilraum* des $\mathbf{R}^n$ auffassen, bei dem der Nullpunkt entsprechend verschoben ist. Besser ist es jedoch, den Tangentialraum überhaupt nicht als Teilraum des $\mathbf{R}^n$ zu interpretieren, sondern einfach als einen d -dimensionalen Vektorraum, der sich für jeden Punkt p einer Mannigfaltigkeit definieren lässt.

5.8 Kegelschnitte

Gleichsam als Anwendung einiger der in diesem Kapitel eingeführten Konzepte wollen wir eine wichtige Klasse von Kurven betrachten, die man als *Kegelschnitte* bezeichnet. Diese Kurven erhält man als Schnittkurven aus der Mantelfläche eines Doppelkegels im $\mathbf{R}^3$ mit einer Ebene. Neben den Ellipsen und Hyperbeln erhält man als Grenzfälle noch Kreise und Parabeln (sowie streng genommen noch einzelne Punkte sowie sich schneidende Geraden). Alle diese Kurven können beispielsweise bei der Bewegung eines Körpers im Gravitationsfeld eines anderen Körpers (das so genannte *Kepler-Problem*) auftreten.

Ein Doppelkegel (der Einfachheit halber mit Steigung 1) lässt sich durch folgende Bedingung beschreiben:

$$x^2 + y^2 = z^2, \quad (5.32)$$

bzw. in der Parameterdarstellung:

$$x(u, v) = u \quad y(u, v) = v \quad z(u, v) = \pm \sqrt{u^2 + v^2}.$$

(Wir verwenden die Komponentenschreibweise (x, y, z) statt (x^1, x^2, x^3) , da wir in kartesischen Koordinaten ohnehin nicht zwischen oben- und untenstehenden Indizes unterscheiden müssen und die Indexschreibweise in konkreten Formeln manchmal verwirrend sein kann.) Eine vollkommen andere aber äquivalente Parameterdarstellung wäre die folgende ($r \geq 0$, $0 \leq \varphi < 2\pi$):

$$x(r, \varphi) = r \cos \varphi \quad y(r, \varphi) = r \sin \varphi \quad z(r, \varphi) = r.$$

Eine beliebige Ebene durch den Nullpunkt können wir durch folgende Gleichung beschreiben:

$$\vec{a} \cdot \vec{x} = 0 \quad \text{bzw.} \quad a_1 x + a_2 y + a_3 z = 0,$$

wobei $\vec{a}$ ein beliebiger (von Null verschiedener) dualer Vektor ist. Da wir ein kartesisches Skalarprodukt verwenden, können wir $\vec{a}$ auch mit einem gewöhnlichen Vektor identifizieren. Die durch obige Gleichung beschriebene Ebene besteht aus allen Vektoren, die auf $\vec{a}$ senkrecht stehen. Soll die Ebene nicht durch den Nullpunkt sondern durch den Punkt $\vec{b} = (b^1, b^2, b^3)$ gehen, folgt:

$$\vec{a} \cdot (\vec{x} - \vec{b}) = 0 \quad \text{bzw.} \quad \vec{a} \cdot \vec{x} = \vec{a} \cdot \vec{b}.$$

Als allgemeine Darstellung einer Ebene erhalten wir somit:

$$a_1x + a_2y + a_3z = b. \quad (5.33)$$

A: Diese Gleichung ist noch überbestimmt. Sie enthält vier freie Parameter, zur Charakterisierung einer allgemeinen Ebene sind jedoch nur drei Parameter notwendig. Für $b \neq 0$ können wir obige Gleichung noch durch b dividieren und finden, dass eine allgemeine Ebene sogar durch die Bedingung

$$a'_1x + a'_2y + a'_3z = 1$$

beschrieben wird. Ist $b = 0$, geht die Ebene also durch den Nullpunkt, können wir durch ein nicht-verschwindendes a_i dividieren, d.h., einen der Koeffizienten vor x, y oder z eins setzen.

Wir betrachten zunächst den Fall $a_3 \neq 0$ (das bedeutet, die Ebene ist nicht parallel zur 3-Achse); dann können wir für die Ebene auch schreiben

$$z = a_1x + a_2y + b.$$

Ohne Einschränkung der Allgemeinheit setzen wir noch $a_2 = 0$ und schreiben $a_1 = a$:

$$z = ax + b. \quad (5.34)$$

Die Ebene ist in diesem Fall parallel zur 2-Achse. Dies bedeutet keine Einschränkung, da das Problem ohnehin symmetrisch unter Drehungen um die 3-Achse ist.

Einsetzen in die Kegelgleichung (5.32) liefert:

$$x^2 + y^2 = (ax + b)^2, \quad (5.35)$$

bzw.

$$(1 - a^2)x^2 - 2abx + y^2 = b^2. \quad (5.36)$$

A: Wir haben immer noch zwei Bestimmungsgleichungen: Gleichung 5.34 erlaubt es uns, die Koordinate z durch x auszudrücken, und Gleichung 5.36 erlaubt es uns, y durch x auszudrücken. Wir erhalten also eine Kurve im $\mathbf{R}^3$, wobei wir x als Kurvenparameter verwenden. Die Parameterdarstellung $\vec{x}(t)$ dieser Kurve wäre somit:

$$x(t) = t \quad y(t) = \pm \sqrt{b^2 - (1 - a^2)t^2 + 2abt} \quad z(t) = at + b.$$

Im Folgenden betrachten wir nur die Projektionen dieser Kurve auf die 1-2-Ebene. An den qualitativen Eigenschaften der Lösungen ändert sich dadurch nichts.

Allgemein können wir folgende Fälle unterscheiden:

1. *Kreise:*

Der Spezialfall $a = 0$ (waagerechte Ebene) führt auf die Gleichung

$$x^2 + y^2 = b^2.$$

Dies ist die Kreisgleichung. Für $b = 0$ erhalten wir als Lösung nur den Punkt $(0, 0, 0)$.

2. *Ellipsen:*

Die Bedingung $a^2 < 1$ führt auf eine Gleichung der Form ($\alpha > 0$):

$$\alpha(x - \gamma)^2 + y^2 = \beta^2.$$

Dies ist die Gleichung einer Ellipse mit Zentrum $(\gamma, 0)$.

3. *Parabeln:*

Für $a^2 = 1$ erhalten wir:

$$y^2 = b^2 \pm 2bx.$$

Dies ist die Gleichung einer Parabel.

4. *Hyperbeln:* Für $a^2 > 1$ folgt eine Gleichung der Form ($\alpha > 0$):

$$-\alpha(x - \gamma)^2 + y^2 = \beta^2.$$

Dies entspricht einer Hyperbelgleichung.

Wir müssen uns nun noch überlegen, welche Gleichung wir für $a_3 = a_2 = 0$ erhalten. In diesem Fall lautet die Gleichung für die Ebene

$$ax = b$$

Einsetzen in die Kegelgleichung 5.32 liefert:

$$\frac{b^2}{a^2} + y^2 = z^2 \quad \text{oder} \quad z^2 - y^2 = \gamma^2.$$

Dies ist die Gleichung einer Hyperbel, die sich asymptotisch den Geraden $z = \pm y$ annähert. Für $b = 0$ erhalten wir diese beiden sich schneidenden Geraden als Lösung.

5.9 Koordinatensysteme

Wir haben schon mehrfach erwähnt, dass man die Punkte des $\mathbf{R}^2$ oder $\mathbf{R}^3$ nicht nur als Vektoren auffassen kann, die sich bezüglich einer Basis des Vektorraums in Komponenten zerlegen lassen, sondern dass man sie auch allgemeiner durch Koordinaten beschreiben kann, die mit der Eigenschaft des $\mathbf{R}^2$ bzw. $\mathbf{R}^3$ als Vektorraum nichts mehr zu tun haben. Beispielsweise lässt sich die Summe zweier Vektoren im Allgemeinen auch nicht mehr als einfache Summe dieser Koordinaten ausdrücken. Außerdem ist der *Parameterbereich* dieser Koordinaten (der Wertebereich, den diese Koordinaten zur eindeutigen Beschreibung von Punkten annehmen können) nicht immer gleich dem $\mathbf{R}^2$ bzw. $\mathbf{R}^3$, sondern nur gleich einer Teilmenge.

Strenggenommen fasst man den $\mathbf{R}^2$ bzw. $\mathbf{R}^3$ in diesem Fall nicht mehr als Vektorraum auf, sondern als so genannte Mannigfaltigkeit. Das macht sich oft auch daran bemerkbar, dass die allgemeinen Koordinaten für bestimmte Punkte nicht mehr eindeutig sind oder singulär werden. Wir betrachten zunächst einige bekannte Beispiele für Koordinaten im $\mathbf{R}^2$ und $\mathbf{R}^3$

5.9.1 Die Polarkoordinaten im $\mathbf{R}^2$

In zwei Dimensionen können wir einen Punkt durch seinen Abstand r vom Nullpunkt und den Winkel φ zwischen der Verbindungsline zum Nullpunkt und der x -Achse beschreiben. Das führt auf die Darstellung eines Punktes in so genannten *Polarkoordinaten*:

$$(x^1, x^2) \simeq (r \cos \varphi, r \sin \varphi). \quad (5.37)$$

Die Komponenten $\{x^i\}$ lassen sich somit durch r und φ ausdrücken:

$$x^1 = r \cos \varphi \quad x^2 = r \sin \varphi .$$

Diese Bedingungen sind umkehrbar, d.h., r und φ lassen sich durch x^1 und x^2 ausdrücken:

$$r = \sqrt{(x^1)^2 + (x^2)^2} \quad \varphi = \arctan \frac{x^2}{x^1} .$$

Der Nullpunkt $(0,0)$ ist ein singulärer Punkt der Polarkoordinaten, da in diesem Punkt φ nicht eindeutig durch x^i gegeben ist. Streng genommen gelten die Polarkoordinaten also nur für den $\mathbf{R}^2$ ohne den Punkt $\{(0,0)\}$.

Die Koordinaten (r, φ) nehmen folgende Werte an: $r \geq 0$ und $\varphi \in [0, 2\pi)$. Der Parameterbereich $U \in \mathbf{R}^2$ der Koordinaten ist also ein Streifen (Breite 2π) oberhalb der positiven reellen Achse.

5.9.2 Zylinderkoordinaten im $\mathbf{R}^3$

Übertragen wir die Polarkoordinaten des $\mathbf{R}^2$ auf den $\mathbf{R}^3$ ohne die z -Komponente mit einzubeziehen, so erhalten wir die *Zylinderkoordinaten*:

$$(x^1, x^2, x^3) \simeq (r \cos \varphi, r \sin \varphi, z) . \quad (5.38)$$

Die alten Koordinaten x^i lassen sich also durch die neuen Koordinaten ausdrücken:

$$x^1 = r \cos \varphi \quad x^2 = r \sin \varphi \quad x^3 = z ,$$

ebenso lassen sich die neuen Koordinaten durch die alten Koordinaten ausdrücken:

$$r = \sqrt{(x^1)^2 + (x^2)^2} \quad \varphi = \arctan \frac{x^2}{x^1} \quad z = x^3 .$$

Die Zylinderkoordinaten sind eindeutig sofern $r \neq 0$, also außer auf der z -Achse, wo der Winkel φ nicht festliegt.

Der Parameterbereich für r und φ ist gleich ihrem Bereich für Polarkoordinaten, z kann beliebige reelle Werte annehmen.

5.9.3 Kugelkoordinaten

Wie schon mehrfach erwähnt, lässt sich ein Punkt im $\mathbf{R}^3$ auch durch seinen Abstand r vom Ursprung sowie zwei Winkel charakterisieren. Der Winkel θ entspricht dem Winkel zwischen der Verbindungslinie des Punkts zum Ursprung und

der z -Achse, und der Winkel φ entspricht dem Winkel zwischen der Projektion dieser Verbindungslinie auf die xy -Ebene und der x -Achse. Dann gilt:

$$(x^1, x^2, x^3) \simeq (r \cos \varphi \sin \theta, r \sin \varphi \sin \theta, r \cos \theta). \quad (5.39)$$

Auch diese Beziehungen lassen sich umkehren:

$$r = \sqrt{(x^1)^2 + (x^2)^2 + (x^3)^2} \quad \varphi = \arctan \frac{x^2}{x^1} \quad \theta = \arctan \frac{\sqrt{(x^1)^2 + (x^2)^2}}{x^3}.$$

Beide Winkel sind für $r = 0$ nicht definiert. Außerdem ist φ nicht definiert für $\theta = 0$ und $\theta = \pi$.

Die Definitionsbereiche für die Parameter sind: $r \geq 0$, $\varphi \in [0, 2\pi)$ und $\theta \in [0, \pi]$.

5.9.4 Allgemeine Koordinatentransformationen

Eine Koordinatentransformation auf dem $\mathbf{R}^n$ ist eine Abbildung $U \subset \mathbf{R}^n \rightarrow \mathbf{R}^n$ derart, dass sich jeder Punkt $p \simeq \vec{x} \in \mathbf{R}^n$ durch die Koordinaten $(u^1, u^2, \dots, u^n)$ ausdrücken lässt. Es gilt somit:

$$\vec{x} \simeq (x^1, x^2, \dots, x^n) \simeq (x^1(u^1, u^2, \dots, u^n), x^2(u^1, u^1, \dots, u^n), \dots, x^n(u^1, u^2, \dots, u^n)).$$

Jede Komponente x^i ist also eine Funktion der Koordinaten $\{u^i\}$:

$$x^i = x^i(u^1, u^2, \dots, u^n).$$

Umgekehrt soll sich auch jede Koordinate u^i als Funktion der $\{x^i\}$ schreiben lassen:

$$u^i = u^i(x^1, x^2, \dots, x^n).$$

Eine Koordinatentransformation soll also bijektiv sein und daher eine eindeutige inverse Abbildung besitzen. Punkte, an denen dies nicht möglich ist, bezeichnet man als singuläre Punkte der Koordinatentransformation. In vielen Fällen wird man nur einen bestimmten Ausschnitt des $\mathbf{R}^n$ (beispielsweise den $\mathbf{R}^n$ ohne gewisse Punkte, Linien oder Flächen) durch die neuen Koordinaten eindeutig ausdrücken können.

In jedem Punkt $p \in \mathbf{R}^n$ schneiden sich n Kurven γ_i , die man erhält, indem man die Koordinaten $\{u^j\}$ ($j \neq i$) festhält und $\vec{x}(u^i)$ nur noch als Funktion des

einen Parameters u^i auffasst. Zu jeder dieser Kurven können wir im Punkte p den Tangentialvektor bilden:

$$\vec{f}_i(p) = \frac{\partial \vec{x}(u^1, \dots, u^n)}{\partial u^i}.$$

Sofern der Punkt p nicht zu einer Koordinatensingularität gehört, sind diese n Vektoren von null verschieden und linear unabhängig. Sie spannen also einen n -dimensionalen Vektorraum auf. Nach den Überlegungen der vorherigen Abschnitte handelt es sich bei diesem Vektorraum strenggenommen um den Tangentialraum im Punkte p . Es ist der „Raum der Geschwindigkeiten“, mit denen der Punkt p durchlaufen werden kann.

Der affine Raum $\mathbf{R}^n$, der die möglichen Lagen von Punkten p charakterisiert, und der Vektorraum $\mathbf{R}^n$, der an jedem Punkt p als Tangentialraum (Raum der Geschwindigkeiten) konstruiert werden kann, sind zunächst zwei verschiedene Räume. Auch wenn eine Identifikation oft vorgenommen wird, sollte man sich der Unterscheidung bewusst sein. Die Physik ist in diesem Fall übrigens „genauer“ als die Mathematik: Während die Koordinaten von Raumpunkten meist die Dimension einer Länge haben, haben Geschwindigkeiten die Dimension [Länge pro Zeit]. Die beiden Konzepte werden daher deutlicher unterschieden.

5.10 Das metrische Feld

Wir sind in den vergangenen Kapiteln mehrfach auf das Konzept des Tangentialraums gestoßen und haben dabei betont, dass man den Tangentialraum (beispielsweise an eine Fläche) nicht mit dem einbettenden Raum identifizieren sollte. Daher haben wir zunächst noch keine Möglichkeit, Winkel und Längen für die Vektoren in einem Tangentialraum zu bestimmen. Trotzdem haben wir schon den Betrag von Geschwindigkeitsvektoren berechnet oder auch (bei den Tangentialebenen an eine Fläche) das Kreuzprodukt gebildet. Implizit haben wir dabei von der Einbettung Gebrauch gemacht. Dieses Konzept soll nun verdeutlicht werden. Ziel dieses Abschnitts ist es, in den Tangentialräumen ein positiv definites Skalarprodukt zu definieren. Ein solches Skalarprodukt bezeichnet man auch manchmal als *Metrik*. Mithilfe der Metrik können wir die Längen von Vektoren berechnen bzw. die Winkel zwischen Vektoren bestimmen, wenn die Komponenten dieser Vektoren in Bezug auf die durch das Koordinatensystem definierten Basisvektoren eines Tangentialraumes angegeben werden. Da wir streng genommen an jedem Punkt des Raumes einen anderen Tangentialraum haben, wird diese Metrik nun vom Raumpunkt abhängen. Daher spricht man auch von einem *metrischen Feld*.

5.10.1 Allgemeine Definition

Wir haben gesehen, dass zwischen der Parameterdarstellung der Einbettung einer Kurve, Fläche, oder allgemeiner eines d -dimensionalen Raumes in einen $\mathbf{R}^n$ und einer Koordinatentransformation auf dem $\mathbf{R}^n$ kein wesentlicher Unterschied besteht. Daher betrachten wir im Folgenden Einbettungen der Form:

$$(u^1, \dots, u^d) \mapsto \vec{x}(u^1, \dots, u^d),$$

wobei wir neben $d < n$ auch den Fall $d = n$ zulassen.

Die Parameterdarstellung definiert in einem Punkt p , der durch die Koordinaten $(u^1, \dots, u^d)$ charakterisiert wird, d linear unabhängige Tangentialvektoren:

$$\vec{f}_\alpha(p) = \frac{\partial \vec{x}(u^1, \dots, u^d)}{\partial u^\alpha}.$$

Diese Vektoren sind im Allgemeinen weder Einheitsvektoren noch orthogonal. Wir können natürlich diese Vektoren noch durch ihren Betrag dividieren und auch zur Konstruktion eines Orthonormalsystems verwenden (ähnlich, wie wir dies beim begleitenden Dreibein schon gemacht haben). In den meisten Fällen ist es aber einfacher, direkt mit den Vektoren $\vec{f}_\alpha(p)$ als Basisvektoren des Tangentialraums im Punkte p zu rechnen. Wenn wir allerdings Vektoren in diesem Tangentialraum durch ihre Komponenten bezüglich der Basis $\{\vec{f}_\alpha(p)\}$ ausdrücken, müssen wir bei der Berechnung des Skalarprodukts (wie es durch den einbettenden $\mathbf{R}^n$ definiert ist) vorsichtig sein: Wir müssen ein Skalarprodukt $g_{\alpha\beta}(p)$ definieren, dass uns bezüglich dieser Komponenten dasselbe Ergebnis liefert, wie das kartesische Skalarprodukt im $\mathbf{R}^n$.

Seien $\vec{y}$ und $\vec{z}$ zwei Vektoren im Tangentialraum des Punktes p . Diese beiden Vektoren lassen sich bezüglich der Tangentialvektoren $\{\vec{f}_\alpha\}$ zerlegen:

$$\vec{y} = \sum_{\alpha} y^{\alpha} \vec{f}_{\alpha} \quad \vec{z} = \sum_{\beta} z^{\beta} \vec{f}_{\beta}.$$

Das Skalarprodukt zwischen diesen Vektoren im einbettenden Raum $\mathbf{R}^n$ ist somit:

$$\vec{y} \cdot \vec{z} = \left(\sum_{\alpha} y^{\alpha} \vec{f}_{\alpha}(p) \right) \cdot \left(\sum_{\beta} z^{\beta} \vec{f}_{\beta}(p) \right) = \sum_{\alpha, \beta} y^{\alpha} z^{\beta} \vec{f}_{\alpha} \cdot \vec{f}_{\beta} = \sum_{\alpha, \beta} g_{\alpha\beta} y^{\alpha} z^{\beta}.$$

Hierbei haben wir definiert:

$$g_{\alpha\beta} = \vec{f}_{\alpha} \cdot \vec{f}_{\beta} = \frac{\partial \vec{x}(u^1, \dots, u^d)}{\partial u^{\alpha}} \cdot \frac{\partial \vec{x}(u^1, \dots, u^d)}{\partial u^{\beta}} = \sum_{i,j=1}^n \delta_{ij} \frac{\partial x^i(u^1, \dots, u^d)}{\partial u^{\alpha}} \frac{\partial x^j(u^1, \dots, u^d)}{\partial u^{\beta}}. \quad (5.40)$$

Die Matrixelemente dieses Skalarprodukts bestehen also aus den Skalarprodukten (bezüglich des kartesischen Skalarprodukts im $\mathbf{R}^n$) der durch das Koordinatensystem definierten Tangentialvektoren $\vec{f}_\alpha$.

Da die Tangentialvektoren $\vec{f}_\alpha(p)$ vom Punkt p abhängen, an dem sie berechnet werden, hängen auch die Elemente des Skalarprodukts von diesem Punkt p ab: $g_{\alpha\beta}(p)$. Man spricht in diesem Fall vom *metrischen Feld*. Da das Skalarprodukt in jedem Punkt ein Tensor (2. Stufe) ist, spricht man auch vom *metrischen Tensor* bzw. *metrischen Feldtensor*.

5.10.2 Polar- und Kugelkoordinaten

Sei in der Ebene $p \in \mathbf{R}^2$ durch den Abstand r vom Nullpunkt und den Winkel φ zur x -Achse charakterisiert, also $p \simeq (r \cos \varphi, r \sin \varphi)$. Die Polarkoordinaten definieren durch den Punkte p zwei ausgezeichnete Kurven:

$$\gamma_1(\varphi) = (r \cos \varphi, r \sin \varphi) \quad (r \text{ fest}) \quad \gamma_2(r) = (r \cos \varphi, r \sin \varphi) \quad (\varphi \text{ fest}). \quad (5.41)$$

$\gamma_1(\varphi)$ beschreibt einen Kreis vom Radius r um den Nullpunkt, und $\gamma_2(r)$ beschreibt eine Gerade, die sich radial vom Nullpunkt nach unendlich erstreckt. Diese beiden Kurven haben in p jeweils die Tangentialvektoren:

$$\begin{aligned} \vec{f}_\varphi(p) &= \frac{\partial \vec{x}(r, \varphi)}{\partial \varphi} = (-r \sin \varphi, r \cos \varphi) \\ \vec{f}_r(p) &= \frac{\partial \vec{x}(r)}{\partial r} = (\cos \varphi, \sin \varphi). \end{aligned}$$

Die Notation $\vec{f}_\varphi(p)$ bzw. $\vec{f}_r(p)$ soll andeuten, dass diese beiden Vektoren im Allgemeinen von dem Punkt p , ausgedrückt durch seine Koordinaten r und φ , abhängen.

Wir können uns leicht davon überzeugen, dass diese beiden Vektoren orthogonal sind. Normieren wir die beiden Vektoren noch auf die Länge eins, so erhalten wir ein *Orthonormalsystem* von Basisvektoren für den Tangentialraum im Punkte p :

$$\vec{e}_\varphi(p) = (-\sin \varphi, \cos \varphi) \quad \vec{e}_r(p) = (\cos \varphi, \sin \varphi).$$

Wollen wir jedoch Vektoren im Tangentialraum durch die Komponenten bezüglich $\vec{f}_r$ und $\vec{f}_\varphi$ ausdrücken, so müssen wir das verallgemeinerte Skalarprodukt verwenden:

$$g(p) \simeq \begin{pmatrix} g_{\varphi\varphi}(p) & g_{r\varphi}(p) \\ g_{\varphi r}(p) & g_{rr}(p) \end{pmatrix} = \begin{pmatrix} r^2 & 0 \\ 0 & 1 \end{pmatrix}. \quad (5.42)$$

Entsprechend finden wir für einen Punkt $p \in \mathbf{R}^3$ in Kugelkoordinaten:

$$\begin{aligned}\vec{f}_r(p) &= \frac{\partial \vec{x}(r, \varphi, \theta)}{\partial r} = (\cos \varphi \sin \theta, \sin \varphi \sin \theta, \cos \theta) \\ \vec{f}_\varphi(p) &= \frac{\partial \vec{x}(r, \varphi, \theta)}{\partial \varphi} = (-r \sin \varphi \sin \theta, r \cos \varphi \sin \theta, 0) \\ \vec{f}_\theta(p) &= \frac{\partial \vec{x}(r, \varphi, \theta)}{\partial \theta} = (r \cos \varphi \cos \theta, r \sin \varphi \cos \theta, -r \sin \theta).\end{aligned}$$

Auch in Kugelkoordinaten sind die drei Tangentialvektoren zu den Parametern r , φ und θ orthogonal. Für die Komponenten des metrischen Tensors erhalten wir:

$$g = \begin{pmatrix} g_{rr} & g_{r\varphi} & g_{r\theta} \\ g_{\varphi r} & g_{\varphi\varphi} & g_{\varphi\theta} \\ g_{\theta r} & g_{\theta\varphi} & g_{\theta\theta} \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & r^2 \sin^2 \theta & 0 \\ 0 & 0 & r^2 \end{pmatrix} \quad (5.43)$$

5.10.3 Die Kugeloberfläche

Die bisherigen Beispiele bezogen sich auf Koordinatentransformationen, nun wollen wir mit der Kugeloberfläche noch ein Beispiel für die Einbettung einer Fläche in den $\mathbf{R}^3$ betrachten. Eine mögliche Beschreibung der Kugeloberfläche folgt aus den Kugelkoordinaten, wobei wir nun den Radius R nicht mehr als Koordinate sondern als festen Parameter interpretieren. Die Koordinaten der Kugeloberfläche sind somit (φ, θ) . Die beiden Basisvektoren für den Tangentialraum an einem Punkt auf der Kugeloberfläche (ausgedrückt durch seine Koordinaten (φ, θ)) sind daher:

$$\begin{aligned}\vec{f}_\varphi(p) &= \frac{\partial \vec{x}(r, \varphi, \theta)}{\partial \varphi} = (-R \sin \varphi \sin \theta, R \cos \varphi \sin \theta, 0) \\ \vec{f}_\theta(p) &= \frac{\partial \vec{x}(r, \varphi, \theta)}{\partial \theta} = (R \cos \varphi \cos \theta, R \sin \varphi \cos \theta, -R \sin \theta).\end{aligned}$$

Die Komponenten des metrischen Tensors bezüglich dieser Basisvektoren des Tangentialraums lauten:

$$g = \begin{pmatrix} R^2 \sin^2 \theta & 0 \\ 0 & R^2 \end{pmatrix}. \quad (5.44)$$

Wir können die Kugeloberfläche auch durch die Koordinaten beschreiben, die wir durch Auflösen von der Bedingungsgleichung

$$(x^1)^2 + (x^2)^2 + (x^3)^2 = R^2$$

nach einer Komponente erhalten. Wir wählen die Parameterdarstellung für die „nördliche“ Hemisphäre:

$$(u, v) \rightarrow (u, v, \sqrt{R^2 - u^2 - v^2}).$$

Die beiden Tangentialvektoren an den Punkt p (gekennzeichnet durch die Koordinaten (u, v)) sind:

$$\begin{aligned}\vec{f}_u(p) &= \frac{\partial \vec{x}(u, v)}{\partial u} = \left(1, 0, -\frac{u}{\sqrt{R^2 - u^2 - v^2}}\right) \\ \vec{f}_v(p) &= \frac{\partial \vec{x}(u, v)}{\partial v} = \left(0, 1, -\frac{v}{\sqrt{R^2 - u^2 - v^2}}\right).\end{aligned}$$

Nun sind die beiden durch das Koordinatensystem definierten Basisvektoren im Allgemeinen weder Einheitsvektoren noch orthogonal und wir erhalten für die Metrik:

$$g = \begin{pmatrix} g_{uu} & g_{uv} \\ g_{vu} & g_{vv} \end{pmatrix} = \begin{pmatrix} 1 + \frac{u^2}{R^2 - u^2 - v^2} & \frac{uv}{R^2 - u^2 - v^2} \\ \frac{vu}{R^2 - u^2 - v^2} & 1 + \frac{v^2}{R^2 - u^2 - v^2} \end{pmatrix}.$$

Offensichtlich sind die Elemente dieser Metrik für $u^2 + v^2 = R^2$ - also auf dem Äquator - nicht mehr definiert (bzw. gehen gegen unendlich). Dies ist ein typisches Beispiel für eine Koordinatensingularität, denn der Äquator ist natürlich keine singuläre Punktmenge der Kugeloberfläche, sondern das gewählte Koordinatensystem ist dort nicht mehr sinnvoll.

5.10.4 Intrinsische Bedeutung der Metrik

Wir wollen nun einen Aspekt der Metrik ansprechen, der anschaulich weniger den Tangentialraum an eine Mannigfaltigkeit betrifft sondern die Mannigfaltigkeit selber. Betrachten wir dazu als Beispiel eine Kugeloberfläche.

Seit jeher ist den Kartographen bekannt, dass man nicht die gesamte Oberfläche der Erde auf eine planare Karte zeichnen kann, ohne dass irgendwelche Längen verzerrt oder gerade Linien (dem Ausschnitt eines Großkreises folgend) krumm dargestellt werden. Wir haben eine Fläche wie die Kugeloberfläche in der Parameterdarstellung als Abbildung eines Bereichs $U \subset \mathbf{R}^2$ in den einbettenden Raum beschrieben: $(u, v) \rightarrow \vec{x}(u, v)$. Der Parameterbereich U kann dabei als eine „Karte“ im herkömmlichen Sinn verstanden werden. Jede Struktur auf der Fläche

(beispielsweise die Umrisse der Kontinente auf der Erde) lässt sich als Kurve beschreiben, und diese Kurve definiert auch eine Kurve in der Parameterfläche U . Eine bekannte Darstellung der Erdoberfläche ist beispielsweise die Darstellung in Winkelvariablen (φ, θ) , bei der die Breitengrade als waagerechte Linien und die Längengrade als senkrechte Linien dargestellt werden.

Wenn wir auf der Karte Entfernungen zwischen zwei Parameterpunkten ausmessen, so hat diese Entfernung im Allgemeinen nichts mit der tatsächlichen Entfernung zu tun, die zwischen den durch die Parameter beschriebenen Punkten herrscht. Für nicht zu große Ausschnitte der Erdoberfläche genügt es meistens, einen „Maßstab“ anzugeben. Auf einem Stadtplan von Freiburg mit Maßstab 1:8000 entspricht einem Abstand von einem Zentimeter auf der Karte in Wirklichkeit ein Abstand von 8000 Zentimeter oder 80 Meter.

Eine solche Maßstabsangabe ist jedoch bei Weltkarten streng genommen nicht mehr sinnvoll. Es würde bedeuten, dass beispielsweise auf einer Weltkarte mit Maßstab 1:300000000 ein Zentimeter Abstand zwischen zwei Punkten in Wirklichkeit 300 km entsprechen. Betrachtet man sich die Karte jedoch genauer und vergleicht mit einem Globus, so erkennt man, dass diese Angabe nicht überall wirklich stimmt. Meist weichen die tatsächlichen Entfernungen in der Nähe der Pole (oder der Ränder der Karte) von dieser aus dem Maßstab berechneten Entfernung ab. Das liegt genau daran, dass man die Kugeloberfläche nicht maßstabsgetreu auf einer Ebene abbilden kann. Wir müssen also streng genommen an jedem Punkt einen Maßstab definieren, und zwar für jede der beiden Koordinatenrichtungen. Es kann nämlich sein, dass die Abstände in Nord-Süd-Richtung maßstabsgetreu wiedergegeben sind (d.h., überall auf der Karte mit demselben Maßstab), die Abstände in Ost-West-Richtung jedoch an den Rändern der Karte maßstabsverzerrt sind. Genau diesen orts- und richtungsabhängigen Maßstab beschreibt das metrische Feld!

Wir betrachten zwei Punkte $p \simeq \vec{x}(u^1, u^2)$ und $p' \simeq \vec{x}(u'^1, u'^2)$ auf der Fläche, die nahe beieinander liegen sollen. Der Abstand zwischen diesen beiden Punkten,

$$ds = \sqrt{(\vec{x}(u'^1, u'^2) - \vec{x}(u^1, u^2))^2}$$

soll also klein sein. In diesem Fall erwarten wir, dass auch die Koordinaten, durch welche die beiden Punkte beschrieben werden, nahe beieinander liegen (sofern nicht eine Koordinatensingularität vorliegt):

$$(u'^1, u'^2) = (u^1 + du^1, u^2 + du^2).$$

Wir erhalten somit:

$$\begin{aligned}\vec{x}(u^1, u^2) - \vec{x}(u^1, u^2) &= \vec{x}(u^1 + du^1, u^2 + du^2) - \vec{x}(u^1, u^2) \\ &= \frac{\partial \vec{x}(u^1, u^2)}{\partial u^1} du^1 + \frac{\partial \vec{x}(u^1, u^2)}{\partial u^2} du^2 + O((du^i)^2).\end{aligned}$$

Die partiellen Ableitungen nach den Parametern sind gerade die Basisvektoren im Tangentialraum:

$$\vec{f}_\alpha(p) = \frac{\partial \vec{x}(u^1, u^2)}{\partial u^\alpha}.$$

Für den „infinitesimalen“ Abstand ds erhalten wir somit:

$$ds = \sqrt{\left(\vec{f}_1(p)du^1 + \vec{f}_2(p)du^2\right)^2} = \sqrt{\sum_{\alpha,\beta} g_{\alpha\beta}(p) du^\alpha du^\beta}. \quad (5.45)$$

Mithilfe des metrischen Feldes können wir also an jedem Punkt (u^1, u^2) unserer Karte bestimmen, welchen Abstand zwei nahe benachbarte Punkte in Wirklichkeit haben. Das metrische Feld ist ein „orts- und richtungsabhängiger Maßstab“.

In der allgemeinen Relativitätstheorie gibt man oft den metrischen Feldtensor in der Form

$$ds^2 = \sum_{\alpha,\beta} g_{\alpha\beta} du^\alpha du^\beta \quad (5.46)$$

an. Dabei bezeichnet ds^2 das Quadrat des physikalischen Abstands zwischen zwei (physikalisch eng benachbarten) Ereignissen, den man mit tatsächlichen Maßstäben oder Uhren ausmessen kann. du^α bezeichnet die Differenz in der Koordinate u^α , durch welche man die beiden Ereignisse in einer Karte ausdrückt. Das metrische Feld gibt die Beziehung zwischen dem physikalischen Abstand und den Koordinatenabständen zweier Ereignisse an.

Kapitel 6

Elemente der Vektoranalysis

Wir wollen uns in den nächsten Kapiteln mit Feldern beschäftigen. Nach ihrem Transformationsverhalten unter einem (aktiven oder passiven) Koordinatenwechsel unterscheidet man skalare Felder, Vektorfelder, Pseudo-Vektorfelder, Tensorfelder, etc. Wir werden zunächst dieses Transformationsverhalten nochmals genauer untersuchen. Anschließend werden wir einige Grundlagen der Vektoranalysis einführen, dazu zählen die Begriffe „Gradient, Rotation, Divergenz“, das Linienintegral sowie Flächen- und Volumenintegrale. Schließlich behandeln wir die wichtigen Integralsätze: Stokes'scher Satz und Gauß'scher Satz. Diese beiden Integralsätze sind eine mehrdimensionale Verallgemeinerung der bekannten Relation, wonach das Integral über die Ableitung einer Funktion durch den Wert der Funktion an den Integrationsgrenzen gegeben ist.

6.1 Tensorfelder

Ganz allgemein kann man Felder als Abbildungen einer Mannigfaltigkeit M , die meist dem Ortsraum der möglichen Lagen von Körpern entspricht, in eine Menge auffassen, die oft ein Teilraum eines Vektorraums V ist. M muss dabei nicht notwendigerweise der $\mathbf{R}^3$ sein, insbesondere handelt es sich beim Ortsraum in der Newton'schen Mechanik nicht um den Vektorraum $\mathbf{R}^3$ (die Menge aller 3-Tupel reeller Zahlen), sondern um den affinen euklidischen Raum E_3 . In diesem affinen Raum ist kein Nullpunkt ausgezeichnet, und Punkte in diesem Raum können daher zunächst auch nicht mit Vektoren („gerichteten Verbindungslinien“) identifiziert werden. Die Verbindungslinien zwischen zwei Punkten lassen sich allerdings mit Vektoren in Beziehung setzen, d.h., durch die willkürliche Auszeichnung eines Nullpunkts O wird der affine Raum zu einem Vektorraum.

Die Bildmenge V kann nahezu jede beliebige Menge sein. Bei einem Temperaturfeld ist V der $\mathbf{R}^+$, ähnlich bei einem Druckfeld. Beim elektrischen und magnetischen Feld ist der Bildraum selber wieder ein $\mathbf{R}^3$, der Raum der möglichen Feldstärken. Andere Beispiele für Felder sind das Geschwindigkeitsfeld in der Strömung einer Flüssigkeit oder eines Gases, das Gravitationsfeld oder das metrische Feld in der Relativitätstheorie.

Da wir in diesem Kapitel keine Dynamik von Feldern betrachten, schreiben wir auch keine explizite Zeitabhängigkeit. Dynamische Felder φ entsprechen Abbildungen von $M \times \mathbf{R}$ in V bzw. $(p, t) \rightarrow \varphi(p, t)$.

Wählt man auf M ein Bezugssystem (einen Koordinatenursprung) und eine (orthonormale) Basis, so entspricht das einer bijektiven Abbildung $\phi : M \rightarrow \mathbf{R}^3$. Wir sollten daher zwischen dem Feld φ über M

$$\varphi : M \longrightarrow V \quad p \in M \mapsto \varphi(p) \in V$$

und dem Feld in einer Koordinatenbasis

$$\varphi_\phi := \varphi \circ \phi^{-1} : \mathbf{R}^3 \longrightarrow V \quad (x_1, x_2, x_3) \mapsto \varphi_\phi(x_1, x_2, x_3)$$

unterscheiden. Ein affiner Wechsel des Bezugssystems in M kann allgemein als Translation um einen Vektor $\vec{a}$ und eine Rotation des (orthonormalen) Koordinatensystems $\{\vec{e}_i\}$ verstanden werden. (Wir betrachten in diesem Kapitel keine krummlinigen Koordinatensysteme.) Sei $\phi' : M \rightarrow \mathbf{R}^3$ ein anderes Bezugssystem, so ist $\phi' \circ \phi^{-1} : \mathbf{R}^3 \rightarrow \mathbf{R}^3$ eine affine Abbildung auf dem $\mathbf{R}^3$, also im allgemeinen die Kombination aus einer Rotation und einer Translation. Wir beschränken uns im Folgenden auf Rotationen und halten den Nullpunkt des Bezugssystems fest. Das Transformationsverhalten der Werte des Feldes im Bildraum V als Folge einer Basistransformation im Urbildraum $\mathbf{R}^3$ bestimmt den Tensorcharakter des Feldes.

6.1.1 Skalare Felder

Wir betrachten zunächst ein Feld, das als Abbildung vom Raum M in die reellen Zahlen definiert ist:

$$\varphi : \mathbf{R}^3 \longrightarrow \mathbf{R} \quad \vec{x} \mapsto \varphi(\vec{x}) .$$

Beispiele sind (auf nicht zu kleinen Skalen) das Temperaturfeld und das Dichtefeld in der Meteorologie, das Energiefeld von elektrischen bzw. magnetischen

Feldern (also $|\vec{E}(\vec{x})|^2$ und $|\vec{B}(\vec{x})|^2$) oder der Betrag des Geschwindigkeitsfeldes in der Hydrodynamik ($|\vec{v}(\vec{x})|$).

Bei einem skalaren Feld wird somit einem bestimmten Raumpunkt p ein fester Wert zugeordnet. Dieser Zahlenwert hängt nicht davon ab, in welchem Bezugssystem man den Punkt p durch einen Vektor $\vec{x} = Op$ ausdrückt, bzw. in welcher Basis man die Komponenten $\{x_i\}$ dieses Punktes schreibt. Der Bildraum ($\mathbf{R}$) ist von der Wahl eines Koordinatensystems für den Punkt p unabhängig.

In einem ausgezeichneten Koordinatensystem wird das skalare Feld zu einer Funktion von drei Parametern $\{x_i\}$. In einer anderen Basis werden demselben Vektor $\vec{x}$ andere Koordinaten $\{x'_i\}$ zugeordnet, und dadurch ändert sich die funktionale Abhängigkeit des skalaren Feldes von diesen drei Koordinaten. Für die neue Funktion φ' soll gelten

$$\varphi_{\phi'}(x'_1, x'_2, x'_3) = \varphi(p) = \varphi_{\phi}(x_1, x_2, x_3) . \quad (6.1)$$

Im vorliegenden Fall haben wir eine passive Basistransformation betrachtet. Wir können aber auch aktive Transformationen untersuchen, d.h., der Punkt p wird in einen Punkt $p' = Rp$ überführt, der aus p durch eine Rotation R des Raums um den Bezugspunkt hervorgeht. In diesem Fall wird die Situation durch ein neues skalares Feld φ' beschrieben. Damit physikalisch dieselbe Aussage gemacht wird, muss der Wert von φ' an dem neuen Punkt p' gleich dem Wert des alten skalaren Feldes an dem alten Punkt p sein:

$$\varphi'(p') = \varphi(p) \quad \text{bzw.} \quad \varphi'(Rp) = \varphi(p) \quad \text{bzw.} \quad \varphi'(p) = \varphi(R^{-1}p) .$$

Diese Gleichung definiert das Transformationsverhalten eines skalaren Feldes unter aktiven Transformationen des Raumes.

6.1.2 Vektorfelder

Vektorfelder sind Abbildungen vom Konfigurationsraum M in den dreidimensionalen Vektorraum $\mathbf{R}^3$:

$$\vec{F} : M \longrightarrow \mathbf{R}^3 \quad p \mapsto \vec{F}(p) .$$

In diesem Fall transformiert sich bei einer aktiven Transformation (Rotation) des Raumes auch das Vektorfeld mit, d.h., wir erhalten ein neues Feld $\vec{F}'$, das durch die Bedingung

$$\vec{F}'(p) = R\vec{F}(R^{-1}p)$$

definiert ist. Ein Feld, dessen Komponenten $\{F_i\}$ sich unter einer Transformation des Raumes nach obiger Formel transformieren, heißt *Vektorfeld*. Beispiele sind das elektrische Feld $\vec{E}(\vec{x})$ in der Elektrodynamik und das Geschwindigkeitsfeld $\vec{v}(\vec{x})$ in der Hydrodynamik.

Ein Vektorfeld ändert bei einer Punktspiegelung (orthogonale Transformation mit Determinante -1) auch sein Vorzeichen. Ein Pseudo-Vektorfeld hingegen ändert bei einer Punktspiegelung sein Vorzeichen nicht; Beispiel ist das Magnetfeld $\vec{B}(\vec{x})$.

6.2 Besondere Ableitungen von Feldern

6.2.1 Gradient, Divergenz und Rotation

Den *Gradient* [*gradient*] eines Feldes haben wir schon mehrfach angesprochen. Sei $\Phi(\vec{x})$ ein skalares Feld auf dem $\mathbf{R}^n$, dann ist der (duale) Vektor $\vec{\nabla}\Phi(\vec{x})$ mit den Komponenten

$$(\vec{\nabla}\Phi(\vec{x}))_i = \partial_i\Phi(\vec{x}) = \frac{\partial\Phi(\vec{x})}{\partial x^i}$$

der *Gradient* des Feldes Φ .

Für Vektorfelder $\vec{A}(\vec{x})$ kann man im Prinzip die Ableitung jeder Komponente des Feldes nach jeder Komponente der Koordinaten bilden. Dies definiert eine Ableitungsmatrix. Allgemein gilt für eine Abbildung $A : \mathbf{R}^n \rightarrow \mathbf{R}^m$:

$$A(\vec{x} + \vec{h}) = A(\vec{x}) + \left(DA(\vec{x}) \right) (\vec{h}) + O(h^2). \quad (6.2)$$

$DA(\vec{x})$ ist eine lineare Abbildung vom $\mathbf{R}^n$ in den $\mathbf{R}^m$. Angewandt auf den Vektor $\vec{h} \in \mathbf{R}^n$ erhält man somit ein Element des $\mathbf{R}^m$. In Komponentenschreibweise folgt:

$$A(\vec{x} + \vec{h})^i = A(\vec{x})^i + \sum_j (DA(\vec{x}))_j^i h^j + O(h^2).$$

Bei Vektorfeldern auf dem $\mathbf{R}^3$ (also Abbildungen vom $\mathbf{R}^3$ in den $\mathbf{R}^3$) erhält man allgemein für $DA(\vec{x})$ eine 3×3 -Matrix mit Komponenten

$$(DA(\vec{x}))_j^i = \frac{\partial A^i}{\partial x^j}.$$

Von diesen neun Komponenten sind zwei Kombinationen von besonderer Bedeutung: die Divergenz und die Rotation.

Die *Divergenz* [*divergence*] eines Vektorfeldes $\vec{A}$ ist definiert als:

$$\vec{\nabla} \cdot \vec{A} = \sum_i \partial_i A^i = \sum_i \frac{\partial A^i(\vec{x})}{\partial x^i}. \quad (6.3)$$

Die Divergenz eines Vektorfeldes ist eine skalare Größe. Sie ändert sich unter Koordinatentransformationen nicht.

Die *Rotation* [*rotation*] eines Vektorfeldes im $\mathbf{R}^3$ ist definiert als:

$$(\vec{\nabla} \times \vec{A})^i = \sum_{jk} \epsilon_{ijk} \frac{\partial A^k}{\partial x^j}. \quad (6.4)$$

Die Rotation eines Vektorfeldes ist somit wieder ein Vektorfeld. Ausgeschrieben in Komponenten erhalten wir:

$$\begin{aligned} (\vec{\nabla} \times \vec{A})^1 &= \frac{\partial A^3}{\partial x^2} - \frac{\partial A^2}{\partial x^3} \\ (\vec{\nabla} \times \vec{A})^2 &= \frac{\partial A^1}{\partial x^3} - \frac{\partial A^3}{\partial x^1} \\ (\vec{\nabla} \times \vec{A})^3 &= \frac{\partial A^2}{\partial x^1} - \frac{\partial A^1}{\partial x^2}. \end{aligned}$$

Die von diesen beiden Ableitungsformen eines Feldes nicht berücksichtigten Komponenten der Ableitungen eines Vektorfeldes lassen sich in einem symmetrischen Tensor („Matrix“) zusammenfassen. Diese Matrix der symmetrisierten Ableitungen spielt beispielsweise beim so genannten Verschiebungsfeld in der Festkörperphysik eine wichtige Rolle. Im Folgenden werden wir diesen Anteil der Ableitungen von Vektorfeldern jedoch nicht weiter untersuchen.

6.2.2 Der Laplace-Operator

Bilden wir die Divergenz eines Gradienten, so erhalten wir:

$$\vec{\nabla} \cdot \vec{\nabla} \Phi(\vec{x}) = \sum_i \frac{\partial^2 \Phi(\vec{x})}{\partial (x^i)^2} = \frac{\partial^2 \Phi(\vec{x})}{\partial (x^1)^2} + \frac{\partial^2 \Phi(\vec{x})}{\partial (x^2)^2} + \frac{\partial^2 \Phi(\vec{x})}{\partial (x^3)^2} = \Delta \Phi(\vec{x}). \quad (6.5)$$

Den Operator

$$\Delta = \sum_i \partial_i^2 = \frac{\partial^2}{\partial (x^1)^2} + \frac{\partial^2}{\partial (x^2)^2} + \frac{\partial^2}{\partial (x^3)^2} \quad (6.6)$$

bezeichnet man als *Laplace-Operator* [*Laplacian*]. Er kann auf skalare Felder wie auch auf Vektorfelder angewandt werden. Bei Vektorfeldern wirkt der Laplace-Operator auf jede der Komponenten unabhängig.

6.2.3 Beziehungen zwischen zweiten Ableitungen von Feldern

Die Rotation lässt sich auf Vektorfelder anwenden. Da der Gradient eines skalaren Feldes ein Vektorfeld ist, können wir die Rotation eines Gradientenfeldes bestimmen. Wir erhalten:

$$\vec{\nabla} \times \vec{\nabla} \Phi(\vec{x}) = 0. \quad (6.7)$$

Beweis: Wir berechnen die i -te Komponente der Rotation des Gradientenfeldes

$$(\vec{\nabla} \times \vec{\nabla} \Phi(\vec{x}))^i = \epsilon^{ijk} \partial_j \partial_k \Phi(\vec{x}) = 0.$$

Der Grund für das Verschwinden der rechten Seite liegt darin, dass $\partial_j \partial_k$ symmetrisch in den Indizes j und k ist (dieses Argument setzt voraus, dass die Funktion $\Phi(\vec{x})$ mindestens zweimal stetig nach ihren Argumenten differenzierbar ist), wohingegen der ϵ -Tensor total antisymmetrisch ist. Die einzelnen Terme heben sich also paarweise weg.

Für die Divergenz einer Rotation folgt das gleiche Ergebnis:

$$\vec{\nabla} \cdot (\vec{\nabla} \times \vec{A}(\vec{x})) = 0. \quad (6.8)$$

Beweis:

$$\vec{\nabla} \cdot (\vec{\nabla} \times \vec{A}(\vec{x})) = \sum_{ijk} \partial_i \epsilon^{ij}{}_k \partial_j A^k(\vec{x}) = 0.$$

Das Verschwinden der rechten Seite gilt aus demselben Grund wie vorher: $\partial_i \partial_j$ ist ein in den Indizes i und j symmetrischer Ausdruck, wohingegen der ϵ -Tensor antisymmetrisch ist.

Für die Rotation einer Rotation erhalten wir:

$$\vec{\nabla} \times (\vec{\nabla} \times \vec{A}(\vec{x})) = \vec{\nabla}(\vec{\nabla} \cdot \vec{A}(\vec{x})) - \Delta \vec{A}. \quad (6.9)$$

Für den Beweis benötigen wir die folgende Identität für das Produkt zweier ϵ -Tensoren:

$$\sum_k \epsilon^{ij}{}_k \epsilon^{kl}{}_m = \delta^{il} \delta^j{}_m - \delta^i{}_m \delta^{jl}.$$

Zum Beweis berechnen wir wieder die i -te Komponente der Rotation der Rotation

(das Argument $\vec{x}$ lassen wir der Einfachheit halber weg):

$$\begin{aligned}
 (\vec{\nabla} \times (\vec{\nabla} \times \vec{A}))^i &= \sum_{jk} \epsilon^{ij}{}_k \partial_j (\vec{\nabla} \times \vec{A})^k \\
 &= \sum_{jklm} \epsilon^{ij}{}_k \epsilon^{kl}{}_m \partial_j \partial_l A^m \\
 &= \sum_{jlm} \left(\delta^{il} \delta^j{}_m - \delta^i{}_m \delta^{jl} \right) \partial_j \partial_l A^m \\
 &= \sum_j \left(\partial_j \partial_i A^j - \partial_j \partial^j A^i \right).
 \end{aligned}$$

Der erste Term entspricht dem Gradienten der Divergenz und der zweite Term dem Laplace-Operator angewandt auf die i -te Komponente von $\vec{A}$.

6.3 Kurvenintegrale

Kurvenintegrale lassen sich grundsätzlich sowohl über skalare Felder als auch Vektorfelder ausführen. Beim Kurvenintegral über Vektorfelder ist von besonderem Interesse, wenn über die Projektion von Vektorfeld auf Tangentialvektor integriert wird. Dies ist beispielsweise bei der Berechnung der Arbeit der Fall.

6.3.1 Die Bogenlänge

Als einfachstes Beispiel für eine Kurvenintegration über ein Skalarfeld betrachten wir zunächst die Bogenlänge [*arc length*] einer Kurve. Das Skalarfeld ist in diesem Fall identisch 1.

Gegeben sei eine Kurve $\gamma \simeq \vec{x}(t)$. Wir wollen die Länge der Kurve zwischen zwei Punkten $\vec{x}(t_0)$ und $\vec{x}(t_1)$ bestimmen. Für zwei infinitesimal benachbarte Punkte $\vec{x}(t + dt)$ und $\vec{x}(t)$ auf der Kurve gilt:

$$\Delta \vec{s} = (\vec{x}(t + \Delta t) - \vec{x}(t)) = \frac{d\vec{x}}{dt} \Delta t + O(\Delta t^2), \quad (6.10)$$

bzw.

$$d\vec{s} = \frac{d\vec{x}}{dt} dt. \quad (6.11)$$

Für den Betrag dieses Elements folgt:

$$ds = \sqrt{d\vec{s} \cdot d\vec{s}} = \sqrt{\frac{d\vec{x}(t)}{dt} \cdot \frac{d\vec{x}(t)}{dt}} dt.$$

Die Bogenlänge ergibt sich aus der Integration über die Bogenlängenelemente:

$$L = \int_{\gamma} ds = \int_{t_0}^{t_1} \left| \frac{d\vec{x}(t)}{dt} \right| dt. \quad (6.12)$$

Nun handelt es sich um ein gewöhnliches Integral über das Argument t .

Für die Bogenlängenparametrisierung ist der Betrag des Tangentialvektors eins und man erkennt, dass die Länge der Kurve genau der Differenz im Kurvenparameter entspricht. Daher auch die Bezeichnung *Bogenlängenparametrisierung*.

Beispiel: Eine Funktion $f(x)$ lässt sich auch als Graph (d.h., als Kurve) in der xy -Ebene darstellen: $y = f(x)$. Wählen wir $t = x$ als den Parameter dieser Kurve, so können wir in der Parameterdarstellung für diesen Graph auch schreiben:

$$\vec{x}(t) = (t, f(t)).$$

Damit folgt für das Bogenlängenelement:

$$ds = \sqrt{\left(\frac{d\vec{x}(t)}{dt}\right)^2} dt = \sqrt{\left(\frac{dx^1(t)}{dt}\right)^2 + \left(\frac{dx^2(t)}{dt}\right)^2} dt = \sqrt{1 + f'(t)^2} dt.$$

Die Länge der Kurve ergibt sich somit zu:

$$L = \int_{t_0}^{t_1} \sqrt{1 + f'(t)^2} dt.$$

Statt des Parameters t können wir natürlich ebenso gut x als Parameter schreiben:

$$L = \int_{x_0}^{x_1} \sqrt{1 + f'(x)^2} dx.$$

6.3.2 Kurvenintegration über skalare Felder

Wir betrachten nun ein skalares Feld $U(\vec{x})$ über dem $\mathbf{R}^3$ und eine Kurve γ in der Parameterdarstellung $\gamma(t) = \vec{x}(t)$. Für die Integration folgt:

$$I = \int_{\gamma} \Phi(\vec{x}) ds = \int_{t_0}^{t_1} \Phi(\vec{x}(t)) \left| \frac{d\vec{x}(t)}{dt} \right| dt.$$

Dieses Integral entspricht wiederum einem gewöhnlichen Integral über den Parameter t .

6.3.3 Berechnung der Arbeit

In der Physik ist die Arbeit, die mit einer Verschiebung eines Körpers verbunden ist, definiert als das Produkt aus der Kraft F , gegen welche die Verschiebung erfolgt, und der zurückgelegten Wegstrecke s .

$$W = F \cdot s.$$

Dabei spielt nur die Komponente der Kraft in Verschiebungsrichtung eine Rolle. Ist die Kraft entlang des Weges nicht konstant oder ändert sich die Wegrichtung, so muss über den Weg integriert werden:

$$W = \int_{\gamma} \vec{F} \cdot d\vec{s}. \quad (6.13)$$

Ist der Weg γ durch eine allgemeine Kurvenparametrisierung gegeben, so können wir wieder

$$d\vec{s} = \frac{d\vec{x}(t)}{dt} dt$$

schreiben und erhalten:

$$W = \int_{t_0}^{t_1} \left[\vec{F}(\vec{x}(t)) \cdot \frac{d\vec{x}(t)}{dt} \right] dt. \quad (6.14)$$

Auch dieses Integral ist nun ein gewöhnliches Integral über eine skalare Funktion des Kurvenparameters t .

Viele Kräfte lassen sich als so genannte *Gradientenfelder* schreiben, d.h., es gibt ein Potenzial $U(\vec{x})$, dessen (negativer) Gradient gleich der Kraft $\vec{F}$ ist. In diesem Fall erhalten wir für die Arbeit:

$$W = - \int_{t_0}^{t_1} \left[\vec{\nabla} U(\vec{x}) \cdot \frac{d\vec{x}(t)}{dt} \right] dt. \quad (6.15)$$

Zu integrieren ist also über das Skalarprodukt des Gradienten von $U(\vec{x})$ mit dem Geschwindigkeitsvektor an die Kurve $\vec{x}(t)$. In Kap. 5.4 hatten wir schon gezeigt, dass

$$\frac{dU(\vec{x}(t))}{dt} = \vec{\nabla} U(\vec{x}) \cdot \frac{d\vec{x}(t)}{dt}$$

ist. Wir erhalten für die Arbeit somit:

$$W = - \int_{t_0}^{t_1} \frac{dU(\vec{x}(t))}{dt} dt = U(\vec{x}(t_0)) - U(\vec{x}(t_1)). \quad (6.16)$$

Ist die Kraft also ein Gradientenfeld, hängt die Arbeit, die bei einer Verschiebung gegen diese Kraft geleistet werden muss, nur vom Anfangs- und Endpunkt des Weges ab.

Es gilt aber auch der umgekehrte Sachverhalt: Sei $\vec{F}(\vec{x})$ ein Vektorfeld, für das jedes Wegintegral unabhängig von der Form des Weges nur vom Anfangs- und Endpunkt abhängt; dann ist $\vec{F}$ ein Gradientenfeld. Das zugehörige skalare Feld lässt sich explizit über ein Wegintegral konstruieren:

$$U(\vec{x}) = - \int_{\vec{x}_0}^{\vec{x}} \vec{F}(\vec{x}) \cdot d\vec{x} .$$

U ist bis auf eine Integrationskonstante (die durch die Wahl des Anfangspunktes des Weges festgelegt werden kann) eindeutig.

Wir hatten bereits bewiesen, dass die Rotation eines Gradientenfeldes verschwindet:

$$\nabla \times \nabla \Phi = 0 .$$

In einem einfach zusammenhängenden Raumgebiet lässt sich jedes Vektorfeld, dessen Rotation verschwindet, als Gradientenfeld schreiben. Diesen Satz werden wir weiter unten beweisen.

6.4 Flächenintegrale und Stokes'scher Satz

6.4.1 Flächenelement und Flächenintegral im $\mathbf{R}^3$

Es soll nun beschrieben werden, wie man ein Skalarfeld bzw. ein Vektorfeld über eine Fläche $\mathcal{F}$ integriert. Die Fläche $\mathcal{F}$ sei durch eine Parameterdarstellung definiert: $(u, v) \ni U \subset \mathbf{R}^2 \rightarrow \mathbf{R}^3$. Das Bild dieser Abbildung beschreibt gerade die Fläche. Der Einfachheit halber sei der Wertebereich U von (u, v) das Quadrat $0 \leq u, v \leq 1$. Das ist aber keine wesentliche Einschränkung.

Als nächsten Schritt benötigen wir ein Maß auf der Fläche, also das Flächenelement $d\vec{f}$. Dieses Flächenelement $d\vec{f}$ sollte senkrecht auf der Fläche stehen und der Betrag sollte proportional zum Flächeninhalt sein. Es gilt:

$$d\vec{f} = \left(\frac{\partial \vec{x}(u, v)}{\partial u} \times \frac{\partial \vec{x}(u, v)}{\partial v} \right) du dv .$$

Offensichtlich ist

$$d\vec{x}_u = \frac{\partial \vec{x}(u, v)}{\partial u} du$$

die Wegstrecke entlang der Linie $\vec{x}(u, v = \text{konst.})$, die mit der Geschwindigkeit $\partial\vec{x}(u, v)/\partial u$ für die „Zeit“ du durchlaufen wird. Entsprechend ist

$$d\vec{x}_v = \frac{\partial\vec{x}(u, v)}{\partial v} dv$$

die Wegstrecke entlang der Linie $\vec{x}(u = \text{konst.}, v)$, die mit der Geschwindigkeit $\partial\vec{x}(u, v)/\partial v$ für die „Zeit“ dv durchlaufen wird. $d\vec{x}_u$ und $d\vec{x}_v$ spannen ein Parallelogramm auf, dessen Fläche durch

$$|d\vec{f}| = |d\vec{x}_u \times d\vec{x}_v|$$

gegeben ist. Der Vektor $d\vec{f}$ steht senkrecht auf diesem Parallelogramm und zeigt (bei geeigneter Wahl der Reihenfolge von u und v) nach außen. Im Grenzfall du und dv gegen Null nähert sich die Fläche des Parallelogramms immer mehr dem Flächeninhalt der Fläche an, die zwischen den Punkten $\vec{x}(u, v)$, $\vec{x}(u + du, v)$, $\vec{x}(u, v + dv)$, $\vec{x}(u + du, v + dv)$ liegt. Daher ist $d\vec{f}$ unser gesuchtes Flächenelement.

Wir definieren nun das Flächenintegral über ein skalar Feld über die Fläche $\mathcal{F}$ als

$$\int_{\mathcal{F}} \Phi(\vec{x}) |d\vec{f}| := \int_0^1 \int_0^1 \Phi(\vec{x}(u, v)) \left| \frac{\partial\vec{x}(u, v)}{\partial u} \times \frac{\partial\vec{x}(u, v)}{\partial v} \right| du dv .$$

Nun handelt es sich um ein gewöhnliches Integral über zwei Variable u und v .

Entsprechend definierten wir den *Fluss* eines Vektorfeldes $\vec{F}$ durch eine Fläche $\mathcal{F}$ als

$$\int_{\mathcal{F}} \vec{F}(\vec{x}) \cdot d\vec{f} := \int_0^1 \int_0^1 \left[\vec{F}(\vec{x}(u, v)) \cdot \left(\frac{\partial\vec{x}(u, v)}{\partial u} \times \frac{\partial\vec{x}(u, v)}{\partial v} \right) \right] du dv .$$

Wir projizieren also an jedem Punkt der Fläche den Vektor $\vec{F}$ auf das Flächenelement $d\vec{f}$ und bilden anschließend das Integral über u und v .

6.4.2 Beispiel: Die Kugeloberfläche

Als Beispiel betrachten wir die Oberfläche einer Kugel vom Radius R . Als Parametrisierung wählen wir die beiden Winkel $\theta \in [0, \pi]$ und $\varphi \in [0, 2\pi]$ der Kugelkoordinaten. Die Kugeloberfläche wird dann beschrieben durch:

$$(\theta, \varphi) \mapsto \vec{x}(\theta, \varphi) = R(\cos \varphi \sin \theta, \sin \varphi \sin \theta, \cos \theta) .$$

Für das Flächenelement an einem Punkt $\vec{x}(\theta, \varphi)$ erhalten wir:

$$\begin{aligned}\frac{\partial \vec{x}}{\partial \theta} &= R(\cos \varphi \cos \theta, \sin \varphi \cos \theta, -\sin \theta) \\ \frac{\partial \vec{x}}{\partial \varphi} &= R(-\sin \varphi \sin \theta, \cos \varphi \sin \theta, 0) \\ \Rightarrow \left(\frac{\partial \vec{x}}{\partial \theta} \times \frac{\partial \vec{x}}{\partial \varphi} \right) &= R^2(\cos \varphi \sin^2 \theta, \sin \varphi \sin^2 \theta, \sin \theta \cos \theta) \\ &= R^2 \sin \theta \vec{n} \quad \text{mit } \vec{n} = \frac{\vec{x}}{|\vec{x}|} .\end{aligned}$$

Am Punkt $\vec{x}(\theta, \varphi)$ ist $\vec{n}$ ein Einheitsvektor, der radial nach außen zeigt, und das Flächenelement ist somit:

$$d\vec{f}(\theta, \varphi) = R^2 \sin \theta d\theta d\varphi \vec{n} .$$

Man beachte, dass der Faktor $R^2 \sin \theta$ gleich der Wurzel der Determinante der Metrik auf der Kugeloberfläche ist:

$$R^2 \sin \theta = \sqrt{\det g_{\alpha\beta}} = \sqrt{g_{11}g_{22} - g_{12}^2} .$$

Diese Beziehung gilt allgemein bei Flächen- bzw. Volumenintegrationen und wird später noch bewiesen.

Der Flächeninhalt der Kugeloberfläche ergibt sich damit zu:

$$A = \int_{\gamma} |d\vec{f}| = R^2 \int_0^\pi \sin \theta d\theta \int_0^{2\pi} d\varphi = 2\pi R^2 \int_{-1}^{+1} d(\cos \theta) = 4\pi R^2 . \quad (6.17)$$

6.4.3 Der Stokes'sche Satz

Die Rotation eines Vektorfeldes war durch folgende Vorschrift definiert:

$$(\vec{\nabla} \times \vec{F})^i = \epsilon^{ij}{}_k \partial_j F^k ,$$

bzw.

$$(\vec{\nabla} \times \vec{F}) \simeq \begin{pmatrix} \partial_2 F^3 - \partial_3 F^2 \\ \partial_3 F^1 - \partial_1 F^3 \\ \partial_1 F^2 - \partial_2 F^1 \end{pmatrix} ,$$

mit $\partial_i = \partial/\partial x^i$. Man bezeichnet die Rotation eines Vektorfeldes auch als seine *Wirbeldichte*. Um diese Bezeichnung zu rechtfertigen, wollen wir zunächst den Stokes'schen Satz beweisen.

Satz von Stokes: Sei $\vec{F}$ ein (stetig differenzierbares) Vektorfeld im $\mathbf{R}^3$ und $\mathcal{F}$ eine (differenzierbare) Fläche im $\mathbf{R}^3$ mit Rand $\partial\mathcal{F}$, dann gilt:

$$\int_{\mathcal{F}} (\vec{\nabla} \times \vec{F}) \cdot d\vec{f} = \int_{\partial\mathcal{F}} \vec{F} \cdot d\vec{x}. \quad (6.18)$$

Das Flächenintegral über die Rotation eines Vektorfeldes ist also gleich dem Integral des Vektorfeldes über den Rand der Fläche.

Zum Beweis benutzen wir folgende Relation:

$$(\vec{\nabla} \times \vec{F}) \cdot \left(\frac{\partial \vec{x}}{\partial u} \times \frac{\partial \vec{x}}{\partial v} \right) = \frac{\partial}{\partial u} \left(\vec{F} \cdot \frac{\partial \vec{x}}{\partial v} \right) - \frac{\partial}{\partial v} \left(\vec{F} \cdot \frac{\partial \vec{x}}{\partial u} \right).$$

Wir beweisen zunächst diese Relation:

$$\begin{aligned} (\vec{\nabla} \times \vec{F}) \cdot \left(\frac{\partial \vec{x}}{\partial u} \times \frac{\partial \vec{x}}{\partial v} \right) &= \epsilon^{ij}_k (\partial_j F^k) \epsilon_{ilm} \frac{\partial x^l}{\partial u} \frac{\partial x^m}{\partial v} \\ &= (\delta^j_l \delta_{km} - \delta^j_m \delta_{kl}) (\partial_j F^k) \frac{\partial x^l}{\partial u} \frac{\partial x^m}{\partial v} \\ &= (\partial_j F^k) \frac{\partial x^j}{\partial u} \frac{\partial x_k}{\partial v} - (\partial_j F^k) \frac{\partial x_k}{\partial u} \frac{\partial x^j}{\partial v} \\ &= \frac{\partial F^k}{\partial u} \frac{\partial x_k}{\partial v} - \frac{\partial F^k}{\partial v} \frac{\partial x_k}{\partial u} \\ &= \frac{\partial}{\partial u} \left(F^k \frac{\partial x_k}{\partial v} \right) - \frac{\partial}{\partial v} \left(F^k \frac{\partial x_k}{\partial u} \right). \end{aligned}$$

Beim letzten Schritt heben sich die Terme mit zweifachen Ableitungen von $\vec{x}$ gerade weg.

Mit diesem Satz und der Definition des Flächenintegrals erhalten wir also

$$\begin{aligned} \int_{\mathcal{F}} (\vec{\nabla} \times \vec{F}) \cdot d\vec{f} &= \int_0^1 \int_0^1 \left[\frac{\partial}{\partial u} \left(\vec{F} \cdot \frac{\partial \vec{x}}{\partial v} \right) - \frac{\partial}{\partial v} \left(\vec{F} \cdot \frac{\partial \vec{x}}{\partial u} \right) \right] du dv \\ &= \int_0^1 \left(\vec{F} \cdot \frac{\partial \vec{x}}{\partial v} \right) \Big|_{u=0}^{u=1} dv - \int_0^1 \left(\vec{F} \cdot \frac{\partial \vec{x}}{\partial u} \right) \Big|_{v=0}^{v=1} du. \end{aligned}$$

Auf der rechten Seite steht nun ein Integral über den Rand des Rechtecks, das den Wertebereich von (u, v) bildet. Dieses Rechteck wird auf den Rand der Fläche

$(\partial\mathcal{F})$ abgebildet. Daher steht auf der rechten Seite gerade ein Linienintegral über die Funktion $\vec{F}$ über den Rand der Fläche $\mathcal{F}$. Damit ist der Stokes'sche Satz bewiesen.

Wir wollen nun die Vorstellung von der Rotation eines Vektorfeldes als eine Wirbeldichte konkretisieren. Das Integral eines Vektorfeldes um einen geschlossenen Weg γ entspricht der Wirbelstärke dieses Vektorfeldes entlang dieses Weges. Wir wählen nun diesen Weg in einer Fläche senkrecht zu einem Normalenvektor $\vec{n}$. Außerdem sei dieser Weg sehr kurz und umschließe eine Fläche mit Flächeninhalt df . Dann gilt nach dem Stokes'schen Satz:

$$(\vec{\nabla} \times \vec{F}) \cdot \vec{n} = \lim_{df \rightarrow 0} \frac{\int_{\gamma} \vec{F} \cdot d\vec{x}}{df},$$

d.h. die Rotation, projiziert auf den Einheitsvektor $\vec{n}$, ist gleich der Wirbelstärke entlang des Weges γ dividiert durch die Fläche, die von γ umrandet wird, im Grenzfall verschwindender Fläche – also gleich der Wirbeldichte. Man beachte, dass es sich hierbei um eine *Flächendichte* handelt, d.h., das Integral von $\vec{\nabla} \times \vec{F}$ über eine Fläche ergibt die Wirbelstärke.

6.5 Volumenintegrale und der Satz von Gauß

6.5.1 Volumenelement und Volumenintegration

Während ein Linienelement $d\vec{x}$ und ein Flächenelement $d\vec{f}$ im $\mathbf{R}^3$ gerichtete (vektorielle) Größen sind, ist ein Volumenelement im $\mathbf{R}^3$ eine (pseudo-) skalare Größe. Daher betrachten wir auch nur Volumenintegrale über Skalarfelder. Bildet man das Volumenintegral über ein Vektorfeld, so erhält man einen Vektor, d.h., das Integral ist für jede Komponente gesondert zu bilden.

Das Volumenelement im $\mathbf{R}^3$ ist $d^3x = dx^1 dx^2 dx^3$. Will man jedoch über ein Volumen $\mathcal{V} \in \mathbf{R}^3$ integrieren, so ist es oft schwierig, die Integrationsgrenzen an $\{x^i\}$ zu berücksichtigen, da – mit Ausnahme der Integration über einen Quader – die Integrationsgrenzen von den Integrationsvariablen abhängen.

Daher ist es meist sinnvoll, eine Abbildung $\mathbf{R}^3 \rightarrow \mathbf{R}^3$ zu finden, die einen Quader $(u^1, u^2, u^3) \in [0, a] \times [0, b] \times [0, c]$ auf das Volumen $\mathcal{V}$ abbildet, wobei der Rand des Quaders auf den Rand des Volumens – d.h. die Oberfläche $\partial\mathcal{V}$ – abgebildet wird. Der Übergang von den Variablen $\{x^i\}$ zu den neuen Variablen

$\{u^i\}$ beinhaltet die Jakobi-Determinante:

$$J = \left| \frac{\partial(x^1, x^2, x^3)}{\partial(u^1, u^2, u^3)} \right| = \begin{vmatrix} \frac{\partial x^1}{\partial u^1} & \frac{\partial x^2}{\partial u^1} & \frac{\partial x^3}{\partial u^1} \\ \frac{\partial x^1}{\partial u^2} & \frac{\partial x^2}{\partial u^2} & \frac{\partial x^3}{\partial u^2} \\ \frac{\partial x^1}{\partial u^3} & \frac{\partial x^2}{\partial u^3} & \frac{\partial x^3}{\partial u^3} \end{vmatrix} = \det \left(\frac{\partial \vec{x}}{\partial u^1}, \frac{\partial \vec{x}}{\partial u^2}, \frac{\partial \vec{x}}{\partial u^3} \right).$$

Wir erhalten also:

$$\int_V \Phi(\vec{x}) d^3x = \int_0^a \int_0^b \int_0^c \Phi(\vec{x}(u^1, u^2, u^3)) J du^1 du^2 du^3.$$

Der Vorteil einer solchen Transformation $\vec{x} \rightarrow \vec{u}$ besteht darin, dass nun die Integrationsgrenzen unabhängig von den Integrationsparametern sind.

Die Jakobi-Determinante lässt sich auch leicht als das Volumen eines Parallelepipedes deuten, das von den drei infinitesimalen Vektoren

$$d\vec{x}_{u^i} = \frac{\partial \vec{x}(u^1, u^2, u^3)}{\partial u^i} du^i$$

aufgespannt wird:

$$J = \frac{\partial \vec{x}}{\partial u^1} \cdot \left(\frac{\partial \vec{x}}{\partial u^2} \times \frac{\partial \vec{x}}{\partial u^3} \right).$$

6.5.2 Beispiel: Die Kugel vom Radius R

Wir betrachten als Beispiel die Integration über das Volumen einer Kugel vom Radius R . Als Parameter wählen wir den Betrag eines Vektors $r \in [0, R]$, sowie die beiden Winkel $\theta \in [0, \pi]$ und $\varphi \in [0, 2\pi]$, die wir schon bei der Beschreibung der Kugeloberfläche benutzt haben. Wir erhalten somit als Parametrisierung der Kugel:

$$(r, \theta, \varphi) \longrightarrow \vec{x}(r, \theta, \varphi) = r(\cos \varphi \sin \theta, \sin \varphi \sin \theta, \cos \theta)$$

mit den angegebenen Bereichsgrenzen. Das neue Volumenelement bzw. die Jakobi-Determinante J lässt sich leicht berechnen, wenn man berücksichtigt, dass

$$\frac{\partial \vec{x}}{\partial r} = \vec{n} = (\cos \varphi \sin \theta, \sin \varphi \sin \theta, \cos \theta)$$

ist, und

$$\left(\frac{\partial \vec{x}}{\partial \theta} \times \frac{\partial \vec{x}}{\partial \varphi} \right) = r^2 \sin \theta \vec{n} .$$

Damit folgt

$$J = \vec{n} \cdot (r^2 \sin \theta) \vec{n} = r^2 \sin \theta .$$

Wir erhalten somit als neues Volumenelement

$$dV = r^2 \sin \theta dr d\theta d\varphi .$$

Für das Volumen folgt daher:

$$V = \int_{\mathcal{V}} dV = \int_0^R r^2 dr \int_0^\pi \sin \theta d\theta \int_0^{2\pi} d\varphi = \frac{4\pi}{3} R^3 . \quad (6.19)$$

6.5.3 Der Satz von Gauß – Die Divergenz als Quellendichte

Der Gauß'sche Satz besagt, dass das Volumenintegral über die Divergenz eines Vektorfeldes gleich dem Integral über die Oberfläche $\partial\mathcal{V}$ über das Vektorfeld selber ist:

$$\int_{\mathcal{V}} (\vec{\nabla} \cdot \vec{F}) d^3x = \int_{\partial\mathcal{V}} \vec{F} \cdot d\vec{f} .$$

Der Beweis erfolgt ähnlich, wie der Beweis des Stokes'schen Satzes. Wir gehen von folgender Identität aus, die hier nicht bewiesen werden soll, aber durch direktes Nachrechnen leicht überprüfbar ist:

$$\begin{aligned} \vec{\nabla} \cdot \vec{F}(\vec{x}(u^1, u^2, u^3)) \left[\frac{\partial \vec{x}}{\partial u^1} \cdot \left(\frac{\partial \vec{x}}{\partial u^2} \times \frac{\partial \vec{x}}{\partial u^3} \right) \right] &= \frac{\partial}{\partial u^1} \left[\vec{F} \cdot \left(\frac{\partial \vec{x}}{\partial u^2} \times \frac{\partial \vec{x}}{\partial u^3} \right) \right] + \\ &+ \frac{\partial}{\partial u^2} \left[\vec{F} \cdot \left(\frac{\partial \vec{x}}{\partial u^3} \times \frac{\partial \vec{x}}{\partial u^1} \right) \right] + \frac{\partial}{\partial u^3} \left[\vec{F} \cdot \left(\frac{\partial \vec{x}}{\partial u^1} \times \frac{\partial \vec{x}}{\partial u^2} \right) \right] . \end{aligned}$$

Nun setzen wir diese Identität in das Volumenintegral über die Divergenz eines Vektorfeldes ein und erhalten ein Oberflächenintegral, parametrisiert durch den Rand des Quaders von (u^1, u^2, u^3) , über das Vektorfeld.

Mit Hilfe des Gauß'schen Satzes können wir der Divergenz eines Vektorfeldes auch die Interpretation einer sogenannten *Quellendichte* geben. Es gilt:

$$\vec{\nabla} \cdot \vec{F} = \lim_{V \rightarrow 0} \frac{1}{V} \int_{\mathcal{V}} (\vec{\nabla} \cdot \vec{F}) dV = \lim_{V \rightarrow 0} \frac{1}{V} \int_{\partial\mathcal{V}} \vec{F} \cdot d\vec{f} .$$

Das Integral über die geschlossene Oberfläche eines Volumens über ein Vektorfeld ergibt aber gerade den Gesamtfluss dieses Vektorfeldes durch diese Oberfläche, also die „Gesamtzahl der Quellen“ des Vektorfeldes innerhalb des Volumens. Die Divergenz ist daher die Quellendichte. In diesem Fall handelt es sich um eine Volumendichte.

6.6 Die Kontinuitätsgleichung

Die angegebenen Integralsätze ermöglichen oft den Übergang zwischen einer lokalen Aussage, die an jedem Raumpunkt erfüllt sein muss, und einer globalen Aussage, die für beliebige Volumina oder Flächen erfüllt sein muss. Als Beispiel betrachten wir die Kontinuitätsgleichung für elektrische Ladungen.

$\rho(\vec{x}, t)$ beschreibe eine elektrische Ladungsdichte bzw. eine Ladungsverteilung zum Zeitpunkt t . D.h.

$$Q_V(t) = \int_V \rho(\vec{x}, t) d^3x$$

ist die Gesamtladung in einem Volumen V zum Zeitpunkt t . Wir fragen uns nun nach der Ladungsänderung in diesem Volumen, d.h.

$$\frac{d}{dt}Q_V(t) = \int_V \dot{\rho}(\vec{x}, t) d^3x .$$

Es ist jedoch kein Prozess in der Natur bekannt, der die Gesamtladung eines abgeschlossenen Systems verändert, d.h., Ladung ist eine Erhaltungsgröße. Die Ladungsänderung in einem Volumen muss mit einem Ladungsfluss durch die Oberfläche des Volumens verbunden sein. Ladungsfluss drücken wir durch die Stromdichte $\vec{j}(\vec{x})$ aus. Die Stromdichte zu einer Ladungsverteilung $\rho(\vec{x})$, die sich am Punkte $\vec{x}$ mit der Geschwindigkeit $\vec{v}(\vec{x})$ bewegt, ist

$$\vec{j}(\vec{x}) = \vec{v}(\vec{x})\rho(\vec{x}) .$$

Der Fluss an Ladung durch eine geschlossene Oberfläche ∂V pro Zeiteinheit ist:

$$\frac{d}{dt}Q_V(t) = - \int_{\partial V} \vec{j}(\vec{x}) \cdot d\vec{f} .$$

Das negative Vorzeichen gibt an, dass $\dot{Q}$ die Änderung der Ladungsmenge in V ist. Fließt beispielsweise Ladung aus V ab, so ist $\dot{Q}$ negativ. Der Strom ist nach Außen gerichtet, d.h., das Integral über die Stromdichte gibt die nach außen abfließende

Ladungsmenge an, die gleich dem negativen der Ladungsänderung im Inneren ist. Wir erhalten also die Aussage:

$$\int_{\mathcal{V}} \dot{\rho}(\vec{x}, t) \, d^3x = - \int_{\partial\mathcal{V}} \vec{j}(\vec{x}) \cdot d\vec{f}.$$

Das Volumenintegral über die zeitliche Änderung von ρ ist gleich dem Oberflächenintegral über den Strom. Diese Aussage gilt für jedes Volumen $\mathcal{V}$. Nach dem Gauß'schen Satz gilt aber auch

$$\int_{\mathcal{V}} \dot{\rho}(\vec{x}, t) \, d^3x = - \int_{\mathcal{V}} (\vec{\nabla} \cdot \vec{j}(\vec{x})) \, d^3x.$$

Da auch diese Aussage für jedes Volumen $\mathcal{V}$ gilt, können wir $\mathcal{V}$ beliebig um einen Punkt konzentrieren und gelangen so zu einer *lokalen* Aussage:

$$\dot{\rho}(\vec{x}, t) = - \vec{\nabla} \cdot \vec{j}(\vec{x}).$$

Die Gleichung bezeichnet man als *Kontinuitätsgleichung*. Abgeleitet haben wir sie aus einer Art Bilanzgleichung (für die Ladung) in einem beliebigen Volumen $\mathcal{V}$. Ganz ähnliche Bilanzgleichungen (in diesem Fall für den Impuls und die Energie) nutzt man in der Hydrodynamik zur Herleitung der Bernoulli-Gleichung bzw. der Navier-Stokes-Gleichung.

6.7 Jacobi-Determinanten und Metrik

Zum Abschluss soll noch eine Identität bewiesen werden, die in den vergangenen Abschnitten mehrfach an Spezialfällen aufgetreten ist, die aber allgemein gilt. Wir betrachten ein kartesisches Koordinatensystem im $\mathbf{R}^n$ (die Komponenten von Punkten sind $\{x^i\}$) und wir wechseln zu einem krummlinigen Koordinatensystem (mit Komponenten $\{u^i\}$). Dann müssen wir in einem Volumenintegral die Jacobi-Determinante dieser Variablentransformation berücksichtigen. Die Tangentialvektoren in einem Punkt p sind:

$$\vec{f}(p)_j = \frac{\partial \vec{x}(\{u\})}{\partial u^j}.$$

Diese Tangentialvektoren spannen (sofern keine Koordinatensingularität vorliegt) ein n -dimensionales Volumen auf und es gilt:

$$dx^1 dx^2 \cdots dx^n = \text{Vol} \left(\frac{\partial \vec{x}}{\partial u^1}, \frac{\partial \vec{x}}{\partial u^2}, \dots, \frac{\partial \vec{x}}{\partial u^n} \right) du^1 du^2 \cdots du^n.$$

Das Volumen ist aber gerade die Determinante der Matrix zu den n -Tangentialvektoren:

$$\text{Vol} \left(\frac{\partial \vec{x}}{\partial u^1}, \frac{\partial \vec{x}}{\partial u^2}, \dots, \frac{\partial \vec{x}}{\partial u^n} \right) = \det J$$

mit

$$J = \left(\frac{\partial \vec{x}}{\partial u^1}, \frac{\partial \vec{x}}{\partial u^2}, \dots, \frac{\partial \vec{x}}{\partial u^n} \right).$$

Die Matrixelemente der Matrix J sind:

$$J^i_j = \frac{\partial x^i}{\partial u^j}.$$

Die Matrixelemente der transponierten Matrix sind:

$$(J^T)_{ji} = \frac{\partial x_j}{\partial u^i}.$$

Wir bilden nun das Produkt der Matrizen J^T und J :

$$(J^T \cdot J)_{ij} = \sum_k (J^T)_{ik} J^k_j = \sum_k \frac{\partial x_k}{\partial u^i} \frac{\partial x^k}{\partial u^j} = \vec{f}(p)_i \cdot \vec{f}(p)_j = g(p)_{ij}.$$

Das Produkt der Jacobi-Matrix J mit ihrer Transponierten ist somit gleich der Metrik im Tangentialraum am Punkt p .

Die Determinante einer Matrix ist gleich der Determinante ihrer Transponierten und die Determinante des Produkts zweier Matrizen ist gleich dem Produkt der Determinanten. Damit folgt:

$$(\det J)^2 = \det g \quad \text{oder} \quad \det J = \sqrt{\det g}.$$

Das Volumenelement für krummlinige Koordinaten ist also gleich der Wurzel aus der Determinante der Metrik.

Kapitel 7

Bewegungsgleichungen, Symmetrien und Erhaltungsgrößen

In diesem Kapitel werden wir uns konkret mit einigen wichtigen Bewegungsgleichungen der Physik beschäftigen und dabei allgemeine Methoden zu ihrer Aufstellung und Lösung behandeln. Werden physikalische Größen mit Indizes geschrieben, so deutet die Stellung der Indizes an, ob es sich um Vektoren oder duale Vektoren handelt, ohne dass in allen Fällen eine Begründung dafür geliefert wird. Vielfach wird die Begründung erst aus dem Formalismus der theoretischen Mechanik deutlich.

7.1 Die Newton'schen Gesetze

In seiner *Philosophiae naturalis principia mathematica* (Mathematische Grundlagen der Naturphilosophie), kurz „Principia“, formulierte Newton im Jahre 1686 drei Grundgesetze der Physik, die auch heute noch gelehrt und (teilweise) ihre Gültigkeit haben. In dieser Principia bringt Newton unter anderem seine Vorstellungen von Raum und Zeit zum Ausdruck, und da diese Vorstellungen das Bild der Physik zum Teil bis auf den heutigen Tag prägen, handelt es sich nach wie vor um ein lesenswertes Werk.

1. Newton'sches Gesetz:

Jeder Körper verharrt in seinem Zustand der Ruhe oder der gleichförmig-geradlinigen Bewegung, sofern er nicht durch eingedrückte Kräfte zur Änderung seines Zustands gezwungen wird.

Dieses Gesetz bezeichnet man manchmal als das Trägheitsgesetz. Es behält auch in der speziellen Relativitätstheorie seine Gültigkeit, und selbst in der allgemeinen Relativitätstheorie gilt es noch, sofern der Begriff der „gleichförmig-geradlinigen“ Bewegung geeignet verallgemeinert wird.

2. Newton'sches Gesetz:

Die Bewegungsänderung ist der eingedrückten Bewegungskraft proportional und geschieht in der Richtung der geraden Linie, in der jene Kraft eingedrückt wird.

Dieses Gesetz wird heute häufig in der Form

$$\vec{F} = m \cdot \vec{a} \quad (7.1)$$

formuliert. Durch die vektorielle Form wird deutlich, dass die Beschleunigung in dieselbe Richtung zeigt wie die Kraft, wie es in obigem Gesetz zum Ausdruck kommt. In dieser Form gilt das Gesetz nicht mehr in der speziellen Relativitätstheorie. Allerdings verstand Newton unter der „Bewegung“ nicht einfach die Geschwindigkeit eines Körpers, sondern eher seinen Impuls ($\vec{p} = m\vec{v}$). Damit würde man das 2. Newton'sche Gesetz in folgender Form schreiben:

$$\vec{F} = \frac{d\vec{p}}{dt}. \quad (7.2)$$

Die Bewegungsänderung (Änderung des Impulses) ist direkt proportional zur Kraft. In dieser Form bleibt das Gesetz auch in der Relativitätstheorie gültig.

Während das zweite Newton'sche Gesetz nur die Proportionalität von Kraft und Bewegungsänderung fordert, werden diese beiden Größen in Gl. 7.2 sogar gleich gesetzt. Tatsächlich ist Newtons Formulierung hier genauer, denn physikalisch verifizieren lässt sich nur die Proportionalität der Größen auf beiden Seiten der Gleichung. Ihre Gleichsetzung ist eine Konvention, die wir durch die geeignete Wahl physikalischer Einheiten erzwungen haben und die nicht Teil des Naturgesetzes ist.

3. Newton'sches Gesetz:

Der Einwirkung ist die Rückwirkung immer entgegengesetzt und gleich, oder: die Einwirkungen zweier Körper aufeinander sind immer gleich und wenden sich jeweils in die Gegenrichtung.

In Kurzform sagt man heute meist: *Kraft gleich Gegenkraft*:

$$\vec{F}_{12} = -\vec{F}_{21}. \quad (7.3)$$

Auch dieses Gesetz bleibt in der Relativitätstheorie gültig, allerdings ist es nicht vollkommen unabhängig vom 2. Newton'schen Gesetz, da es zur operationalen

Definition der trägen Masse m notwendig ist. Auf diese Feinheiten soll hier aber nicht näher eingegangen werden. (Eine lesenswerte Behandlung dieses Problems findet man beispielsweise in dem klassischen Werk „Die Mechanik in ihrer Entwicklung“ von Ernst Mach, ursprünglich aus dem Jahre 1883, das aber immer wieder in Neuauflagen erschienen ist.)

Abschließend noch eine Bemerkung zur trägen Masse m im zweiten Newton'schen Gesetz: In diesem Gesetz bringt m eine Korpereigenschaft zum Ausdruck, die ein Maß für die Trägheit des Körpers ist (daher auch *träge Masse*), also für die „Kraft“, mit der sich der Körper einer von außen einwirkenden Kraft „widersetzt“.

Von dieser trägen Masse muss man konzeptionell die so genannte *schwere Masse* unterscheiden. Sie tritt im Gravitationsgesetz auf:

$$F = G \frac{m_1 m_2}{|\vec{x}_1 - \vec{x}_2|^2}.$$

(Die Richtungsabhängigkeit dieser Kraft soll hier nicht interessieren.) m_1 bzw. m_2 bezeichnen die Eigenschaften von Körpern, über die sie gegenseitig eine Kraft ausüben, ähnlich wie die elektrischen Ladungen im Coulomb-Gesetz. Daher bezeichnet man manchmal die schwere Masse auch als die „Ladung der Gravitation“. Konzeptionell könnten die schwere und die träge Masse vollkommen verschiedene Größen sein (ähnlich wie die Ladung und die träge Masse in der Elektrodynamik). Die experimentelle Tatsache, dass träge und schwere Masse proportional sind (und daher gleichgesetzt werden dürfen), bezeichnet man als das *Äquivalenzprinzip*. Es bildet die Grundlage der allgemeinen Relativitätstheorie.

7.2 Die Newton'schen Bewegungsgleichungen

Im Folgenden verwenden wir das zweite Newton'sche Gesetz, ausgedrückt durch Gl. 7.1, als Gesetz zur Bestimmung der Bahnkurve von Körpern. Da wir uns vorläufig nur mit Körpern beschäftigen, deren Ausdehnung vernachlässigt und auf die Bahnkurve keinen nennenswerten Einfluss hat, sprechen wir auch oft von Teilchen oder *Punktteilchen*. Die Kraft $\vec{F}$ soll dabei einer von außen auf das Teilchen einwirkenden Kraft entsprechen.

Die Parameterdarstellung der (zu bestimmenden) Bahnkurve des Teilchens sei $\vec{x}(t)$, wobei wir als Parameter der Kurve die Zeit t wählen. Die Geschwindigkeit des Teilchens ist dann durch $\vec{v}(t) = \dot{\vec{x}}(t)$ gegeben. Damit Gl. 7.1 sinnvoll angewandt werden kann, müssen wir die Kraft, die zu einem beliebigen Zeitpunkt t auf

das Teilchen wirkt, kennen. Ganz allgemein kann diese Kraft nicht nur vom Ort abhängen, sondern auch von der Geschwindigkeit des Teilchens (beispielsweise bei der Berücksichtigung von Widerstandskräften oder bei der Bewegung geladener Teilchen in einem Magnetfeld) und sogar explizit von der Zeit (z.B. bei der Bewegung geladener Teilchen in zeitabhängigen elektromagnetischen Feldern):

$$\vec{F} = \vec{F}(\vec{x}(t), \vec{v}(t), t).$$

Es hat den Anschein, als ob die Kenntnis der Kraft schon die Kenntnis der Bahnkurve voraussetzt, die allerdings erst aus der Bewegungsgleichung bestimmt werden soll. In den meisten Fällen ist die Kraft jedoch an den *möglichen* Orten des Teilchens (beispielsweise im gesamten $\mathbf{R}^3$) und für alle möglichen Geschwindigkeiten bekannt:

$$\vec{F} = \vec{F}(\vec{x}, \vec{v}, t).$$

Gleichung 7.1 ist nun eine Bedingung an die Bahnkurve $\vec{x}(t)$, die zu jedem Zeitpunkt erfüllt sein muss:

$$\vec{F}(\vec{x}(t), \vec{v}(t), t) = m\ddot{\vec{x}}(t). \quad (7.4)$$

In dieser Gleichung werden zu jedem Zeitpunkt der Ort der Bahnkurve sowie die erste Ableitung nach der Zeit (die Geschwindigkeit) und die zweite Ableitung nach der Zeit (die Beschleunigung) in Beziehung gesetzt. Eine solche Gleichung bezeichnet man als *gewöhnliche Differentialgleichung 2. Ordnung*. („Gewöhnlich“, weil nur nach einem Parameter — der Zeit t — abgeleitet wird und nur eine Funktion dieses Parameters gesucht ist, „2. Ordnung“, weil die höchste auftretende Ableitung die zweite Ableitung nach der Zeit ist.)

Gleichung 7.4 ist eine vektorielle Gleichung, d.h., es handelt sich eigentlich um drei Gleichungen (für jede Komponente eine). Das bedeutet jedoch nicht, dass jede dieser drei Gleichungen nur von jeweils einer Komponente der Bahnkurve abhängen kann. Beispielsweise könnte die Kraft vom Abstand $r = \sqrt{(x^1)^2 + (x^2)^2 + (x^3)^2}$ des Teilchens vom Kraftzentrum abhängen. In diesem Fall spricht man von einem *gekoppelten System* von Differentialgleichungen:

$$\begin{aligned} F_1(x^1(t), x^2(t), x^3(t), \dot{x}^1(t), \dot{x}^2(t), \dot{x}^3(t), t) &= m\ddot{x}^1(t) \\ F_2(x^1(t), x^2(t), x^3(t), \dot{x}^1(t), \dot{x}^2(t), \dot{x}^3(t), t) &= m\ddot{x}^2(t) \\ F_3(x^1(t), x^2(t), x^3(t), \dot{x}^1(t), \dot{x}^2(t), \dot{x}^3(t), t) &= m\ddot{x}^3(t). \end{aligned}$$

Solche Gleichungen lassen sich im Allgemeinen nicht in geschlossener Form lösen. Glücklicherweise sind viele Bewegungsgleichungen in der Physik nicht von dieser allgemeinen Form und lassen geschlossene Lösungen zu.

Hat man es nicht mit einem sondern mit mehreren Teilchen zu tun, so verallgemeinert sich die Anzahl der Komponenten entsprechend, die grundsätzliche Struktur der Gleichung bleibt jedoch gleich:

$$m \frac{d^2 \vec{x}_i(t)}{dt^2} = \vec{F}_i(\vec{x}_1(t), \dots, \vec{x}_N(t), \dot{\vec{x}}_1(t), \dots, \dot{\vec{x}}_N(t); t)$$

Der Index $i = 1, \dots, N$ bezieht sich nun auf eine Nummerierung der Teilchen. Die Kraft auf ein Teilchen kann allgemein eine Funktion der Orte und Geschwindigkeiten aller Teilchen sein.

7.2.1 Anfangsbedingungen

Eine Differentialgleichung wie Gl. 7.4 legt die Lösung im Allgemeinen nicht eindeutig fest. Es handelt sich lediglich um eine Bedingung, die zwischen der Bahnkurve $\vec{x}(t)$ und ihren ersten und zweiten Ableitungen zu jedem Zeitpunkt erfüllt sein muss. Für eine eindeutige Lösung müssen zusätzliche Bedingungen festgelegt werden. In den meisten Fällen handelt es sich dabei um so genannte *Anfangsbedingungen*. Im konkreten Fall bedeutet dies, dass man den Ort $\vec{x}(t_0)$ und die Geschwindigkeit $\dot{\vec{x}}(t_0)$ zu einem Zeitpunkt t_0 (nahezu) beliebig vorgeben kann und nun zu einer eindeutigen Lösung der Differentialgleichung gelangt.

Das folgende Argument soll plausibel machen, weshalb die Vorgabe des Orts und der Geschwindigkeit zu einem Zeitpunkt $t = 0$ die Lösung einer gewöhnlichen Differentialgleichung 2. Ordnung eindeutig festlegt. Gegeben sei eine Differentialgleichung der Form

$$f(\ddot{x}(t), \dot{x}(t), x(t)) = 0. \quad (7.5)$$

(Der Einfachheit halber betrachten wir nur eine Komponente. Die Verallgemeinerung auf mehrere Komponenten stellt keine wesentliche Schwierigkeit dar.) Will man eine solche Gleichung numerisch lösen, also beispielsweise auf einem Computer, so ersetzt man die erste und zweite Ableitung durch den ersten und zweiten Differenzenquotienten:

$$\dot{x}(t) \rightarrow \frac{x(t + \Delta t) - x(t)}{\Delta t} \quad \text{und} \quad \ddot{x}(t) \rightarrow \frac{x(t + 2\Delta t) - 2x(t + \Delta t) + x(t)}{\Delta t^2}.$$

Damit wird aus Gl. 7.5:

$$\tilde{f}(x(t + 2\Delta t), x(t + \Delta t), x(t)) = 0.$$

(Die Funktion $\tilde{f}$ hängt noch explizit von Δt ab, was für die numerische Behandlung zwar wichtig ist, im Folgenden jedoch nicht weiter berücksichtigt wird.)

Wenn man von singulären Fällen absieht, lässt sich diese Gleichung nach $x(t + 2\Delta t)$ auflösen:

$$x(t + 2\Delta t) = \hat{f}(x(t + \Delta t), x(t)).$$

Nun handelt es sich um eine *iterative Gleichung*: Sind zu irgendeinem Zeitpunkt t_0 (beispielsweise $t_0 = 0$) die beiden Werte $x(t_0)$ und $x(t_0 + \Delta t)$ bekannt, so lässt sich $x(t_0 + 2\Delta t)$ berechnen. Damit wiederum können wir $x(t_0 + 3\Delta t)$ bestimmen etc. Iterativ berechnen wir somit die Folge von Werten $x(t_0 + 2\Delta t), x(t_0 + 3\Delta t), \dots, x(t_0 + n\Delta t)$. Diese Folge liegt eindeutig fest, sofern die beiden ersten Werte $(x(t_0), x(t_0 + \Delta t))$ vorgegeben sind. Für kleine Werte von Δt wird man erwarten, dass die Folge von Punkten $x(t_0 + n\Delta t)$ die Lösung der kontinuierlichen Kurve $x(t)$ gut nähert. Im konkreten Fall ist numerisch zu überprüfen, ob Änderungen in Δt die Form der Lösung noch wesentlich beeinflusst.

In der iterativen Form der Gleichung wird deutlich, dass die Vorgabe von zwei Anfangswerten $(x(t_0)$ und $x(t_0 + \Delta t))$ die weiteren Punkte eindeutig festlegt. Statt $x(t_0)$ und $x(t_0 + \Delta t)$ kann man äquivalent auch $x(t_0)$ und $v(t_0) = (x(t_0 + \Delta t) - x(t_0))/\Delta t$ vorgeben. Die Lösung einer gewöhnlichen Differentialgleichung 2. Ordnung ist also durch die Angabe des Orts und der Geschwindigkeit zu einem bestimmten Zeitpunkt festgelegt.

7.2.2 Randbedingungen

Schwieriger wird die Situation, wenn man andere Bedingungen an die Lösung vorgibt, z.B. Randbedingungen. In diesem Fall werden zwei Werte der Kurve zu verschiedenen Zeitpunkten, beispielsweise $x(t_0)$ und $x(t_1)$, vorgegeben. In solchen Fällen ist es schwieriger zu beweisen, ob eine Lösung existiert bzw. ob diese Lösung eindeutig festgelegt ist.

Wie wir im nächsten Kapitel sehen werden, hat die allgemeine Lösung der Differentialgleichung des harmonischen Oszillators

$$\ddot{x}(t) = -\omega^2 x(t)$$

die Form

$$x(t) = a \cos \omega t + b \sin \omega t.$$

a und b können frei gewählt, beispielsweise durch die Anfangsbedingungen $x(0)$ und $v(0)$ ausgedrückt werden. $T = 2\pi/\omega$ bezeichnet man als die Periode des harmonischen Oszillators.

Gibt man nun als Randbedingungen $x(0)$ und $x(T)$ vor mit $x(0) \neq x(T)$, so besitzt die Gleichung überhaupt keine Lösungen. Jede Lösung der Differentialgleichung des harmonischen Oszillators hat nämlich die Eigenschaft, nach einer vollen Periode wieder an demselben Ort zu sein. Wählt man hingegen $x(0) = x(T)$, so liegt die Lösung nicht eindeutig fest, sondern es gibt immer noch eine einparametrische Schar von Lösungen. Man erkennt an diesem Beispiel, dass die Vorgabe von Randbedingungen (im Gegensatz zu der Vorgabe von Anfangsbedingungen) an die Lösungen der Newton'schen Bewegungsgleichungen problematisch sein kann.

7.2.3 Zustandsraum bzw. Phasenraum

Nach dem bisher Gesagten liegt die Lösungskurve der Newton'schen Bewegungsgleichungen für ein Teilchen fest, wenn der Ort und die Geschwindigkeit zu einem beliebigen Zeitpunkt bekannt sind. Daher sagt man auch, der *Zustand* eines Teilchens ist in der Newton'schen Mechanik durch Ort und Geschwindigkeit gegeben. Der Zustandsraum eines einzelnen Teilchens ist somit ein 6-dimensionaler Raum (drei Orts- und drei Geschwindigkeitskomponenten).

Für mehrere Teilchen ist der Zustandsraum entsprechend der Produktraum der Zustandsräume für die einzelnen Teilchen. Ein System aus N Teilchen hat einen $6N$ -dimensionalen Zustandsraum. Der Zustand der Teilchen wird durch einen Punkt in diesem Zustandsraum festgelegt. Statt der N Bahnkurven in einem $\mathbf{R}^3$ kann man auch die einzelne Bahnkurve im $\mathbf{R}^{6N}$ betrachten. Da diese Kurve im Zustandsraum bereits durch die Angabe des Zustands zu einem bestimmten Zeitpunkt festliegt, lässt sich diese Kurve sogar durch eine Differentialgleichung 1. Ordnung in den Zustandsvariablen beschreiben.

Dieser letzte Punkt wird verständlich wenn man bedenkt, dass sich jede Differentialgleichung 2. Ordnung

$$\ddot{x}(t) = f(x(t), \dot{x}(t))$$

auch als gekoppeltes System von Differentialgleichungen erster Ordnung schreiben lässt:

$$\dot{v}(t) = f(x(t), v(t)) \quad \text{und} \quad \dot{x}(t) = v(t).$$

Statt von Zustandsraum spricht man in der Mechanik auch oft von *Phasenraum*. Streng genommen besteht jedoch ein Unterschied: Der Zustandsraum ist der Raum der möglichen Anfangsbedingungen, durch die eine Lösungskurve eindeutig festgelegt ist. Der Phasenraum hat jedoch eine zusätzliche Struktur, die aus der speziellen Form der Newton'schen Gleichungen folgt. Besonders deutlich

wird diese Struktur, wenn man nicht den Ort und die Geschwindigkeit, sondern den Ort und den Impuls zur Beschreibung des Zustands eines Teilchens wählt. Diese Unterschiede werden im Rahmen der theoretischen Mechanik deutlicher.

7.2.4 Beispiel: Die freie Bewegungsgleichung

Wenn keine äußere Kraft vorhanden ist, lautet die Newton'sche Bewegungsgleichung

$$m \frac{d^2 \vec{x}(t)}{dt^2} = 0.$$

Die Masse fällt heraus und in Komponentenschreibweise gilt:

$$\ddot{x}^i(t) = 0. \quad (7.6)$$

Die allgemeine Lösung dieser Gleichung lautet:

$$x^i(t) = a^i t + b^i \quad (7.7)$$

für beliebige Konstanten a^i und b^i . Für die drei Komponenten erhalten wir also insgesamt sechs freie Konstanten, sodass die Lösungen (Gl. 7.7) die Differentialgleichung (7.6) erfüllen. Die Bedeutung der Konstanten ist leicht geklärt:

$$x^i(0) = b^i \quad \text{und} \quad v^i(0) = \dot{x}^i(0) = a^i.$$

b^i sind also die drei Komponenten des Orts des Teilchens zum Zeitpunkt $t = 0$, und a^i sind die drei Komponenten der Geschwindigkeit des Teilchens zum Zeitpunkt $t = 0$. Als Lösung erhalten wir somit:

$$\vec{x}(t) = \vec{v}(0) t + \vec{x}(0). \quad (7.8)$$

7.2.5 Beispiel: Der schiefe Wurf

Die Bewegung des schiefen Wurfs betrachten wir in der (x, y) -Ebene, wobei sich y auf die Höhe des Teilchens über dem Erdboden und x auf eine horizontale Komponente bezieht. Die Kraft auf das Teilchen ist die Gravitationskraft. Sie wirkt nur entlang der y -Richtung (sie ist zum Erdmittelpunkt gerichtet, also $F^x = 0$) und außerdem nehmen wir sie in Erdnähe als konstant an, d.h., sie hängt nicht von der y -Koordinate ab und ist durch $F^y = mg$ gegeben, wobei g die Erdbeschleunigung $g = 9,81 \text{ m/s}^2$ ist. Die Bewegungsgleichung wird zu:

$$m\ddot{x}(t) = 0 \quad \text{und} \quad m\ddot{y}(t) = -mg. \quad (7.9)$$

(Das Minuszeichen besagt, dass die Kraft zu kleineren Werten von y gerichtet ist.) Die beiden Gleichungen sind immer noch entkoppelt (die Komponenten mischen nicht). Die erste Gleichung haben wir im vorherigen Abschnitt schon gelöst:

$$x(t) = v^x(0)t + x(0) .$$

Die zweite Gleichung in (7.9) können wir zunächst durch m dividieren und erkennen, dass die Bewegung unabhängig von der Masse ist. Man beachte jedoch, dass dies eine unmittelbare Folgerung des Äquivalenzprinzips ist: Auf der linken Seite der Gleichung steht die träge Masse, auf der rechten Seite die schwere Masse. Gälte das Äquivalenzprinzip nicht, hinge die Bewegung von dem Verhältnis von schwerer zu träger Masse ab.

Die Gleichung

$$\ddot{y}(t) = -g$$

hat die allgemeine Lösung:

$$y(t) = -\frac{1}{2}gt^2 + v^y(0)t + y(0) , \quad (7.10)$$

wobei sich wieder leicht zeigen lässt, dass $y(0)$ die y -Komponente des Orts des Teilchens zum Zeitpunkt $t = 0$ und $v^y(0)$ die y -Komponente seiner Geschwindigkeit zum Zeitpunkt $t = 0$ sind.

Die allgemeinste Lösung von Gl. 7.9 lautet somit:

$$\vec{x}(t) \simeq (x(t), y(t)) = \left(v^x(0)t + x(0) , -\frac{1}{2}gt^2 + v^y(0)t + y(0) \right) ,$$

bzw.

$$\vec{x}(t) = -\frac{1}{2}gt^2\vec{e}_x + \vec{v}(0)t + \vec{x}(0) .$$

Die Bewegung entlang der x - und y -Richtung sind entkoppelt. Entlang der x -Achse handelt es sich um eine gleichförmig-geradlinige Bewegung, entlang der y -Achse um eine Parabel. Ist $v^y(0)$ positiv, hat diese Parabel einen Scheitelpunkt mit verschwindender Geschwindigkeit.

Eine sehr wichtige Klasse von Differentialgleichungen, die Gleichungen von harmonischen Oszillatoren, werden wir im nächsten Kapitel behandeln.

7.3 Symmetrien und Erhaltungsgrößen

Symmetrien und Erhaltungsgrößen spielen in der Physik eine wesentliche Rolle. Durch sie wird die Lösung vieler Bewegungsgleichungen überhaupt erst möglich. Dabei gibt es einen sehr engen Zusammenhang zwischen den Symmetrien eines physikalischen Systems und seinen Erhaltungsgrößen. Dieser Zusammenhang wird in der Theoretischen Mechanik unter dem Stichwort „Noether’sches Theorem“ genauer behandelt. Hier können wir diese Beziehung nur an einfachen Beispielen andeuten.

7.3.1 Erhaltungsgrößen

Eine Erhaltungsgröße $T(\vec{x}, \vec{v})$ ist eine Funktion auf dem Zustandsraum. Ausgewertet auf einer Bahnkurve $\vec{x}(t)$ ist die Erhaltungsgröße $T(\vec{x}(t), \vec{v}(t))$ (aufgefasst als Funktion von t) konstant. Einmal vorgegeben durch die Anfangsbedingungen ändert sie ihren Wert im Verlauf der Zeit nicht — sie bleibt erhalten. Beispiele für Erhaltungsgrößen sind die Energieerhaltung, die Drehimpulserhaltung, die Ladungserhaltung und — je nach System — viele weitere Größen. Dabei ist wichtig zu bemerken, dass eine Erhaltungsgröße $T(\vec{x}(t), \vec{v}(t))$ nicht für jede beliebige Kurve im Zustandsraum konstant ist (das wäre eine triviale Funktion), sondern nur für solche Bahnkurven, die den Bewegungsgleichungen genügen. Umgekehrt bedeutet dies, dass die Bahnkurve auf den Äquipotentialflächen der Erhaltungsgrößen verläuft. Jede Erhaltungsgröße schränkt die Form der Bewegung ein und erleichtert somit die Lösung des Problems.

7.3.2 Die Energieerhaltung

Im Folgenden nehmen wir eine wesentliche Einschränkung an die Form der Kraft vor, die aber in der Physik sehr häufig erfüllt ist. Wir nehmen an, die Kraft $\vec{F}$ lässt sich als Gradient einer skalaren Funktion $U(\vec{x})$ schreiben:

$$\vec{F}(\vec{x}) = -\vec{\nabla}U(\vec{x}). \quad (7.11)$$

Kräfte, die dieser Bedingung genügen, bezeichnet man als *konservative Kräfte*. Als Bewegungsgleichung erhalten wir in diesem Fall:

$$m \frac{d^2 \vec{x}(t)}{dt^2} = -\vec{\nabla}U(\vec{x}(t)), \quad (7.12)$$

bzw.

$$m\dot{\vec{v}}(t) = -\vec{\nabla}U(\vec{x}(t)). \quad (7.13)$$

Als erstes Beispiel für eine Erhaltungsgröße betrachten wir die Energie. Wir bilden auf beiden Seiten von Gl. 7.13 das Skalarprodukt mit $\vec{v}(t)$:

$$m\dot{\vec{v}}(t) \cdot \vec{v}(t) = -\vec{\nabla}U(\vec{x}(t)) \cdot \vec{v}(t). \quad (7.14)$$

Auf der linken Seite dieser Gleichung steht die Ableitung von $(\vec{v}(t))^2$ nach der Zeit:

$$\frac{d}{dt} \left(\frac{1}{2} m (\vec{v}(t))^2 \right) = m \vec{v}(t) \cdot \dot{\vec{v}}(t).$$

Auf der rechten Seite von Gl. 7.14 erkennen wir durch Vergleich mit Gl. 5.20 in Kapitel 5.4 ebenfalls eine Zeitableitung:

$$\frac{d}{dt} U(\vec{x}(t)) = \vec{\nabla}U(\vec{x}) \cdot \dot{\vec{x}}(t).$$

Wir können Gl. 7.14 daher in folgender Form schreiben:

$$\frac{d}{dt} \left(\frac{1}{2} m (\vec{v}(t))^2 + U(\vec{x}(t)) \right) = 0.$$

Da die Zeitableitung verschwindet, ist die Größe in Klammern zeitlich konstant:

$$E = \frac{1}{2} m (\vec{v}(t))^2 + U(\vec{x}(t)) = \text{konst.} \quad (7.15)$$

Diese Größe bezeichnet man als die Energie. Der erste Term entspricht der kinetischen Energie, der zweite Term der potenziellen Energie. Eine wesentliche Eigenschaft, die bei der Herleitung dieser Gleichung einging, war die explizite Zeitunabhängigkeit der Funktion $U(\vec{x}(t))$. Anderenfalls wäre bei der Zeitableitung von U noch ein zusätzlicher Term aufgetreten. Die Energieerhaltung ist unter anderem eine Folge der Tatsache, dass die fundamentalen physikalischen Gesetze nicht explizit zeitabhängig sind, dass ein Experiment heute also dieselben Ergebnisse liefert wie das gleiche Experiment vor zwei Wochen.

7.3.3 Die Drehimpulserhaltung

Wir machen nun zusätzlich die Annahme, dass das Potenzial $U(\vec{x}(t))$ nur eine Funktion von $r(t) = |\vec{x}(t)|$ ist. Das Potenzial zeigt also keine explizite Richtungsabhängigkeit. Die damit verbundene Symmetrie bezeichnet man als *Rotations-symmetrie*.

In diesem Fall gilt für den Gradienten des Potentials:

$$\vec{\nabla}U(r(t)) = U'(r(t)) \vec{\nabla}|\vec{x}(t)| = U'(r(t)) \frac{\vec{x}(t)}{|\vec{x}(t)|}.$$

Der Gradient von U zeigt also in radialer Richtung. Das Kreuzprodukt von $\vec{\nabla}U$ mit $\vec{x}$ verschwindet daher:

$$\vec{\nabla}U(|\vec{x}(t)|) \times \vec{x}(t) = \frac{U'(r(t))}{r(t)} (\vec{x}(t) \times \vec{x}(t)) = 0.$$

Bilden wir nun auf beiden Seiten von Gleichung 7.13 das Kreuzprodukt mit $\vec{x}(t)$, so erhalten wir auf der rechten Seite null, die linke Seite muss also verschwinden:

$$m\vec{a}(t) \times \vec{x}(t) = 0. \quad (7.16)$$

Wir zeigen nun, dass sich der Ausdruck auf der linken Seite als totale Zeitableitung schreiben lässt. Es gilt:

$$\frac{d}{dt} (\dot{\vec{x}}(t) \times \vec{x}(t)) = \ddot{\vec{x}}(t) \times \vec{x}(t) + \dot{\vec{x}}(t) \times \dot{\vec{x}}(t).$$

Der zweite Term auf der rechten Seite verschwindet, da das Kreuzprodukt von einem Vektor mit sich selber verschwindet. Der erste Term auf der rechten Seite entspricht aber der linken Seite von Gl. 7.16. Wir finden daher aus unserer Bewegungsgleichung eine neue Erhaltungsgröße:

$$\frac{d}{dt} \vec{L} = 0 \quad \text{mit} \quad \vec{L} = m\vec{x}(t) \times \vec{v}(t) = \vec{x}(t) \times \vec{p}(t).$$

Die Größe $\vec{L}$ bezeichnet man als *Drehimpuls*. Nach den Überlegungen aus Kapitel 3.4 handelt es sich beim Drehimpuls um einen Pseudovektor. Die Drehimpulserhaltung ist eine Folgerung aus der Rotationsinvarianz der Bewegungsgleichung.

Streng genommen haben wir weniger benutzt, als die Rotationsinvarianz: Wir können die Drehimpulserhaltung direkt aus der Newton'schen Gleichung (mit Kraft $\vec{F}(x, v, t)$) ableiten, sofern diese Kraft in radialer Richtung zeigt, also von der Form

$$\vec{F}(\vec{x}, \vec{v}, t) = f(\vec{x}, \vec{v}, t) \cdot \frac{\vec{x}}{|\vec{x}|} \quad (7.17)$$

ist, wobei f eine skalare Funktion von $\vec{x}$, $\vec{v}$ und t ist. In diesem Fall verschwindet das Kreuzprodukt von $\vec{F}$ mit $\vec{x}$. Für rotationsinvariante Systeme zeigt die

Kraft immer in radialer Richtung. Die Kraft kann sogar explizit zeitabhängig sein. Allerdings handelt es sich in den seltensten Fällen um konservative Kräfte. Beispielsweise sieht man sofort, dass für eine winkelabhängige Funktion $f(r, \theta, \varphi)$ das Integral für die Arbeit entlang geschlossener Wege nicht verschwinden muss: Das Integral für Wege auf einer Kugelfläche verschwindet, da $\vec{F} \perp d\vec{s}$, andererseits könnte man bei winkelabhängigem f zwei Kugelschalen bei unterschiedlichen Winkeln (und daher unterschiedlichen Kräften) verbinden und erhält verschiedene Anteile zur Arbeit.

Die Bahnkurve eines Teilchens, für das der Drehimpuls eine Erhaltungsgröße ist, verläuft in einer Ebene. Die Konstanz des Drehimpulses bedeutet nämlich insbesondere auch die Unveränderlichkeit seiner Richtung. Das bedeutet aber, dass $\vec{x}(t)$ und $\vec{v}(t)$ immer senkrecht zu dieser unveränderten Richtung stehen und daher die Bewegung des Teilchens auf diese Ebene beschränkt ist.

7.3.4 Symmetrien

Wir haben zwar bisher oft von der „Symmetrie“ bzw. „Invarianz“ eines Systems gesprochen, aber es wurde noch nicht genau definiert, was eine Symmetrie bzw. Invarianz ist.

Definition: Eine *Symmetrie* ist eine Gruppe G , die auf dem Raum der Lösungen X einer Gleichung wirkt (d.h., es gibt eine Darstellung der Gruppe G auf dem Raum der Lösungen X , bzw. es gibt eine Abbildung $G \times X \rightarrow X$ mit $(g, x) \rightarrow x^g$).

Oftmals wirkt die Gruppe G auf einen größeren Raum, von dem die Lösungen einer Gleichung nur eine Untermenge sind. Dann soll gelten: Ist x eine Lösung einer Gleichung, dann ist für jedes $g \in G$ auch x^g eine Lösung der Gleichung.

Betrachten wir als Beispiel die Newton'schen Bewegungsgleichungen. Die meisten Transformationsgruppen wirken zunächst auf dem Raum aller möglichen Bahnkurven $x(t)$. Handelt es sich bei den Transformationen jedoch um eine Symmetrie, so gilt: Ist $x(t)$ Lösung der Newton'schen Gleichung, dann ist auch $x^g(t)$ Lösung der Newton'schen Gleichung. Wir betrachten dazu einige Beispiele:

Die *Zeittranslationssymmetrie*: Wenn die Kraft nicht explizit von der Zeit abhängt, spielt der Anfangszeitpunkt einer Bewegung keine Rolle. Das bedeutet, dass mit jeder Lösung $x(t)$ auch $x(t + t_0)$ eine Lösung der Newton'schen Gleichungen ist, wobei $t_0 \in \mathbf{R}$ beliebig ist. Die Gruppe G ist die Gruppe der Zeittranslationen $t \rightarrow t + t_0$. Sie wirkt auf dem Raum aller Bahnkurven in der Form: $x(t) \rightarrow x^{t_0}(t) = x(t + t_0)$.

Die *Rotationssymmetrie*: Hat die Kraft auf ein Teilchen ein Kraftzentrum und hängt die Kraft nicht von der Richtung sondern nur vom Abstand von diesem

Kraftzentrum ab, so kann man das System beliebig um das Kraftzentrum rotieren, ohne die physikalischen Gesetze zu ändern. Die Gruppe ist nun die Rotationsgruppe G , dargestellt beispielsweise durch die 3×3 -Rotationsmatrizen R . Zu einer Bahnkurve $\vec{x}(t)$ erhalten wir die transformierte Bahnkurve durch die Anwendung der Matrix auf die Komponenten der Bahnkurve: $\vec{x}(t) \rightarrow \vec{x}^g(t) = (R\vec{x})(t)$. Ist ein System rotationsinvariant so folgt, dass mit jeder Lösung $\vec{x}(t)$ auch $(R\vec{x})(t)$ eine Lösung ist.

Zeitumkehrinvarianz: Die Gruppe G ist diesmal die diskrete Gruppe $G = T = \{+1, -1\}$. Sie besteht nur aus zwei Elementen. Die Darstellung des Elements $\{-1\}$ auf einer Bahnkurve besteht in der Umkehrung der Zeitrichtung: $\vec{x}(t) \rightarrow \vec{x}^-(t) = \vec{x}(-t)$. Wenn mit jeder Lösung $\vec{x}(t)$ auch $\vec{x}(-t)$ eine Lösung der Bewegungsgleichung ist, so bezeichnet man die Bewegungsgleichung als zeitumkehrinvariant. Die fundamentalen Gleichungen der Physik sind zeitumkehrinvariant, trotzdem beobachten wir im Alltag eine deutliche Abweichung von dieser Symmetrie. Diese Abweichung wird durch den zweiten Hauptsatz der Thermodynamik ausgedrückt, wonach die Entropie in geschlossenen Systemen niemals abnehmen kann und bei Nichtgleichgewichtssystemen in der Zeit zunimmt. Die Begründung einer ausgezeichneten Zeitrichtung bei zeitumkehrinvarianten Bewegungsgleichungen ist ein schwieriges (in der Sicht mancher Physiker immer noch ungelöstes) Problem der Physik.

Kapitel 8

Oszillatoren

Die Gleichungen für ein oder mehrere harmonische Oszillatoren tauchen in unterschiedlichen Formen in nahezu allen Bereichen der Physik auf. Sie bilden oft den Ausgangspunkt für eine störungstheoretische Behandlung komplexer Probleme, die sich nicht in geschlossener Form lösen lassen. Etwas vereinfacht könnte man sagen: Eine der häufigsten Aufgaben des theoretischen Physikers besteht darin, ein konkretes Problem (zumindest näherungsweise) auf den harmonischen Oszillator bzw. ein System aus harmonischen Oszillatoren zurückzuführen.

8.1 Der freie harmonische Oszillator

8.1.1 Die Bewegungsgleichung und die „geratenen“ Lösungen

Das einfachste Beispiel eines freien harmonischen Oszillators ist die Feder. Das Federgesetz (eine spezielle Form des Hooke'schen Gesetzes) besagt, dass die Kraft F zum Spannen einer Feder proportional zur Auslenkung x der Feder aus ihrer Ruhelage ist, und dass die Richtung dieser Kraft zur Ruhelage (der Auslenkung x entgegengesetzt) zeigt:

$$F = -kx. \quad (8.1)$$

Wir behandeln zunächst die 1-dimensionale Bewegung, eine Verallgemeinerung auf mehrere Freiheitsgrade erfolgt später.

Mit der Federkraft lautet das 2. Newton'sche Gesetz:

$$m\ddot{x}(t) = -kx(t), \quad (8.2)$$

bzw.

$$\ddot{x}(t) = -\omega^2 x(t) \quad \text{mit} \quad \omega^2 = \frac{k}{m}. \quad (8.3)$$

Abgesehen davon, dass es sich um eine gewöhnliche Differentialgleichung zweiter Ordnung handelt, ist die Bewegungsgleichung des harmonischen Oszillators eine *lineare, homogene* Gleichung. „Linear“ bedeutet dabei, dass die gesuchte Funktion $x(t)$ nur linear, also in erster Potenz, auftritt. „Homogen“ bedeutet zusätzlich, dass es keinen additiven Term in der Gleichung gibt, der die gesuchte Funktion $x(t)$ nicht enthält. Ein solcher Term würde einer äußeren Kraftquelle entsprechen. Eine homogene Differentialgleichung enthält somit keine äußeren Quellterme.

Die Linearität der Differentialgleichung ist das charakteristische Merkmal des harmonischen Oszillators. Gleichgültig ob wir mehrere unabhängige oder harmonisch gekoppelte Oszillatoren betrachten, immer sind die zugehörigen Bewegungsgleichungen linear.

Eine besondere Eigenschaft linearer, homogener Differentialgleichungen ist die Möglichkeit der *Superposition* von Lösungen. Das bedeutet im vorliegenden Fall: Seien $x_1(t)$ und $x_2(t)$ zwei beliebige spezielle Lösungen der Differentialgleichung, dann ist

$$x(t) = ax_1(t) + bx_2(t) \quad (8.4)$$

für beliebige Koeffizienten a und b ebenfalls eine Lösung. (Diese Eigenschaft kann man auch als Definition einer „linearen“ Gleichung ansehen.) Andererseits erwarten wir für eine gewöhnliche Differentialgleichung 2. Ordnung pro Freiheitsgrad zwei freie Parameter, welche die Anfangsbedingungen (beispielsweise Ort und Geschwindigkeit) festlegen. Da a und b zwei solche Parameter sind, ist $x(t)$ aus Gl. 8.4 schon die allgemeinste Lösung der Bewegungsgleichung, vorausgesetzt, die beiden speziellen Lösungen sind unabhängig. Das bedeutet, die Funktion $x_1(t)$ ist nicht proportional zu $x_2(t)$, oder anders ausgedrückt, es gibt kein α , sodass :

$$x_1(t) = \alpha x_2(t) \quad \text{für alle } t.$$

Wir werden in den folgenden Abschnitten Verfahren angeben, wie man für die Differentialgleichung 8.3 Lösungen finden kann. Zunächst wollen wir die Lösungen „erraten“ und ihre Eigenschaften untersuchen.

Gleichung 8.3 besagt, dass die zweite Ableitung der gesuchten Funktion $x(t)$ proportional zu der Funktion selber sein soll. Aus der Kenntnis der trigonometrischen Funktionen ist bekannt, dass dies für die Sinus- und Kosinus-Funktion der Fall ist:

$$\frac{d^2 \sin \omega t}{dt^2} = -\omega^2 \sin \omega t \quad \text{und} \quad \frac{d^2 \cos \omega t}{dt^2} = -\omega^2 \cos \omega t. \quad (8.5)$$

Damit erhalten wir zwei unabhängig Lösungen der Bewegungsgleichung 8.3:

$$x_1(t) = \cos \omega t \quad \text{und} \quad x_2(t) = \sin \omega t. \quad (8.6)$$

Die allgemeinste Lösung der Bewegungsgleichung ergibt sich zu:

$$x(t) = a \cos \omega t + b \sin \omega t. \quad (8.7)$$

a und b sollen nun durch die Anfangsbedingungen $x(0)$ und $v(0)$ ausgedrückt werden. Zunächst setzen wir in Gl. 8.7 die Zeit $t = 0$ und erhalten:

$$x(0) = a.$$

Nun berechnen wir die Geschwindigkeit

$$v(t) = \dot{x}(t) = -a\omega \sin \omega t + b\omega \cos \omega t.$$

Für $t = 0$ folgt daher:

$$v(0) = b\omega.$$

Ausgedrückt durch die Anfangsbedingungen $x(0)$ und $v(0)$ lautet die allgemeine Lösung:

$$x(t) = x(0) \cos \omega t + \frac{v(0)}{\omega} \sin \omega t. \quad (8.8)$$

Manchmal gibt man die allgemeine Lösung der Bewegungsgleichung 8.3 auch in folgender Form an:

$$x(t) = A \sin(\omega t - \varphi). \quad (8.9)$$

Die beiden freien Konstanten sind nun A und φ . Man könnte zunächst meinen, φ sei eine weitere freie Konstante, unabhängig von a und b , und wir hätten nun drei linear unabhängige Lösungen gefunden, was für eine gewöhnliche Differentialgleichung 2. Ordnung nicht der Fall sein sollte. Tatsächlich lässt sich diese Lösung (und ebenso die Lösung $x(t) = B \cos(\omega t - \varphi)$) in die andere Lösung überführen. Dazu nutzen wir das Additionstheorem der Sinus-Funktion,

$$\sin(\alpha + \beta) = \sin \alpha \cos \beta + \cos \alpha \sin \beta,$$

aus:

$$x(t) = A \left(\sin \omega t \cos \varphi - \cos \omega t \sin \varphi \right) = (-A \sin \varphi) \cos \omega t + (A \cos \varphi) \sin \omega t.$$

Die Beziehungen zu den ursprünglichen Koeffizienten sind daher:

$$a = -A \sin \varphi \quad b = A \cos \varphi.$$

Ähnliche Beziehungen erhält man für die Lösung $x(t) = B \cos(\omega t - \varphi)$.

8.1.2 Lösung durch den Exponentialansatz

Der folgende Ansatz mag zunächst willkürlich erscheinen, er findet aber seine Begründung in der Tatsache, dass es sich um eine lineare Differentialgleichung handelt, die „invariant unter Zeittranslationen“ ist, d.h., mit $x(t)$ muss auch $x(t+t_0)$ für jedes beliebige t_0 eine Lösung sein. Der folgende Ansatz berücksichtigt diese Bedingung in besonderem Maße. Die Auswertung soll hier jedoch nicht weiter verfolgt werden.

Die Lösung einer Gleichung (insbesondere einer Differentialgleichung) durch einen *Ansatz* besteht darin, „versuchsweise“ eine allgemeine Form für die Lösung anzusetzen, in die Gleichung einzusetzen, und anschließend konkrete Bedingungen für die freien Parameter in diesem Ansatz zu erhalten. Der so genannte Exponentialansatz besteht darin, für $x(t)$ eine Funktion der folgenden Form anzunehmen:

$$x(t) = e^{i\alpha t}. \quad (8.10)$$

Es mag zunächst erstaunen, eine komplexe Funktion als Ansatz zu wählen, doch der Vorteil dieses Verfahrens besteht darin, dass der Imaginärteil und der Realteil unabhängige Lösungen der linearen Differentialgleichung sein müssen, wenn die komplexe Funktion eine Lösung ist. Wir bilden die erste und zweite Ableitung von (8.10),

$$\dot{x}(t) = i\alpha e^{i\alpha t} \quad \ddot{x}(t) = -\alpha^2 e^{i\alpha t},$$

und finden, dass diese Funktion folgende Gleichungen erfüllt:

$$\dot{x}(t) = i\alpha x(t) \quad \text{und} \quad \ddot{x}(t) = -\alpha^2 x(t).$$

Der Vergleich der zweiten Bedingung mit Gl. 8.3 führt auf die Bedingung:

$$\alpha^2 = \omega^2 \quad \text{bzw.} \quad \alpha = \pm\omega.$$

Sowohl

$$x_1(t) = e^{i\omega t}$$

als auch

$$x_2(t) = e^{-i\omega t}$$

sind Lösungen der Differentialgleichung. Eine allgemeine Lösung kann also in der Form:

$$x(t) = ae^{i\omega t} + be^{-i\omega t}$$

geschrieben werden. Da wir komplexe Funktionen als Lösungen zulassen, können auch die Koeffizienten a und b komplex sein. Doch unabhängig davon, wie wir diese Koeffizienten wählen, wenn wir die Funktionen in ihren Real- und Imaginärteil aufspalten, erhalten wir immer nur Linearkombinationen von Sinus- und Kosinus-Funktionen mit demselben Argument ωt . Wir haben also nicht mehr gefunden, als vorher auch. Der Vorteil ist aber, dass wir bereits mit der einen Funktion

$$x_1(t) = e^{i\omega t} = \cos \omega t + i \sin \omega t$$

beide linear unabhängigen Lösungen der Differentialgleichung gefunden haben, wenn wir diese komplexe Lösung in ihren Real- und Imaginärteil aufspalten und berücksichtigen, dass Real- und Imaginärteil getrennt Lösungen der Gleichung sind.

8.1.3 Lösung durch Integration der Erhaltungsgröße

Da wir die Federkraft (Gl. 8.1) als Ableitung (die 1-dimensionale Form des Gradienten) eines Potentials schreiben können,

$$F = -\frac{dU(x)}{dx} \quad \text{mit} \quad U(x) = \frac{1}{2}kx^2,$$

ist die Energie:

$$E = \frac{1}{2}mv(t)^2 + \frac{1}{2}kx(t)^2 \tag{8.11}$$

eine Erhaltungsgröße der Bewegungsgleichung. Wir können uns explizit davon überzeugen, indem wir die bekannte Lösung (8.8) einsetzen, und die Bedingung $\omega^2 = k/m$ sowie die trigonometrische Identität $\sin^2 \omega t + \cos^2 \omega t = 1$ ausnutzen:

$$\begin{aligned} E &= \frac{1}{2}m \left(-x(0)\omega \sin \omega t + v(0) \cos \omega t \right)^2 + \frac{1}{2}k \left(x(0) \cos \omega t + \frac{v(0)}{\omega} \sin \omega t \right)^2 \\ &= \frac{1}{2}mv(0)^2 + \frac{1}{2}kx(0)^2. \end{aligned}$$

Die Anfangsbedingungen legen die Energie also fest und es gibt keine Zeitabhängigkeit mehr.

Doch wenn die Energie eine Konstante ist, können wir dann nicht direkt aus der Energie die Bahnkurve berechnen? Für die Bewegungsgleichung, die wir im 1-dimensionalen Fall haben, ist dies tatsächlich möglich. Bei einem 3-dimensionalen

System müssen wir neben der Energieerhaltung noch die Erhaltung weiterer Größen ausnutzen.

Wir betrachten die Energie E als Konstante und lösen Gleichung 8.11 nach $v(t) = \dot{x}(t)$ auf:

$$\dot{x}(t) = \sqrt{\frac{2E}{m} - \omega^2 x(t)^2}. \quad (8.12)$$

Diese Gleichung erscheint zwar zunächst komplizierter, doch es handelt sich nun um eine gewöhnliche (allerdings nicht mehr lineare) Differentialgleichung 1. Ordnung, die man allgemein durch Integration lösen kann. Dies soll nun gezeigt werden. Die Umformungen folgen dabei den typischen Überlegungen in der Physik, wobei die mathematische Korrektheit etwas zu kurz kommt, was der Richtigkeit des Ergebnisses jedoch keinen Abbruch tut.

Wir schreiben $\dot{x}(t)$ als Differentialquotient und bringen alle Terme, die x enthalten, auf eine Seite und den Term mit dt auf die andere Seite und bilden auf beiden Seiten das Integral:

$$\begin{aligned} \frac{dx}{dt} &= \sqrt{\frac{2E}{m} - \omega^2 x^2} \\ \Rightarrow \frac{dx}{\sqrt{\frac{2E}{m} - \omega^2 x^2}} &= dt \\ \Rightarrow \int_{x(0)}^{x(t)} \frac{dx}{\sqrt{\frac{2E}{m} - \omega^2 x^2}} &= \int_0^t dt. \end{aligned} \quad (8.13)$$

A: In der Mathematik wird dieses Verfahren als *Quadratur* bezeichnet. Allgemein gilt: Sei eine Differentialgleichung 1. Ordnung vom Typ

$$\frac{dy}{dx} = X(x) \cdot Y(y)$$

gegeben, so lässt sich die Lösung in der Form:

$$\int^y \frac{dy}{Y(y)} = \int^x X(x) dx + c$$

schreiben, wobei c eine Integrationskonstante ist. Differenziert man beide Seiten (implizit) nach x , so erhält man wieder obige Differentialgleichung.

Während die rechte Seite von Gl. 8.13 t ergibt, entnimmt man die Lösung des Integrals auf der linken Seite am besten einer Integraltafel, wo man folgende

Formel findet:

$$\int \frac{dx}{\sqrt{a - cx^2}} = \frac{-1}{\sqrt{c}} \arcsin \frac{-2cx}{\sqrt{4ac}}.$$

Setzen wir $a = 2E/m$ und $c = \omega^2$ ein und definieren noch

$$\varphi_0 = \arcsin \left(-\frac{\omega \sqrt{m} x(0)}{\sqrt{2E}} \right),$$

so erhalten wir:

$$-\frac{1}{\omega} \arcsin \left(\frac{-\omega \sqrt{m} x(t)}{\sqrt{2E}} \right) + \frac{1}{\omega} \varphi_0 = t,$$

oder, aufgelöst nach $x(t)$:

$$x(t) = \sqrt{\frac{2E}{m}} \omega \sin(\omega t - \varphi_0).$$

Diese Lösung gleicht zwar Gl. 8.9, allerdings sind die Integrationskonstanten unterschiedlich. Wir haben nun die Energie E und (implizit über die Phase φ_0) den Ort $x(0)$ als Integrationskonstanten der Bewegung gewählt, während wir früher die Anfangsbedingungen $x(0)$ und $v(0)$ betrachtet haben. Unter Ausnutzung von

$$\frac{2E}{m} = v(0)^2 + \omega^2 x(0)^2$$

und der folgenden trigonometrischen Beziehung,

$$\arcsin x = \arccos \sqrt{1 - x^2},$$

lässt sich jedoch nach einigen Rechenschritten zeigen, dass die obige Lösung mit Gl. 8.8 übereinstimmt.

Es mag zunächst als unvergleichbar größerer Aufwand erscheinen, die Lösung über die Integration der Energie zu suchen und für die (lineare) Gleichung des harmonischen Oszillators ist das auch sicherlich richtig, doch der Vorteil dieses Verfahrens liegt darin, dass es auch für Kräfte angewandt werden kann, die nicht linear in der Auslenkung sind.

8.2 Der gedämpfte harmonische Oszillator

8.2.1 Bewegungsgleichung und Exponentialansatz

Als weitere Anwendung für den Exponentialansatz betrachten wir den *gedämpften harmonischen Oszillator*. Neben der reinen Federkraft proportional zur Auslen-

kung der Feder,

$$F(x) = -kx,$$

(k ist die Federkonstante) kommt nun noch eine *Widerstandskraft* proportional zur Geschwindigkeit des Körpers hinzu:

$$F_R(v) = -\mu v. \quad (8.14)$$

Die Bewegungsgleichung lautet nun:

$$m\ddot{x}(t) = -\mu\dot{x}(t) - kx(t). \quad (8.15)$$

Hier besteht somit zu jedem Zeitpunkt eine Beziehung zwischen der Bahnkurve $x(t)$ und ihrer ersten und zweiten Ableitung nach der Zeit. In diesem Fall fällt das „Raten“ der Lösung schon schwerer, daher suchen wir die speziellen Lösungen mithilfe des Exponentialansatzes:

$$x(t) = e^{i\alpha t}.$$

Dies führt auf folgende Gleichung:

$$-m\alpha^2 x(t) = -i\mu\alpha x(t) - kx(t).$$

Da wir nicht an der Lösung interessiert sind, für die $x(t)$ identisch null ist, können wir durch $x(t)$ dividieren und erhalten:

$$\alpha^2 - i\frac{\mu}{m}\alpha - \frac{k}{m} = 0. \quad (8.16)$$

Offensichtlich muss α komplex sein, um diese Gleichung lösen zu können. Aus der allgemeinen Formel für quadratische Gleichungen ergibt sich:

$$\alpha = i\frac{\mu}{2m} \pm \sqrt{\frac{k}{m} - \frac{\mu^2}{4m^2}}. \quad (8.17)$$

Wir treffen nun eine Fallunterscheidung, je nachdem ob der Term unter der Wurzel positiv, negativ oder null ist.

8.2.2 Die gedämpfte Schwingung: $4km > \mu^2$

Wir definieren zunächst die neuen Parameter:

$$\gamma = \frac{\mu}{2m} \quad \text{und} \quad \omega = +\sqrt{\frac{k}{m} - \frac{\mu^2}{4m^2}} = \sqrt{\omega_0^2 - \gamma^2}$$

($\omega_0 = \sqrt{k/m}$ ist die Winkelgeschwindigkeit der freien, ungedämpften Schwingung), sodass

$$\alpha = i\gamma \pm \omega.$$

Eingesetzt in den Exponentialansatz erhalten wir somit zwei spezielle Lösungen:

$$\tilde{x}_1(t) = e^{-\gamma t} e^{i\omega t} \quad \text{und} \quad \tilde{x}_2(t) = e^{-\gamma t} e^{-i\omega t}. \quad (8.18)$$

Spalten wir diese Lösungen in Real- und Imaginärteil auf, so erhalten wir die beiden linear unabhängigen Lösungen:

$$x_1(t) = e^{-\gamma t} \cos \omega t \quad \text{und} \quad x_2(t) = e^{-\gamma t} \sin \omega t. \quad (8.19)$$

Die allgemeine Lösung der Differentialgleichung lautet somit:

$$x(t) = e^{-\gamma t} (a \cos \omega t + b \sin \omega t) \quad (8.20)$$

oder, äquivalent:

$$x(t) = A e^{-\gamma t} \cos(\omega t + \varphi).$$

Die Lösung besteht aus einem Produkt einer abfallenden Exponentialfunktion, welche die Dämpfung der Schwingung beschreibt, und einer reinen Schwingung. Die freien Integrationskonstanten lassen sich wieder mit den Anfangsbedingungen in Beziehung setzen:

$$x(0) = a \quad \text{und} \quad v(0) = -a\gamma + \omega b, \quad (8.21)$$

was auf folgende Darstellung der Lösung führt:

$$x(t) = e^{-\gamma t} \left(x(0) \cos \omega t + \frac{x(0)\gamma + v(0)}{\omega} \sin \omega t \right). \quad (8.22)$$

8.2.3 Die reine Dämpfung: $4km < \mu^2$

In diesem Fall ist α (Gl. 8.17) imaginär,

$$\alpha = i \left(\frac{\mu}{2m} \pm \sqrt{\frac{\mu^2}{4m^2} - \frac{k}{m}} \right) = i(\gamma \pm \sqrt{\gamma^2 - \omega_0^2}), \quad (8.23)$$

und wir erhalten eine reine Dämpfung:

$$x(t) = a e^{-(\gamma + \sqrt{\gamma^2 - \omega_0^2})t} + b e^{-(\gamma - \sqrt{\gamma^2 - \omega_0^2})t}. \quad (8.24)$$

Man spricht in diesem Fall auch von einer *aperiodischen Bewegung*.

Man beachte, dass $\gamma \pm \sqrt{\gamma^2 - \omega_0^2}$ immer positiv ist (sofern $\gamma > 0$), sodass beide Anteile der Lösung eine gedämpfte, d.h. abklingende, Bewegung beschreiben.

8.2.4 Der aperiodische Grenzfall: $4km = \mu^2$

Für den Fall $4km = \mu^2$ bzw. $\gamma^2 = \omega_0^2$ verschwindet der Wurzelterm und wir erhalten:

$$\alpha = i \frac{\mu}{2m}. \quad (8.25)$$

Die zugehörige Lösung entspricht ebenfalls einer reinen Dämpfung:

$$x_1(t) = e^{-\gamma t}. \quad (8.26)$$

Allerdings haben wir nur eine Lösung gefunden. Da es sich bei der Bewegungsgleichung um eine gewöhnliche Differentialgleichung 2. Ordnung handelt wissen wir, dass noch eine zweite unabhängige Lösung existieren muss. Der Exponentialansatz liefert für den aperiodischen Grenzfall offensichtlich nicht alle Lösungen.

Um die Zusammenhänge etwas besser zu verstehen, betrachten wir nochmals die freie Bewegung ($\gamma = k = 0$). Der Exponentialansatz,

$$x(t) = e^{i\alpha t},$$

eingesetzt in die Bewegungsgleichung,

$$m\ddot{x}(t) = 0,$$

führt auf die Bedingung

$$\alpha^2 = 0,$$

mit der einzigen Lösung $\alpha = 0$. Wir finden daher als Lösung:

$$x(t) = x(0) = \text{konst.}$$

Die zweite bereits bekannte Lösung,

$$x(t) = v(0)t,$$

erhalten wir nicht aus einem Exponentialansatz. Wir können jedoch einen kleinen „Trick“ anwenden. Betrachten wir zunächst die Lösung des ungedämpften harmonischen Oszillators zu nicht-verschwindendem ω als Funktion der Anfangswerte $x(0)$ und $v(0)$ (vgl. Gl. 8.8):

$$x(t) = x(0) \cos \omega t + \frac{v(0)}{\omega} \sin \omega t. \quad (8.27)$$

Wir nehmen nun den Grenzfall $\omega \rightarrow 0$, wobei

$$\lim_{\omega \rightarrow 0} \frac{\sin \omega t}{\omega} = \frac{\omega t}{\omega} + O(\omega^2) = t,$$

und finden die allgemeine Lösung:

$$x(t) = x(0) + v(0)t.$$

Der Trick besteht darin, nicht von der Lösung mit beliebigen freien Integrationskonstanten auszugehen, sondern den Grenzfall $\omega \rightarrow 0$ erst zu betrachten, nachdem die Integrationskonstanten durch physikalische Größen, beispielsweise die Anfangsbedingungen $x(0)$ und $v(0)$, ersetzt wurden.

Wir können nun den gleichen Trick bei der gedämpften Schwingung anwenden. Ausgehend von der allgemeinen Lösung 8.22 finden wir im Grenzfall $\omega \rightarrow 0$ die Lösung:

$$x(t) = e^{-\gamma t} (x(0) + (x(0)\gamma + v(0))t). \quad (8.28)$$

Die zweite spezielle Lösung neben (8.26) lautet daher:

$$x_2(t) = te^{-\gamma t}.$$

Diese Lösung hätten wir durch einen reinen Exponentialansatz nicht gefunden.