

Reinhard Weiß

Mathematik

Einige ausgewählte Themen für das Physikstudium

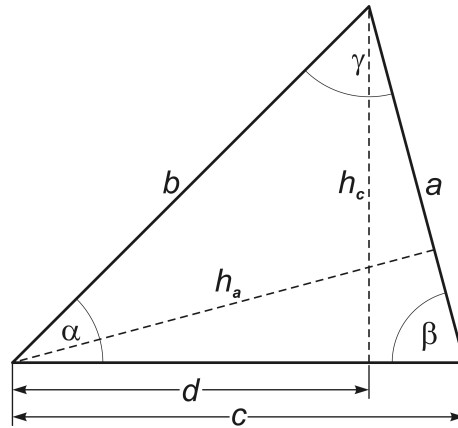
02.11.2025

Inhaltsverzeichnis

1	Sinussatz, Kosinussatz, Additionstheoreme	5
2	Differentialbegriff und Infinitesimalrechnung	8
2.1	Differentialrechnung	8
2.2	Integralrechnung	9
2.3	Beispiele	11
3	Totales Differential und thermodynamisches Potential, totale Ableitung	13
4	Lösung der Integrale $\int_0^{\infty} x^q \cdot e^{-a x^2} dx$, $q \in \mathbb{N}$, $a \in \mathbb{R}^{>0}$	17
4.1	Für ungerade q	17
4.2	Für $q = 0$	19
4.3	Für gerade q	19
5	Polar-, Zylinder- und Kugelkoordinaten auf einen Blick	22
5.1	Polarkoordinaten (r, φ)	22
5.2	Zylinderkoordinaten (r, φ, z)	24
5.3	Kugelkoordinaten (r, ϑ, φ)	26
6	Linien-, Flächen-, Volumenelement	28
6.1	Linienelement	28
6.1.1	Linienelement in kartesischen Koordinaten	28
6.1.2	Linienelement in Polarkoordinaten	29
6.2	Flächenelement	32
6.3	Volumenelement	33
6.4	Gelegentlich hilfreiche Ergänzungen	35
7	Differentialoperatoren in krummlinig-orthogonalen Koordinatensystemen	37
7.1	Gradient	37
7.2	Divergenz	38
7.3	Laplace-Operator	40
7.4	Rotation	41
8	Green'sche Identitäten	43
9	Taylor-Entwicklung eines skalaren Feldes	44
10	Dirac'sche delta-Funktion	48
10.1	Definition der δ -Funktion	48
10.2	Faltungsintegral mit der δ -Funktion	50
10.3	Eigenschaften der δ -Funktion – Rechenregeln	52
10.4	Fourier-Transformation und δ -Funktion	61
10.5	Die δ -Funktion in Kugelkoordinaten	63
10.6	Verallgemeinerung für das Ersetzen der kartesischen Koordinaten in der δ -Funktion (Koordinaten-Transformation)	64

10.7	Herleitung von $\Delta \frac{1}{ \vec{r}-\vec{r}' } = -4\pi \delta(\vec{r}-\vec{r}')$	65
10.7.1	Betrachtungen für $\vec{r} \neq \vec{r}'$	66
10.7.2	Betrachtungen unter Einschluss von $\vec{r} = \vec{r}'$	72
10.8	Herleitung von $(\Delta + k^2) \frac{\exp(\pm ikr)}{r} = -4\pi \delta(\vec{r}-\vec{r}')$	74
10.9	Die δ -Funktion in der Elektrostatik	76
11	Das Wichtigste zu reellen Matrizen	77
11.1	Assoziativgesetz für die Matrizenmultiplikation	77
11.2	Die Transponierte einer Matrix	78
11.3	Multiplikation transponierter Matrizen	78
11.4	Quadratische Matrizen	79
11.4.1	Invertieren von Matrizen	81
11.4.2	Orthogonale Matrizen – Drehmatrizen	82
12	Rechnen mit komplexen Vektoren und Matrizen bzw. Operatoren	84
12.1	Rechenregeln	84
12.2	Veranschaulichung des komplexen Standardskalarprodukts und der Multiplikation komplexer Vektoren mit komplexen Matrizen	89
12.3	Matrixdarstellung von Operatoren	92
13	Eigenwertgleichung einer 2-reihigen reellen Matrix – Verallgemeinerungen für n-reihige Matrizen	95
14	Hilbertraum, Hermitesche Matrizen (Operatoren) und ihre Eigenschaften	100
15	Unitäre Matrizen (Operatoren)	103
15.1	Eigenschaften unitärer Matrizen (Operatoren)	103
15.2	Unitäre Transformation	104
15.3	Diagonalisierung von Matrizen (Operatoren)	106
16	Spur einer Matrix	109
17	Kombinatorik, Verteilungsmöglichkeiten und ihre Wahrscheinlich- keit	111
18	Gesetz der großen Zahlen – Binomialverteilung	120
19	Kapitalwachstum bei stetiger Reinvestition der Zinsen	137

1 Sinussatz, Kosinussatz, Additionstheoreme



Herleitung Sinussatz

$$\left. \begin{array}{l} \sin \alpha = \frac{h_c}{b} \\ \sin \beta = \frac{h_c}{a} \end{array} \right\} \Rightarrow h_c = \underbrace{b \cdot \sin \alpha = a \cdot \sin \beta} \Leftrightarrow \frac{\sin \alpha}{a} = \frac{\sin \beta}{b} ,$$

$$\left. \begin{array}{l} \sin \beta = \frac{h_a}{c} \\ \sin \gamma = \frac{h_a}{b} \end{array} \right\} \Rightarrow h_a = \underbrace{c \cdot \sin \beta = b \cdot \sin \gamma} \Leftrightarrow \frac{\sin \beta}{b} = \frac{\sin \gamma}{c} ,$$

$$\boxed{\frac{\sin \alpha}{a} = \frac{\sin \beta}{b} = \frac{\sin \gamma}{c}} . \quad \text{q.e.d.}$$

Herleitung Kosinussatz

$$\cos \alpha = \frac{d}{b} \Leftrightarrow d = b \cdot \cos \alpha ,$$

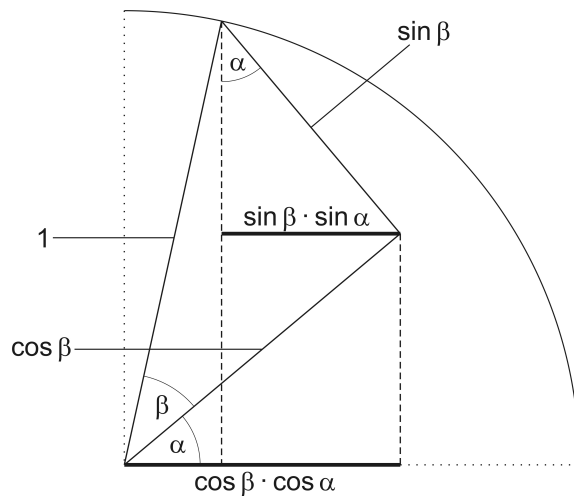
$$b^2 = h_c^2 + d^2 \Leftrightarrow h_c^2 = b^2 - d^2 ,$$

$$\begin{aligned} a^2 &= (c - d)^2 + h_c^2 \\ &= c^2 - 2cd + d^2 + b^2 - d^2 \\ &= c^2 - 2c \cdot b \cos \alpha + b^2 , \end{aligned}$$

$$\boxed{a^2 = b^2 + c^2 - 2bc \cdot \cos \alpha} . \quad \text{q.e.d.}$$

Die Herleitungen zur Berechnung der Seiten b und c sind völlig analog.

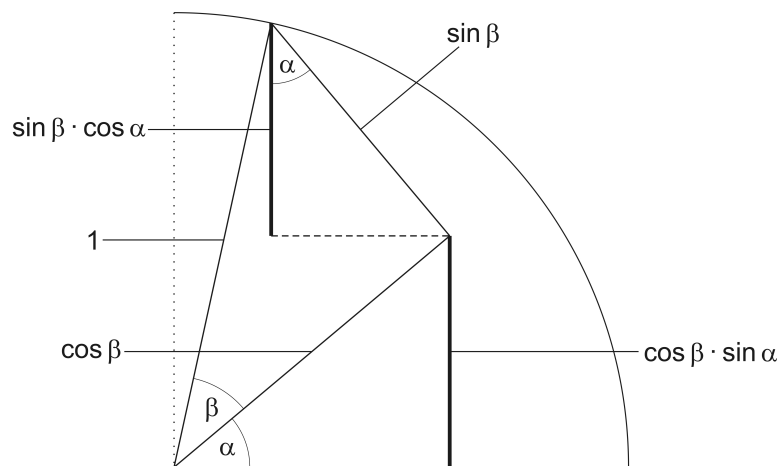
Veranschaulichung der Additionstheoreme im ersten Quadranten des Einheitskreises



$$\cos(\alpha + \beta) = \cos \alpha \cdot \cos \beta - \sin \alpha \cdot \sin \beta$$

und mit $\cos(-\gamma) = \cos \gamma$ sowie mit $\sin(-\gamma) = -\sin \gamma$

$$\cos(\alpha - \beta) = \cos \alpha \cdot \cos \beta + \sin \alpha \cdot \sin \beta .$$



$$\sin(\alpha + \beta) = \sin \alpha \cdot \cos \beta + \sin \beta \cdot \cos \alpha$$

und mit $\cos(-\gamma) = \cos \gamma$ sowie mit $\sin(-\gamma) = -\sin \gamma$

$$\sin(\alpha - \beta) = \sin \alpha \cdot \cos \beta - \sin \beta \cdot \cos \alpha .$$

2 Differentialbegriff und Infinitesimalrechnung

Siehe auch:

Springer-Lehrbuch von Klaus Weltner, Mathematik für Physiker 1, 12. Aufl., Springer-Verlag, Berlin, Heidelberg, New York, 2001, Kapitel 5 Differentialrechnung und Kapitel 6 Integralrechnung.

Ohne Anspruch auf mathematische Vollständigkeit wollen wir in diesem Kapitel die Begriffe Differential, Differenzenquotient und Differentialquotient anhand der Abbildung 2.1 und anschließend die Begriffe bestimmtes und unbestimmtes Integral skizzieren. Dabei wird deutlich werden, dass Differenziale nicht unbedingt infinitesimal kleine Größen sein müssen, auch wenn diese Betrachtungsweise sehr praktisch sein kann und deshalb in der Physik üblich ist gemäß $dx = \lim_{x \rightarrow x_0} (x - x_0) \hat{=} x - x_0 = \Delta x \rightarrow 0$.

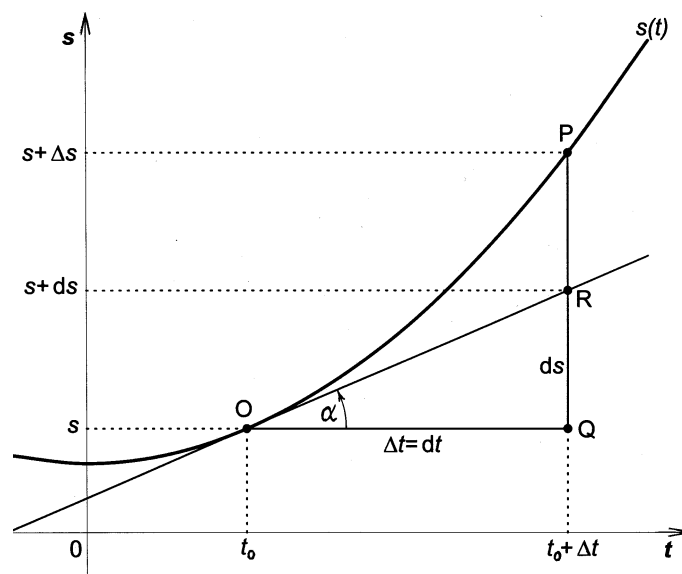


Abb. 2.1 Geometrische Veranschaulichung der Differenziale ds und dt sowie des Differentialquotienten $ds/dt = \tan \alpha$ im Punkt $s(t_0)$ der Funktion $s(t)$.

2.1 Differentialrechnung

Gegeben sei eine glatte Funktion $s(t)$, also eine stetige, von t abhängige Funktion ohne „Knick“. Z. B. kann die unabhängige Variable t die Zeit und die abhängige Variable s die in der Zeit zurückgelegte Weglänge sein. Während die Zeit t ausgehend von t_0 um Δt (Strecke $\overline{OQ}$) fortschreitet, ändert sich die Funktion $s(t)$ um $\Delta s = s(t_0 + \Delta t) - s(t_0)$ (Strecke $\overline{QP}$). Jetzt legen wir im Punkt $s(t_0)$ (Punkt O) eine Tangente an die Kurve $s(t)$. Die Tangente schneidet die Strecke $\overline{OP}$ im Punkt R . Die Strecke $\overline{QR}$ ist die geometrische Darstellung des abhängigen Differentials ds .

Merke:

Das **abhängige Differential** ds (das Differential der Funktion) ist die lineare Näherung für die Änderung Δs . Wir haben dabei das **unabhängige Differential** $dt = \Delta t$ (das Differential des Arguments, Strecke $\overline{OQ}$) zugrunde gelegt.

Die Näherung ist um so besser, je mehr sich der Punkt Q dem Punkt O nähert bzw. je kleiner wir das unabhängige Differential $dt = \Delta t$ wählen. Bei linearen Funktionen $s(t)$ ist die Tangente an jeden Punkt der Kurve identisch mit der Kurve selbst, so dass in diesem Fall für das abhängige Differential $ds = \Delta s$ gilt.

Der Quotient aus den Differentialen ds und dt bezüglich $s(t_0)$, also der *Differentialquotient* ds/dt an der Stelle t_0 , ist geometrisch die Steigung $\tan \alpha$ der Tangente an den Kurvenpunkt $s(t_0)$ (s. Abbildung 2.1). Allgemein gilt, je kleiner wir im *Differenzenquotienten* $\Delta s/\Delta t$ das Intervall Δt wählen, desto mehr nähert sich dieser dem Differentialquotienten an, so dass wir für den Differentialquotienten von $s(t)$ an der Stelle $t = t_0$ folgendes schreiben können:

$$\left. \frac{ds}{dt} \right|_{t_0} = \lim_{\Delta t \rightarrow 0} \left. \frac{\Delta s}{\Delta t} \right|_{t_0} = \lim_{\Delta t \rightarrow 0} \frac{s(t_0 + \Delta t) - s(t_0)}{\Delta t} = s'(t_0) . \quad (2.1)$$

Diesen Grenzwert können wir für beliebige t bilden, so dass wir eine Funktion der Differentialquotienten von t erhalten. Diese Funktion heißt Ableitung

$$s'(t) = \frac{ds(t)}{dt} = \lim_{\Delta t \rightarrow 0} \frac{\Delta s}{\Delta t} = \lim_{\Delta t \rightarrow 0} \frac{s(t + \Delta t) - s(t)}{\Delta t} = v(t) \quad (2.2)$$

der Funktion $s(t)$. Die Ableitung $s'(t)$ der Funktion $s(t)$ liefert uns für jeden Zeitpunkt t die momentane Änderung des Weges s pro zugehöriger Änderung der Zeit t , also die Geschwindigkeit $s'(t) = v(t) = ds/dt$ und mit (2.2) schreiben wir für das abhängige Differential

$$ds = s'(t)dt = v(t)dt . \quad (2.3)$$

2.2 Integralrechnung

Wir wollen jetzt aus der Funktion $s'(t)$ die Änderung Δs in den Grenzen von t_0 und t berechnen. Dazu teilen wir Δt in N Abschnitte Δt_i , so dass $\sum_{i=1}^N \Delta t_i = \Delta t$. Wenn die Teilbereiche Δt_i klein im Vergleich zu Δt gewählt werden, erhält man mit je einem Funktionswert $s'(t_i)$ aus jedem Teilbereich Δt_i die Näherung

$$\sum_{i=1}^N s'(t_i) \cdot \Delta t_i = \sum_{i=1}^N \left. \frac{ds}{dt} \right|_{t_i} \cdot \Delta t_i = \sum_{i=1}^N v(t_i) \cdot \Delta t_i \approx \Delta s . \quad (2.4)$$

Je kleiner wir die Intervalle Δt_i wählen, desto größer wird ihre Anzahl N in den Grenzen von t_0 bis $t_0 + \Delta t$ und desto besser wird die Näherung. Im Grenzfall für $\Delta t_i \rightarrow 0$ bzw. $N \rightarrow \infty$ erhalten wir aus der Näherungsgleichung (2.4) die exakte Änderung Δs :

$$\begin{aligned} \lim_{\Delta t_i \rightarrow 0} \sum_{i=1}^N s'(t_i) \cdot \Delta t_i &= \lim_{\Delta t_i \rightarrow 0} \sum_{i=1}^N v(t_i) \cdot \Delta t_i = \\ &= \int_{t_0}^{t_0 + \Delta t} s'(t) dt = \int_{t_0}^{t_0 + \Delta t} v(t) dt = \int_{s(t_0)}^{s(t_0 + \Delta t)} ds = s(t_0 + \Delta t) - s(t_0) = \Delta s . \end{aligned} \quad (2.6)$$

Der Ausdruck

$$\int_{t_0}^{t_0+\Delta t} s'(t) dt = \Delta s \quad (2.7)$$

heißt **bestimmtes Integral**¹. Es hat nämlich wegen seiner Integrationsgrenzen einen eindeutig bestimmbar Wert.

Wie wir sehen, ist das Integralzeichen der Befehl zur Summation über eine unendliche Anzahl von *infinitesimal kleinen Differentialen* $ds = s'(t)dt = v(t)dt$. Geometrisch ist das bestimmte Integral (2.6) der Flächeninhalt A unter dem Graphen der Funktion $s'(t)$ in den Grenzen von t_0 und $t_0 + \Delta t$ (s. Abb. 2.2 und Abb. 2.3).

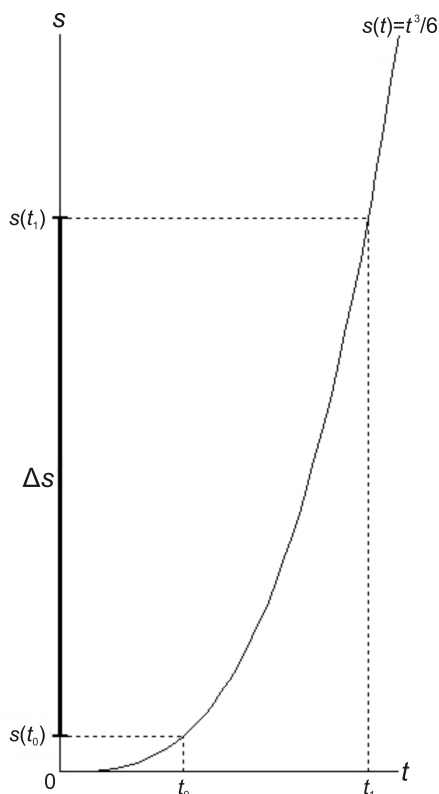


Abb. 2.2 Beispiel Weg-Zeit-Funktion

$$s(t) = \frac{1}{6} t^3.$$

Gemäß der Weg-Zeit-Funktion $s(t)$ ist $\Delta s = s(t_1) - s(t_0)$ der im Zeitintervall $t_1 - t_0$ zurückgelegte Weg. Der Anstieg des Graphen der Funktion $s(t)$ in Abhängigkeit von t wird beschrieben durch die Ableitung $s'(t) = \frac{ds(t)}{dt} = v(t)$.

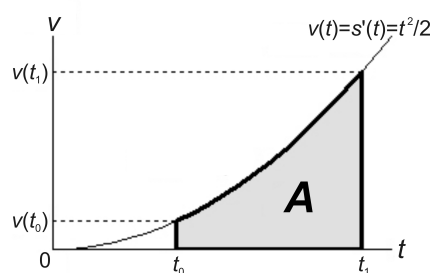


Abb. 2.3 Ableitung von $\frac{1}{6} t^3$ nach t :
 $s'(t) = \frac{1}{2} t^2$.

Für den Flächeninhalt A unter dem Graphen der Ableitungsfunktion $s'(t)$ von t_0 bis t_1 gilt

$$A = \Delta s = \int_{t_0}^{t_1} \frac{ds(t)}{dt} dt = \int_{t_0}^{t_1} s'(t) dt.$$

Das **unbestimmte Integral** über $s'(t)$ hat keine Integrationsgrenzen und liefert uns die Stammfunktion zu $s'(t)$:

$$\int s'(t) dt = s(t) + C. \quad (2.8)$$

¹Dieses bestimmte Integral als Beispiel wird gelesen: „Integral über $s'(t) dt$ in den Grenzen von t_0 bis $t_0 + \Delta t$ gleich Δs “ oder kurz: „Integral $s'(t) dt$ von t_0 bis $t_0 + \Delta t$ gleich Δs “.

Die Integrationskonstante C hängt von den Anfangsbedingungen ab. Bilden wir die Ableitung der Stammfunktion, so verschwindet die Integrationskonstante, weil die Änderung oder Ableitung einer Konstanten gleich Null ist. Wir erhalten deshalb aus (2.8) wieder

$$\frac{d}{dt}(s(t) + C) = \frac{ds(t)}{dt} + \frac{dC}{dt} = s'(t) + 0 = s'(t) . \quad (2.9)$$

Differentiation und Integration sind zueinander Umkehroperationen. Sinngemäß liefern uns die Differentiation einen „Quotienten aus Differenzen“ und die Integration eine „Summe von Produkten“.

2.3 Beispiele

Im Gegensatz zur Tatsache, dass die Integration über unendlich kleine Differentiale erfolgt, sind Differentiale allgemein keine unendlich kleinen Größen. Dessen ungeachtet ist es in der Physik üblich, Differentiale als unendlich kleine Änderungen von Größen anzusehen und mit ihnen zu rechnen, d. h., auf sie die Grundrechenarten anzuwenden. Das führt oft auf bequeme Weise zu exakten Ergebnissen, wenn man die Differenzierbarkeit bzw. die Integrabilität der betrachteten Funktionen voraussetzt. Dies werde für die folgenden veranschaulichenden Beispiele angenommen.

- Gegeben sei die Weg-Zeit-Funktion $x(t) = t^2$. Die Differentiationsregel für Potenzfunktionen ist

$$\frac{d}{dx}(a \cdot x^n + C) = a \cdot nx^{n-1} . \quad (2.10)$$

Dann ist die Geschwindigkeit-Zeit-Funktion

$$v(t) = \frac{dx}{dt} = \frac{d}{dt}t^2 = 2t \quad (2.11)$$

und das abhängige Differential $dx = v(t)dt = 2t dt$.

- Gegeben sei die verkettete Funktion $v \circ x = v(x(t))$ mit $v(x) = x/2$ und $x(t) = t^2$. Die Ableitung von $v(x(t)) = t^2/2$ nach t erhalten wir durch Anwendung der Kettenregel:

$$\frac{dv(x)}{dx} \cdot \frac{dx}{dt} = \frac{1}{2} \cdot 2t = \frac{dv(x(t))}{dt} = t . \quad (2.12)$$

Man kommt direkt zum Ergebnis, wenn man dx „herauskürzt“ und für v die Funktion $v(x(t))$ einsetzt.

- Für das Geschwindigkeitsquadrat im kartesischen Koordinatensystem gilt mit (2.2)

$$v^2 = v_x^2 + v_y^2 + v_z^2 \quad (2.13)$$

$$= \left(\lim_{\Delta t \rightarrow 0} \frac{\Delta x}{\Delta t} \right)^2 + \left(\lim_{\Delta t \rightarrow 0} \frac{\Delta y}{\Delta t} \right)^2 + \left(\lim_{\Delta t \rightarrow 0} \frac{\Delta z}{\Delta t} \right)^2 \quad (2.14)$$

$$= \lim_{\Delta t \rightarrow 0} \frac{(\Delta x)^2 + (\Delta y)^2 + (\Delta z)^2}{(\Delta t)^2} \quad (2.15)$$

$$v^2 = \frac{(dx)^2 + (dy)^2 + (dz)^2}{(dt)^2} . \quad (2.16)$$

Es ist nicht ganz korrekt aber bequem und deshalb üblich, bei Potenzen der Koordinatendifferentiale die Klammern zu vernachlässigen. So erhalten wir aus (2.16) für das Geschwindigkeitsquadrat

$$v^2 = \frac{dx^2 + dy^2 + dz^2}{dt^2} \quad (2.17)$$

und für das Differential des räumlichen Abstandsquadrats

$$dx^2 + dy^2 + dz^2 = v^2 dt^2 . \quad (2.18)$$

- Ein Körper werde ab dem Zeitpunkt $t = 0$ mit $a(t) = dv/dt = t$ beschleunigt. Bei $t = 0$ habe der Körper die Anfangsgeschwindigkeit $C = v_0 = 4$. Welche Geschwindigkeit hat er dann bei $t_1 = 2$ und bei $t_2 = 8$?

Die Integrationsregel für Potenzfunktionen ist

$$\int a \cdot x^n dx = a \cdot \frac{1}{n+1} x^{n+1} + C . \quad (2.19)$$

Aus dem Geschwindigkeitsdifferential $dv(t) = a(t)dt$ erhalten wir

$$v(t) = \int dv(t) = \int a(t)dt = \int t dt = \frac{1}{2}t^2 + v_0 = \frac{1}{2}t^2 + 4 . \quad (2.20)$$

Die gesuchten Geschwindigkeiten sind also $v(t_1) = v_1 = 6$ und $v(t_2) = v_2 = 36$. Der Körper hat also seine Geschwindigkeit in der Zeit von t_1 bis t_2 von 6 auf 36 erhöht. Diese Geschwindigkeitsänderung um 30 liefert uns das bestimmte Integral:

$$v_2 - v_1 = \int_{v_1}^{v_2} dv = \int_{t_1=2}^{t_2=8} t dt = \left. \frac{1}{2}t^2 \right|_{t_1=2}^{t_2=8} = 32 - 2 = 30 . \quad (2.21)$$

Die Integrationskonstante $v_0 = 4$ brauchen wir beim bestimmten Integral nicht zu berücksichtigen, denn sie fällt bei der Subtraktion $v_2 - v_1 = (32 + 4) - (2 + 4)$ heraus.

3 Totales Differential und thermodynamisches Potential, totale Ableitung

Mit dem Gradienten $\text{grad } f$ eines skalaren Feldes $f(\vec{r}) = f(x, y, z)$ in unmittelbarem Zusammenhang steht das vollständige oder totale Differential auf folgende Weise (totale Differenzierbarkeit von f vorausgesetzt):

$$\text{Gradient : } \frac{f(\vec{r} + d\vec{r}) - f(\vec{r})}{d\vec{r}} = \frac{df(\vec{r})}{d\vec{r}} = \text{grad } f(\vec{r}) := \vec{e}_x \frac{\partial f}{\partial x} + \vec{e}_y \frac{\partial f}{\partial y} + \vec{e}_z \frac{\partial f}{\partial z} \Rightarrow$$

$$\begin{aligned} \text{totales Differential : } df(\vec{r}) &= f(\vec{r} + d\vec{r}) - f(\vec{r}) = \frac{df(\vec{r})}{d\vec{r}} \cdot d\vec{r} = \text{grad } f(\vec{r}) \cdot d\vec{r} \\ &= \left(\vec{e}_x \frac{\partial f(x, y, z)}{\partial x} + \vec{e}_y \frac{\partial f(x, y, z)}{\partial y} + \vec{e}_z \frac{\partial f(x, y, z)}{\partial z} \right) \cdot (dx \vec{e}_x + dy \vec{e}_y + dz \vec{e}_z), \end{aligned}$$

$$\boxed{df(\vec{r}) = \text{grad } f(\vec{r}) \cdot d\vec{r} = \frac{\partial f}{\partial x} dx + \frac{\partial f}{\partial y} dy + \frac{\partial f}{\partial z} dz} \quad (3.1)$$

Die Funktion, von der das totale Differential gebildet wird, muss nicht von den drei Raumkoordinaten abhängen, sondern kann von beliebig vielen Koordinaten (Größen) x_i ($i = 1, 2, 3, \dots$) abhängen. Deshalb lautet die allgemeine Definition des totalen Differentials

$$df := \sum_i \frac{\partial f}{\partial x_i} dx_i.$$

Wenn in einem Punkt $\vec{r}_0$ alle partiellen Ableitungen einer Funktion $f(\vec{r})$ existieren und wenn dann dort auch noch alle partiellen Ableitungen stetig sind, dann ist f in $\vec{r}_0$ total differenzierbar. Geometrisch bedeutet das, dass sich die total differenzierbare Funktion lokal durch eine lineare Abbildung (Tangentialebene) approximieren lässt.¹

Die geometrische Bedeutung des totalen Differentials lässt sich für eine Funktion von zwei unabhängigen Variablen wie beispielsweise $f(\vec{r}) = f(x, y) = z$ graphisch veranschaulichen (siehe Abbildung 3.1).

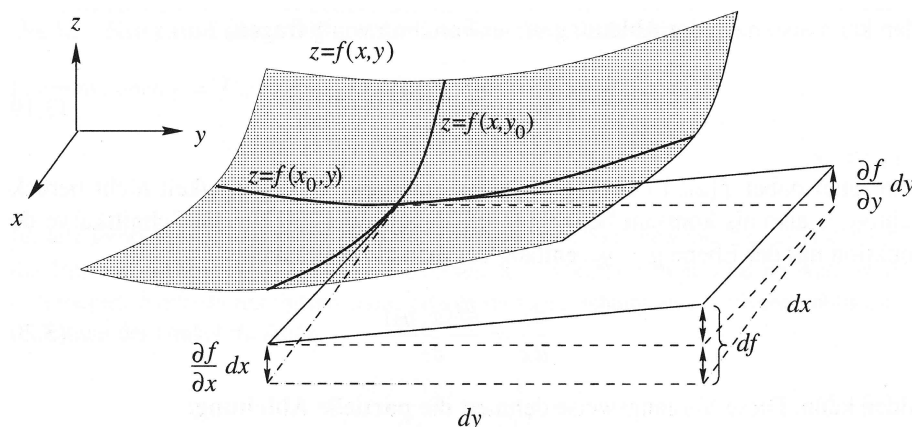


Abb. 3.1 Veranschaulichung der geometrischen Bedeutung des totalen Differentials der Funktion $f(x, y) = z$. Die Funktionswerte z bilden die punktierte gekrümmte Fläche über der (x, y) -Koordinatenebene. Abbildung aus dem Hochschultaschenbuch von Christian B. Lang und Norbert Pucker, Spektrum Akademischer Verlag, Heidelberg, Berlin, 1998, Seite 84.

¹Genauereres dazu findet man unter dem Suchbegriff *Totale Differenzierbarkeit* – Wikipedia.

Mit $dx = x - x_0$, $dy = y - y_0$, $dz = z - z_0 = df(x_0, y_0)$ und $z_0 = f(x_0, y_0)$ liefert das totale Differential die Ebenengleichung für die Tangentialebene an den Graphen der Funktion $f(x, y)$ im Punkt (x_0, y_0) :

$$\begin{aligned} df &= \frac{\partial f}{\partial x} dx + \frac{\partial f}{\partial y} dy \Rightarrow \\ df(x_0, y_0) &= \left. \frac{\partial f}{\partial x} \right|_{(x_0, y_0)} dx + \left. \frac{\partial f}{\partial y} \right|_{(x_0, y_0)} dy \Rightarrow \\ z - z_0 &= \frac{\partial f(x_0, y_0)}{\partial x} (x - x_0) + \frac{\partial f(x_0, y_0)}{\partial y} (y - y_0) \Leftrightarrow \end{aligned}$$

$$\text{Tangentialebene} := z = f(x_0, y_0) + \frac{\partial f(x_0, y_0)}{\partial x} (x - x_0) + \frac{\partial f(x_0, y_0)}{\partial y} (y - y_0)$$

Die Differentiale dx und dy bzw. die Differenzen $x - x_0 = \Delta x$ und $y - y_0 = \Delta y$ stehen hier in der ersten Potenz, beschreiben also einen linearen Zusammenhang, eine Ebene. Vergleichen wir die Tangentialebenengleichung mit der Taylor-Entwicklung von Funktionen mehrerer Variablen, so stellen wir fest, dass das totale Differential (die Ebenengleichung) eine lineare Näherung von f ist, d. h. eine Taylor-Entwicklung bis zur ersten Ordnung.

Handelt es sich bei f um eine Funktion von nur einer Variablen wie beispielsweise $f(x) = y$, so entspricht das zugehörige totale Differential einer Geraden, nämlich der Tangente an die Funktion f im Punkt $f(x_0)$ mit der Tangentengleichung

$$y = f(x) = f(x_0) + \left. \frac{df(x)}{dx} \right|_{x_0} \cdot (x - x_0).$$

Diese Tangentengleichung ist aber gerade die Taylor-Entwicklung der Funktion $f(x)$ an der Stelle x_0 bis zum linearen Glied, also bis zur ersten Ordnung. Und wie man hier sieht, stimmen im 1-dimensionalen Fall und nur im 1-dimensionalen Fall die klassische reelle, die partielle und die totale Differenzierbarkeit überein.

Ist f jedoch von mehr als zwei Variablen abhängig, so spricht man im Zusammenhang mit der geometrischen Bedeutung des zugehörigen totalen Differential trotzdem von einer Tangentialebene. Ein **linearer** Zusammenhang entspricht immer einer Ebene bzw. **ebenen** Verhältnissen.

Eine wichtige Rolle spielen totale Differentiale in der **klassischen Thermodynamik**. Im Gegensatz zu den **Prozessgrößen** Wärmeenergie Q und Arbeit W mit ihren „infinitesimalen“ Änderungen δQ und δW besitzen die **Zustandsgrößen**² totale Differentiale. So besitzt beispielsweise die Zustandsgröße (die Zustandsfunktion, das thermodynamische Potential) innere Energie U das totale Differential

$$dU = \left(\frac{\partial U(V, T)}{\partial V} \right)_{(T)} dV + \left(\frac{\partial U(V, T)}{\partial T} \right)_{(V)} dT.$$

²**Intensive Zustandsgrößen:** z. B. Druck p und Temperatur T .

Extensive Zustandsgrößen: z. B. Volumen V , Teilchenzahl N , Stoffmenge, Entropie und thermodynamische Potentiale.

Thermodynamische Potentiale: innere Energie U , freie Energie, Gibbs-Energie und großkanonisches Potential.

Die tiefgestellten Indizes in Klammern, also (T) und (V) , bedeuten hier, dass T bzw. V jeweils festgehalten werden. In Analogie zu (3.1) können wir $U(V, T)$ als thermodynamisches **Potential** bzw. als skalares Feld über den unabhängigen Variablen V und T ansehen. Der Gradient von U liefert demzufolge das Gradienten- bzw. Vektorfeld

$$\vec{V}(V, T) = \text{grad } U(V, T) = \begin{pmatrix} \frac{\partial U(V, T)}{\partial V} \\ \frac{\partial U(V, T)}{\partial T} \end{pmatrix}.$$

Damit erhält das totale Differential der inneren Energie U die Form

$$dU = \text{grad } U(V, T) \cdot \begin{pmatrix} dV \\ dT \end{pmatrix} = \vec{V} \cdot \begin{pmatrix} dV \\ dT \end{pmatrix} \quad (3.2)$$

Schließlich integrieren wir (3.2) in den Grenzen von V_1, T_1 bis V_2, T_2 :

$$\begin{aligned} \int_{V_1, T_1}^{V_2, T_2} dU &= \int_{V_1, T_1}^{V_2, T_2} \vec{V} \cdot \begin{pmatrix} dV \\ dT \end{pmatrix} \\ &= \int_{V_1, T_1}^{V_2, T_2} \left[\left(\frac{\partial U(V, T)}{\partial V} \right)_{(T)} dV + \left(\frac{\partial U(V, T)}{\partial T} \right)_{(V)} dT \right] \\ &= U(V_2, T_2) - U(V_1, T_1) = \Delta U. \end{aligned}$$

Wir können feststellen, dass das zur inneren Energie gehörende Gradientenfeld $\vec{V}(V, T)$ dem konservativen Vektorfeld $\vec{V}(\vec{r})$ entspricht. Deshalb ist das Integral über $\vec{V}(V, T)$ wegunabhängig und somit das totale Differential dU bzw. die Änderung ΔU der inneren Energie nur vom Anfangs- und Endpunkt der Änderung abhängig. Oft vereinfachen sich dadurch Problemlösungen in der klassischen Thermodynamik oder werden dadurch erst möglich.

Explizite und implizite Zeitabhängigkeit, totale Zeitableitung

Insbesondere in der Mechanik spielen die verschiedenen Zeitabhängigkeiten physikalischer Größen und die sich daraus ergebenden verschiedenen **Ableitungen nach der Zeit** eine große Rolle. Betrachten wir also eine physikalische Größe f als Funktion von der Zeit und von den kartesischen Koordinaten x, y, z im Ortsraum:

- $f = f(\vec{r}, t) = f(x, y, z, t)$

f hängt von den Ortskoordinaten und **explizit** von der Zeit ab, wobei die Ortskoordinaten nicht von der Zeit abhängen. Dabei werden die Orts- und die Zeitabhängigkeit von f an einem **festen** Ort betrachtet. Weil die Ortskoordinaten zeitunabhängig sind, ist die Zeitableitung von f in diesem Fall ausschließlich partiell:

$$\frac{\partial f(\vec{r}, t)}{\partial t} = \frac{\partial f(x, y, z, t)}{\partial t}. \quad (3.3)$$

- $f = f(\vec{r}(t)) = f(x(t), y(t), z(t))$

f hängt über die zeitabhängigen Ortskoordinaten ausschließlich **implizit** von der Zeit ab. Wir betrachten f also an einem mit der Geschwindigkeit $\vec{v}$ **bewegten** Raumpunkt. Die Zeitableitung erfolgt hier nach der Kettenregel:

$$\begin{aligned} \frac{df(\vec{r}(t))}{dt} &= \frac{\partial f(x, y, z)}{\partial x} \cdot \frac{dx}{dt} + \frac{\partial f(x, y, z)}{\partial y} \cdot \frac{dy}{dt} + \frac{\partial f(x, y, z)}{\partial z} \cdot \frac{dz}{dt} \quad (3.4) \\ &= \text{grad } f \cdot \frac{d\vec{r}}{dt} = \vec{v} \cdot \nabla f = (\vec{v} \cdot \vec{\nabla}) f \end{aligned}$$

$$\text{mit } \vec{v} = \begin{pmatrix} \frac{dx}{dt} \\ \frac{dy}{dt} \\ \frac{dz}{dt} \end{pmatrix} = \begin{pmatrix} v_x \\ v_y \\ v_z \end{pmatrix} \quad \text{und dem Operator } (\vec{v} \cdot \vec{\nabla}).$$

Formal liefert die Multiplikation von df/dt mit dt das totale Differential df .

- $f = f(\vec{r}(t), t) = f(x(t), y(t), z(t), t)$

f ist sowohl implizit über die Zeitabhängigkeit der Ortskoordinaten als auch explizit zeitabhängig. Man spricht dann auch von einer zusammengesetzten Funktion. Die Ableitung von f nach der Zeit setzt sich folglich zusammen aus (3.3) und (3.4) und heißt **totale Ableitung**:

$$\begin{aligned} \frac{df(\vec{r}(t), t)}{dt} &= \frac{\partial f}{\partial x} \cdot \frac{dx}{dt} + \frac{\partial f}{\partial y} \cdot \frac{dy}{dt} + \frac{\partial f}{\partial z} \cdot \frac{dz}{dt} + \frac{\partial f}{\partial t} \\ &= \text{grad } f \cdot \frac{d\vec{r}}{dt} + \frac{\partial f}{\partial t} = (\vec{v} \cdot \nabla) f + \frac{\partial f}{\partial t}. \end{aligned}$$

4 Lösung der Integrale $\int_0^{\infty} x^q \cdot e^{-ax^2} dx$, $q \in \mathbb{N}$, $a \in \mathbb{R}^{>0}$

Siehe auch Integrationsmethode nach Feynman oder Feynman-Trick.

$\mathbb{N}$ sei die Menge der natürlichen Zahlen **einschließlich der Null**. Wir werden im Folgenden eine Fallunterscheidung vornehmen und betrachten zuerst den Fall mit

$$\textbf{ungeraden } q \in \mathbb{N} \Rightarrow m = 2k + 1 \text{ mit } k = 0, 1, 2, 3, \dots$$

und anschließend den Fall mit

$$\textbf{geraden } q \in \mathbb{N} \Rightarrow n = 2k \text{ mit } k = 0, 1, 2, 3, \dots$$

Außerdem ist zu berücksichtigen, dass die Integranden $x^q \cdot e^{-ax^2}$ mit ungeradem q bzw. m ungerade Funktionen und mit geradem q bzw. n gerade Funktionen von x sind.

4.1 Für ungerade q

Die Substitution

$$\begin{aligned} x^2 = u &\Rightarrow x^m = u^{\frac{m}{2}}, \\ &\Rightarrow x^{2k+1} = u^{k+\frac{1}{2}} = u^k \cdot u^{\frac{1}{2}}, \\ &\Rightarrow \frac{du}{dx} = 2x, \quad dx = \frac{1}{2x} du \end{aligned}$$

liefert

$$\begin{aligned} \int_0^{\infty} x^m e^{-ax^2} dx &= \int_0^{\infty} x^{2k+1} \cdot \frac{1}{2x} \cdot e^{-au} du = \frac{1}{2} \int_0^{\infty} x^{2k} \cdot e^{-au} du = \frac{1}{2} \int_0^{\infty} u^k \cdot e^{-au} du \Rightarrow \\ \int_0^{\infty} x^m e^{-ax^2} dx &= \frac{1}{2} \int_0^{\infty} \left(-\frac{\partial}{\partial a} \right)^k e^{-au} du, \quad k = \frac{m-1}{2} \text{ für } \textbf{ungerade } m \in \mathbb{N}, \quad u = x^2, \end{aligned} \quad (4.1)$$

denn

$$\begin{aligned} m = 1 &\Rightarrow k = 0 \Rightarrow \left(-\frac{\partial}{\partial a} \right)^0 e^{-au} = e^{-au}, \\ m = 3 &\Rightarrow k = 1 \Rightarrow \left(-\frac{\partial}{\partial a} \right)^1 e^{-au} = u \cdot e^{-au}, \\ m = 5 &\Rightarrow k = 2 \Rightarrow \left(-\frac{\partial}{\partial a} \right)^2 e^{-au} = u^2 \cdot e^{-au}, \\ m = 7 &\Rightarrow k = 3 \Rightarrow \left(-\frac{\partial}{\partial a} \right)^3 e^{-au} = u^3 \cdot e^{-au}, \\ &\vdots \end{aligned}$$

$$\left(-\frac{\partial}{\partial a} \right)^k e^{-au} = (-1)^k \cdot \left(\frac{\partial}{\partial a} \right)^k e^{-au} = u^k \cdot e^{-au}. \quad (4.2)$$

Damit kann man für (4.1)

$$\frac{1}{2} \int_0^{\infty} \left(-\frac{\partial}{\partial a} \right)^k e^{-au} du = \frac{1}{2} \cdot (-1)^k \cdot \frac{\partial^k}{\partial a^k} \int_0^{\infty} e^{-au} du$$

schreiben und mit

$$\int_0^{\infty} e^{-au} du = -\frac{1}{a} \cdot e^{-au} \Big|_{u=0}^{\infty} = \frac{1}{a} = a^{-1}$$

schließlich

$$\boxed{\int_0^{\infty} x^m e^{-ax^2} dx = \frac{1}{2} (-1)^k \frac{\partial^k}{\partial a^k} a^{-1} = \frac{k!}{2 \cdot a^{k+1}} \text{ für } \mathbf{ungerade} \ m \in \mathbb{N} \text{ und } k = \frac{m-1}{2}}. \quad (4.3)$$

Beispiel $m = 1 \Rightarrow k = 0$:

$$\int_0^{\infty} x \cdot e^{-ax^2} dx = \frac{1}{2} \cdot (-1)^0 \cdot a^{-1} = \frac{1}{2} \cdot 0! \cdot \frac{1}{a^1},$$

$$\boxed{\int_0^{\infty} x \cdot e^{-ax^2} dx = \frac{1}{2a}}. \quad (4.4)$$

Beispiel $m = 3 \Rightarrow k = 1$:

$$\int_0^{\infty} x^3 \cdot e^{-ax^2} dx = \frac{1}{2} \cdot (-1)^1 \cdot (-1) \cdot a^{-2} = \frac{1}{2} \cdot 1! \cdot \frac{1}{a^2},$$

$$\boxed{\int_0^{\infty} x^3 \cdot e^{-ax^2} dx = \frac{1}{2a^2}}.$$

Beispiel $m = 5 \Rightarrow k = 2$:

$$\int_0^{\infty} x^5 \cdot e^{-ax^2} dx = \frac{1}{2} \cdot (-1)^2 \cdot (-1) \cdot (-2) \cdot a^{-3} = \frac{1}{2} \cdot 2! \cdot \frac{1}{a^3},$$

$$\boxed{\int_0^{\infty} x^5 \cdot e^{-ax^2} dx = \frac{1}{a^3}}.$$

Beispiel $m = 7 \Rightarrow k = 3$:

$$\int_0^{\infty} x^7 \cdot e^{-ax^2} dx = \frac{1}{2} \cdot (-1)^3 \cdot (-1) \cdot (-2) \cdot (-3) \cdot a^{-4} = \frac{1}{2} \cdot 3! \cdot \frac{1}{a^4},$$

$$\boxed{\int_0^{\infty} x^7 \cdot e^{-ax^2} dx = \frac{3}{a^4}}.$$

4.2 Für $q = 0$

Mit $\int_{-\infty}^{\infty} x^0 e^{-ax^2} dx = \int_{-\infty}^{\infty} y^0 e^{-ay^2} dy = \int_{-\infty}^{\infty} e^{-ax^2} dx$ erhalten wir

$$\int_{-\infty}^{\infty} x^0 e^{-ax^2} dx = \sqrt{\int_{-\infty}^{\infty} e^{-ax^2} dx \cdot \int_{-\infty}^{\infty} e^{-ay^2} dy} = \sqrt{\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} e^{-a(x^2+y^2)} dx dy}.$$

Die Verwendung von ebenen Polarkoordinaten ermöglicht die Substitutionen

$$x^2 + y^2 = r^2, \quad \text{Flächenelement } dx dy = r dr d\varphi,$$

sodass

$$\int_{-\infty}^{\infty} x^0 e^{-ax^2} dx = \sqrt{\int_0^{2\pi} d\varphi \cdot \int_0^{\infty} r e^{-ar^2} dr}.$$

Mit $\int_0^{\infty} r e^{-ar^2} dr = \frac{1}{2a}$ gemäß (4.4) und mit $\int_0^{2\pi} d\varphi = 2\pi$ resultiert daraus schließlich

$$\int_{-\infty}^{\infty} x^0 e^{-ax^2} dx = \sqrt{2\pi \cdot \frac{1}{2a}},$$

$$\boxed{\int_{-\infty}^{\infty} e^{-ax^2} dx = \sqrt{\frac{\pi}{a}} \Rightarrow \int_0^{\infty} e^{-ax^2} dx = \frac{1}{2} \sqrt{\frac{\pi}{a}}, \quad \text{für } n = 0} \quad (4.5)$$

Hierbei haben wir berücksichtigt, dass e^{-ax^2} eine gerade Funktion ist.

4.3 Für gerade q

Mit der Substitution $x^2 = u \Rightarrow x^n = u^{\frac{n}{2}} = u^k$ gilt zunächst

$$\int_0^{\infty} x^n e^{-ax^2} dx = \int_0^{\infty} u^k e^{-au} du.$$

Nach Äquivalenzumformung von (4.2) gemäß

$$u^k \cdot e^{-au} = (-1)^k \cdot \left(\frac{\partial}{\partial a} \right)^k e^{-au} \Leftrightarrow u^k = e^{au} \cdot (-1)^k \cdot \left(\frac{\partial}{\partial a} \right)^k e^{-au}$$

können wir u^k im Integranden ersetzen und erhalten

$$\begin{aligned} \int_0^{\infty} x^n e^{-ax^2} dx &= \int_0^{\infty} e^{au} \cdot (-1)^k \cdot \left[\left(\frac{\partial}{\partial a} \right)^k e^{-au} \right] \cdot e^{-au} du \\ &= (-1)^k \cdot \int_0^{\infty} \underbrace{e^{au} \cdot e^{-au}}_{=1} \cdot \left(\frac{\partial}{\partial a} \right)^k e^{-au} du. \end{aligned}$$

Wir ziehen wieder den Differentiationsoperator vor das Integral und verwenden für das verbleibende Integral die Lösung (4.5):

$$\begin{aligned}\int_0^{\infty} x^n e^{-ax^2} dx &= (-1)^k \cdot \left(\frac{\partial}{\partial a}\right)^k \int_0^{\infty} e^{-au} dx = (-1)^k \cdot \left(\frac{\partial}{\partial a}\right)^k \int_0^{\infty} e^{-ax^2} dx \\ &= (-1)^k \cdot \left(\frac{\partial}{\partial a}\right)^k \frac{1}{2} \sqrt{\frac{\pi}{a}},\end{aligned}$$

$$\boxed{\int_0^{\infty} x^n e^{-ax^2} dx = \frac{1}{2} \cdot (-1)^k \cdot \sqrt{\pi} \cdot \left(\frac{\partial}{\partial a}\right)^k a^{-\frac{1}{2}} \text{ für } \mathbf{gerade} \ n \in \mathbb{N} \text{ und } k = \frac{n}{2}}.$$

Beispiel $n = 2 \Rightarrow k = 1$:

$$\begin{aligned}\int_0^{\infty} x^2 \cdot e^{-ax^2} dx &= \frac{1}{2} \cdot (-1)^1 \cdot \sqrt{\pi} \cdot \left(\frac{\partial}{\partial a}\right)^1 a^{-\frac{1}{2}} = \frac{1}{2} \cdot (-1) \cdot \sqrt{\pi} \cdot \left(-\frac{1}{2}\right) \cdot a^{-\frac{3}{2}}, \\ \boxed{\int_0^{\infty} x^2 \cdot e^{-ax^2} dx &= \frac{1}{2} \cdot \sqrt{\pi} \cdot \frac{1}{2} a^{-\frac{3}{2}} \Rightarrow \int_{-\infty}^{\infty} x^2 \cdot e^{-ax^2} dx = \sqrt{\pi} \cdot \frac{1}{2} a^{-\frac{3}{2}}}.\end{aligned}$$

Beispiel $n = 4 \Rightarrow k = 2$:

$$\begin{aligned}\int_0^{\infty} x^4 \cdot e^{-ax^2} dx &= \frac{1}{2} \cdot (-1)^2 \cdot \sqrt{\pi} \cdot \left(\frac{\partial}{\partial a}\right)^2 a^{-\frac{1}{2}} = \frac{1}{2} \cdot \sqrt{\pi} \cdot \left(-\frac{1}{2}\right) \cdot \left(-\frac{3}{2}\right) \cdot a^{-\frac{5}{2}}, \\ \boxed{\int_0^{\infty} x^4 \cdot e^{-ax^2} dx &= \frac{1}{2} \cdot \sqrt{\pi} \cdot \frac{3}{4} a^{-\frac{5}{2}} \Rightarrow \int_{-\infty}^{\infty} x^4 \cdot e^{-ax^2} dx = \sqrt{\pi} \cdot \frac{3}{4} a^{-\frac{5}{2}}}.\end{aligned}$$

Beispiel $n = 6 \Rightarrow k = 3$:

$$\begin{aligned}\int_0^{\infty} x^6 \cdot e^{-ax^2} dx &= \frac{1}{2} \cdot (-1)^3 \cdot \sqrt{\pi} \cdot \left(\frac{\partial}{\partial a}\right)^3 a^{-\frac{1}{2}} = \frac{1}{2} \cdot \left(-\frac{1}{2}\right) \cdot \sqrt{\pi} \cdot \left(-\frac{1}{2}\right) \cdot \left(-\frac{3}{2}\right) \cdot \left(-\frac{5}{2}\right) \cdot a^{-\frac{7}{2}}, \\ \boxed{\int_0^{\infty} x^6 \cdot e^{-ax^2} dx &= \frac{1}{2} \cdot \sqrt{\pi} \cdot \frac{15}{8} a^{-\frac{7}{2}} \Rightarrow \int_{-\infty}^{\infty} x^6 \cdot e^{-ax^2} dx = \sqrt{\pi} \cdot \frac{15}{8} a^{-\frac{7}{2}}}.\end{aligned}$$

Die letzten drei Beispiele lassen sich verallgemeinern zu

$$\boxed{\int_0^{\infty} x^n e^{-ax^2} dx = \frac{1}{2} \sqrt{\pi} \frac{1 \cdot 3 \cdot 5 \cdot \dots \cdot (n-1)}{2^{\frac{n}{2}}} a^{-\frac{n+1}{2}} \text{ für } \mathbf{gerade} \ n \in \mathbb{N} \setminus 0}.$$

Der Vollständigkeit halber zeigen wir noch die Lösungen von häufig vorkommenden Integralen, die eine gewisse Ähnlichkeit zu den bisher dargestellten Integralen der Form $\int_0^\infty x^q e^{-ax^2} dx$ besitzen.

- e^{-ax} statt e^{-ax^2} und $q = N \in \mathbb{N}$:

$$\int_0^\infty x^N e^{-ax} dx = \frac{N!}{a^{N+1}}, \quad N \in \mathbb{N}, \quad 0! = 1, \quad a \in \mathbb{R}^{>0}$$

lässt sich zeigen mit der Substitution

$$-ax = u \quad \Rightarrow \quad dx = -\frac{1}{a} du,$$

mit den Grenzen

$$\begin{aligned} x = 0 & \Rightarrow u = 0, \\ x \rightarrow \infty & \Rightarrow u \rightarrow -\infty \end{aligned}$$

und durch ggf. wiederholte Anwendung der partiellen Integration.

- e^{-ax} statt e^{-ax^2} und $q = \frac{m}{2}$ mit $m = 1, 3, 5, \dots$:

$$\int_0^\infty x^{\frac{1}{2}} e^{-ax} dx = \frac{1}{2} \sqrt{\frac{\pi}{a^3}} \quad \text{und} \quad \int_0^\infty x^{\frac{3}{2}} e^{-ax} dx = \frac{3}{4} \sqrt{\frac{\pi}{a^5}}, \quad a \in \mathbb{R}^{>0}$$

kann man zeigen mit der Substitution

$$x = u^2 \quad \text{bzw.} \quad x^{\frac{1}{2}} = u \quad \text{und} \quad x^{\frac{3}{2}} = u^3 \quad \Rightarrow \quad dx = 2u du$$

und durch anschließende Anwendung von (4.3).

5 Polar-, Zylinder- und Kugelkoordinaten auf einen Blick

(Siehe Lehrbuch aus der Teubner Studienbücherei Physik: Siegfried Großmann, Mathematischer Einführungskurs für die Physik, Teubner-Verlag, Stuttgart, Leipzig, 8. Auflage, 2000.)

Es ist üblich, weil es am einfachsten und am anschaulichsten ist, die verschiedenen Koordinatensysteme in Beziehung zu kartesischen Koordinaten darzustellen. Das kartesische (x, y, z) -Koordinatensystem ist nämlich ausgezeichnet durch seine globale Orthonormalbasis $\{\vec{e}_x, \vec{e}_y, \vec{e}_z\}$. Die kartesische Orthonormalbasis wird Standardbasis, natürliche Basis oder kanonische Basis genannt. Die Darstellung von Ortsvektoren $\vec{r}$ in nichtkartesischen Koordinaten u, v, w erfolgt meistens in der Form

$$\underbrace{\vec{r} = \begin{pmatrix} x \\ y \\ z \end{pmatrix}}_{\text{kartesisch}} = x\vec{e}_x + y\vec{e}_y + z\vec{e}_z \longrightarrow \vec{r} = \begin{pmatrix} x(u, v, w) \\ y(u, v, w) \\ z(u, v, w) \end{pmatrix} = x(u, v, w)\vec{e}_x + y(u, v, w)\vec{e}_y + z(u, v, w)\vec{e}_z.$$

Besitzt das (u, v, w) -Koordinatensystem die normierte Basis $\{\vec{e}_u, \vec{e}_v, \vec{e}_w\}$, so kann die Darstellung eines Ortsvektors $\vec{r}$ in der Form

$$\vec{r} = u\vec{e}_u + v\vec{e}_v + w\vec{e}_w \longrightarrow \vec{r} = u(x, y, z)\vec{e}_u(x, y, z) + v(x, y, z)\vec{e}_v(x, y, z) + w(x, y, z)\vec{e}_w(x, y, z)$$

erfolgen. Dabei haben wir berücksichtigt, dass die Standardbasis unabhängig von $\vec{r}$, also ortsunabhängig bzw. **global** ist, während die Basisvektoren $\{\vec{e}_u, \vec{e}_v, \vec{e}_w\}$ allgemein abhängig von den Ortskoordinaten und somit **lokal** sind (wie z. B. bei krummlinigen Koordinatensystemen wie den Polarkoordinatensystemen). Wenn man von Polarkoordinaten im Allgemeinen spricht, sind sphärische Polarkoordinaten (Kugelkoordinaten), Zylinderkoordinaten und ebene Polarkoordinaten (kurz Polarkoordinaten) gemeint. **Diese drei Polarkoordinatensysteme sind lokal orthogonal und teilweise krummlinig.**

5.1 Polarkoordinaten (r, φ)

- Polarkoordinaten (für $0 \leq \varphi < 2\pi$), dargestellt in kartesischen Koordinaten:

$$r = \sqrt{x^2 + y^2},$$

$$\varphi = \begin{cases} \arctan\left(\frac{y}{x}\right) & , \text{ wenn } x > 0, \\ \arctan\left(\frac{y}{x}\right) + \pi & , \text{ wenn } x < 0 \wedge y \geq 0, \\ \arctan\left(\frac{y}{x}\right) - \pi & , \text{ wenn } x < 0 \wedge y < 0, \\ \arccos\frac{x}{r} & , \text{ wenn } y \geq 0, \\ 2\pi - \arccos\frac{x}{r} & , \text{ wenn } y < 0. \end{cases}$$

- Ortsvektor $\vec{r}$:

$$\vec{r} = \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} r \cdot \cos \varphi \\ r \cdot \sin \varphi \end{pmatrix} = r \cos \varphi \cdot \vec{e}_x + r \sin \varphi \cdot \vec{e}_y = r \cdot \vec{e}_r \Rightarrow$$

$$|\vec{g}_r| = \left| \frac{\partial \vec{r}}{\partial r} \right| = g_r = 1, \quad |\vec{g}_\varphi| = \left| \frac{\partial \vec{r}}{\partial \varphi} \right| = g_\varphi = r.$$

- Basiseinheitsvektoren:

$$\vec{e}_r = \frac{\frac{\partial \vec{r}}{\partial r}}{\left| \frac{\partial \vec{r}}{\partial r} \right|} = \begin{pmatrix} \cos \varphi \\ \sin \varphi \end{pmatrix}, \quad \vec{e}_\varphi = \frac{\frac{\partial \vec{r}}{\partial \varphi}}{\left| \frac{\partial \vec{r}}{\partial \varphi} \right|} = \begin{pmatrix} -\sin \varphi \\ \cos \varphi \end{pmatrix}.$$

- Jacobi-Matrix J :

$$J = \begin{pmatrix} \cos \varphi & -r \sin \varphi \\ \sin \varphi & r \cos \varphi \end{pmatrix} = (\vec{g}_r, \vec{g}_\varphi) \Leftrightarrow J^{-1} = \begin{pmatrix} \cos \varphi & \sin \varphi \\ -\frac{1}{r} \sin \varphi & \frac{1}{r} \cos \varphi \end{pmatrix} = \begin{pmatrix} \frac{x}{\sqrt{x^2+y^2}} & \frac{y}{\sqrt{x^2+y^2}} \\ \frac{-y}{x^2+y^2} & \frac{x}{x^2+y^2} \end{pmatrix}.$$

- Drehmatrix (Rotationsmatrix) D :

$$J = D \cdot h = \underbrace{\begin{pmatrix} \cos \varphi & -\sin \varphi \\ \sin \varphi & \cos \varphi \end{pmatrix}}_D \cdot \underbrace{\begin{pmatrix} 1 & 0 \\ 0 & r \end{pmatrix}}_h = \begin{pmatrix} \cos \varphi & -r \sin \varphi \\ \sin \varphi & r \cos \varphi \end{pmatrix}.$$

D ist orthogonal $\Rightarrow D^{-1} = D^T$.

- Funktionaldeterminante $\det J = r$.
- Metriktenor: $(g_{kl}) = J^T \cdot J = h^2 = \begin{pmatrix} 1 & 0 \\ 0 & r^2 \end{pmatrix}$.
- Transformation der Basiseinheitsvektoren:

$$(\vec{e}_x, \vec{e}_y) \cdot D = (\vec{e}_r, \vec{e}_\varphi) \Leftrightarrow (\vec{e}_r, \vec{e}_\varphi) \cdot D^T = (\vec{e}_x, \vec{e}_y) \quad \text{mit } D^T = D^{-1} \text{ (} D \text{ ist orthogonal)}.$$

- Transformation von Differentialen:

$$J^{-1} \cdot \begin{pmatrix} dx \\ dy \end{pmatrix} = \begin{pmatrix} dr \\ d\varphi \end{pmatrix} \Leftrightarrow J \cdot \begin{pmatrix} dr \\ d\varphi \end{pmatrix} = \begin{pmatrix} dx \\ dy \end{pmatrix}.$$

- Transformation partieller Ableitungen:

$$\left(\frac{\partial}{\partial x}, \frac{\partial}{\partial y} \right) \cdot J = \left(\frac{\partial}{\partial r}, \frac{\partial}{\partial \varphi} \right) \Leftrightarrow \left(\frac{\partial}{\partial r}, \frac{\partial}{\partial \varphi} \right) \cdot J^{-1} = \left(\frac{\partial}{\partial x}, \frac{\partial}{\partial y} \right).$$

Mit $\cos \varphi = \frac{x}{r} = \frac{x}{\sqrt{x^2+y^2}}$ und $\sin \varphi = \frac{y}{r} = \frac{y}{\sqrt{x^2+y^2}}$ ergibt das

$$\begin{aligned} \frac{\partial}{\partial r} &= \frac{x}{\sqrt{x^2+y^2}} \frac{\partial}{\partial x} + \frac{y}{\sqrt{x^2+y^2}} \frac{\partial}{\partial y}, & \frac{\partial}{\partial \varphi} &= x \frac{\partial}{\partial y} - y \frac{\partial}{\partial x}, \\ \frac{\partial}{\partial x} &= \cos \varphi \frac{\partial}{\partial r} - \frac{\sin \varphi}{r} \frac{\partial}{\partial \varphi}, & \frac{\partial}{\partial y} &= \sin \varphi \frac{\partial}{\partial r} + \frac{\cos \varphi}{r} \frac{\partial}{\partial \varphi}. \end{aligned}$$

- Transformation eines Vektorfeldes $\vec{F}(\vec{r})$:

$$D^T \cdot \begin{pmatrix} F_x \\ F_y \end{pmatrix} = \begin{pmatrix} F_r \\ F_\varphi \end{pmatrix} \Leftrightarrow D \cdot \begin{pmatrix} F_r \\ F_\varphi \end{pmatrix} = \begin{pmatrix} F_x \\ F_y \end{pmatrix},$$

$$\vec{F} = F_x \vec{e}_x + F_y \vec{e}_y = F_r \vec{e}_r + F_\varphi \vec{e}_\varphi.$$

- Flächenelement: $dS(r, \varphi) = |\det J| \cdot dr d\varphi = r dr d\varphi$.
- Linienelement ds :

$$(\mathbf{d}\vec{r})^2 = (ds)^2 = g_r^2 (dr)^2 + g_\varphi^2 (d\varphi)^2 = (dr)^2 + r^2 (d\varphi)^2.$$

- Gradient: $\text{grad } f(r, \varphi) = \frac{\partial f}{\partial r} \vec{e}_r + \frac{1}{r} \frac{\partial f}{\partial \varphi} \vec{e}_\varphi$.
- Divergenz: $\text{div } \vec{F}(r, \varphi) = \frac{1}{r} \frac{\partial}{\partial r} (r \cdot F_r) + \frac{1}{r} \frac{\partial F_\varphi}{\partial \varphi}$.
- Rotation: $\left[\text{rot } \vec{F}(r, \varphi) \right]_z = \frac{1}{r} \frac{\partial}{\partial r} (r \cdot F_\varphi) - \frac{1}{r} \frac{\partial F_r}{\partial \varphi}$.
- Laplacian: $\Delta f(r, \varphi) = \frac{\partial^2 f}{\partial r^2} + \frac{1}{r} \frac{\partial f}{\partial r} + \frac{1}{r^2} \frac{\partial^2 f}{\partial \varphi^2}$,

wobei der Laplace-Operator oder Laplacian $\Delta := \nabla \cdot \nabla = \nabla^2 = \text{div grad}$.

5.2 Zylinderkoordinaten (r, φ, z)

- Zylinderkoordinaten (für $0 \leq \varphi < 2\pi$), dargestellt in kartesischen Koordinaten:

$$r = \sqrt{x^2 + y^2},$$

$$\varphi = \begin{cases} \arctan\left(\frac{y}{x}\right) & , \text{ wenn } x > 0, \\ \arctan\left(\frac{y}{x}\right) + \pi & , \text{ wenn } x < 0 \wedge y \geq 0, \\ \arctan\left(\frac{y}{x}\right) - \pi & , \text{ wenn } x < 0 \wedge y < 0, \\ \arccos\frac{x}{r} & , \text{ wenn } y \geq 0, \\ 2\pi - \arccos\frac{x}{r} & , \text{ wenn } y < 0, \end{cases}$$

$$z = z.$$

- Ortsvektor $\vec{r}$:

$$\vec{r} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} r \cdot \cos \varphi \\ r \cdot \sin \varphi \\ z \end{pmatrix} = r \cos \varphi \cdot \vec{e}_x + r \sin \varphi \cdot \vec{e}_y + z \cdot \vec{e}_z = r \cdot \vec{e}_r + z \cdot \vec{e}_z \Rightarrow$$

$$\vec{g}_r = \frac{\partial \vec{r}}{\partial r} = \begin{pmatrix} \cos \varphi \\ \sin \varphi \\ 0 \end{pmatrix}, \quad \vec{g}_\varphi = \frac{\partial \vec{r}}{\partial \varphi} = \begin{pmatrix} -r \sin \varphi \\ r \cos \varphi \\ 0 \end{pmatrix}, \quad \vec{g}_z = \frac{\partial \vec{r}}{\partial z} = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix},$$

$$|\vec{g}_r| = \left| \frac{\partial \vec{r}}{\partial r} \right| = g_r = 1, \quad |\vec{g}_\varphi| = \left| \frac{\partial \vec{r}}{\partial \varphi} \right| = g_\varphi = r, \quad |\vec{g}_z| = \left| \frac{\partial \vec{r}}{\partial z} \right| = g_z = 1.$$

- Basiseinheitsvektoren:

$$\vec{e}_r = \frac{\frac{\partial \vec{r}}{\partial r}}{\left| \frac{\partial \vec{r}}{\partial r} \right|} = \begin{pmatrix} \cos \varphi \\ \sin \varphi \\ 0 \end{pmatrix}, \quad \vec{e}_\varphi = \frac{\frac{\partial \vec{r}}{\partial \varphi}}{\left| \frac{\partial \vec{r}}{\partial \varphi} \right|} = \begin{pmatrix} -\sin \varphi \\ \cos \varphi \\ 0 \end{pmatrix}, \quad \vec{e}_z = \frac{\frac{\partial \vec{r}}{\partial z}}{\left| \frac{\partial \vec{r}}{\partial z} \right|} = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

- Jacobi-Matrix J :

$$J = \begin{pmatrix} \cos \varphi & -r \sin \varphi & 0 \\ \sin \varphi & r \cos \varphi & 0 \\ 0 & 0 & 1 \end{pmatrix} = (\vec{g}_r, \vec{g}_\varphi, \vec{g}_z) \Leftrightarrow$$

$$J^{-1} = \begin{pmatrix} \cos \varphi & \sin \varphi & 0 \\ -\frac{1}{r} \sin \varphi & \frac{1}{r} \cos \varphi & 0 \\ 0 & 0 & 1 \end{pmatrix} = \begin{pmatrix} \frac{x}{\sqrt{x^2+y^2}} & \frac{y}{\sqrt{x^2+y^2}} & 0 \\ \frac{-y}{x^2+y^2} & \frac{x}{x^2+y^2} & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

- Drehmatrix (Rotationsmatrix) D :

$$J = D \cdot h = \underbrace{\begin{pmatrix} \cos \varphi & -\sin \varphi & 0 \\ \sin \varphi & \cos \varphi & 0 \\ 0 & 0 & 1 \end{pmatrix}}_D \cdot \underbrace{\begin{pmatrix} 1 & 0 & 0 \\ 0 & r & 0 \\ 0 & 0 & 1 \end{pmatrix}}_h.$$

$$D \text{ ist orthogonal} \Rightarrow D^{-1} = D^T.$$

- Funktionaldeterminante $\det J = r$.

- Metriktensor: $(g_{kl}) = J^T \cdot J = h^2 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & r^2 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$

- Transformation der Basiseinheitsvektoren:

$$(\vec{e}_x, \vec{e}_y, \vec{e}_z) \cdot D = (\vec{e}_r, \vec{e}_\varphi, \vec{e}_z) \Leftrightarrow (\vec{e}_r, \vec{e}_\varphi, \vec{e}_z) \cdot D^T = (\vec{e}_x, \vec{e}_y, \vec{e}_z) .$$

- Transformation von Differentialen:

$$J^{-1} \cdot \begin{pmatrix} dx \\ dy \\ dz \end{pmatrix} = \begin{pmatrix} dr \\ d\varphi \\ dz \end{pmatrix} \Leftrightarrow J \cdot \begin{pmatrix} dr \\ d\varphi \\ dz \end{pmatrix} = \begin{pmatrix} dx \\ dy \\ dz \end{pmatrix} .$$

- Transformation partieller Ableitungen:

$$\left(\frac{\partial}{\partial x}, \frac{\partial}{\partial y}, \frac{\partial}{\partial z} \right) \cdot J = \left(\frac{\partial}{\partial r}, \frac{\partial}{\partial \varphi}, \frac{\partial}{\partial z} \right) \Leftrightarrow \left(\frac{\partial}{\partial r}, \frac{\partial}{\partial \varphi}, \frac{\partial}{\partial z} \right) \cdot J^{-1} = \left(\frac{\partial}{\partial x}, \frac{\partial}{\partial y}, \frac{\partial}{\partial z} \right) .$$

- Transformation eines Vektorfeldes $\vec{F}(\vec{r})$:

$$D^T \cdot \begin{pmatrix} F_x \\ F_y \\ F_z \end{pmatrix} = \begin{pmatrix} F_r \\ F_\varphi \\ F_z \end{pmatrix} \Leftrightarrow D \cdot \begin{pmatrix} F_r \\ F_\varphi \\ F_z \end{pmatrix} = \begin{pmatrix} F_x \\ F_y \\ F_z \end{pmatrix} ,$$

$$\vec{F} = F_x \vec{e}_x + F_y \vec{e}_y + F_z \vec{e}_z = F_r \vec{e}_r + F_\varphi \vec{e}_\varphi + F_z \vec{e}_z .$$

- Volumenelement: $dV = |\det J| \cdot dr d\varphi dz = r \cdot dr d\varphi dz$.
- Flächenelement auf dem Zylindermantel mit $r = \text{const}$:

$$dS = \frac{dV}{dr} = |\det J| \cdot d\varphi dz = r \cdot d\varphi dz .$$

Das Flächenelement in den Ebenen $z=\text{const}$ ist gleich dem Flächenelement in ebenen Polarkoordinaten.

- Linienelement ds :

$$(d\vec{r})^2 = (ds)^2 = g_r^2 (dr)^2 + g_\varphi^2 (d\varphi)^2 + g_z^2 (dz)^2 = (dr)^2 + r^2 (d\varphi)^2 + (dz)^2 .$$

- Gradient: $\text{grad } f(r, \varphi, z) = \frac{\partial f}{\partial r} \vec{e}_r + \frac{1}{r} \frac{\partial f}{\partial \varphi} \vec{e}_\varphi + \frac{\partial f}{\partial z} \vec{e}_z$.

- Divergenz: $\text{div } \vec{F}(r, \varphi, z) = \frac{1}{r} \frac{\partial}{\partial r} (r \cdot F_r) + \frac{1}{r} \frac{\partial F_\varphi}{\partial \varphi} + \frac{\partial F_z}{\partial z}$.

- Rotation: $\text{rot } \vec{F}(r, \varphi, z) = \left(\frac{1}{r} \frac{\partial F_z}{\partial \varphi} - \frac{\partial F_\varphi}{\partial z} \right) \vec{e}_r + \left(\frac{\partial F_r}{\partial z} - \frac{\partial F_z}{\partial r} \right) \vec{e}_\varphi + \frac{1}{r} \left(\frac{\partial}{\partial r} (r \cdot F_\varphi) - \frac{\partial F_r}{\partial \varphi} \right) \vec{e}_z$.

- Laplacian: $\Delta f(r, \varphi, z) = \frac{1}{r} \frac{\partial}{\partial r} \left(r \cdot \frac{\partial f}{\partial r} \right) + \frac{1}{r^2} \frac{\partial^2 f}{\partial \varphi^2} + \frac{\partial^2 f}{\partial z^2}$.

5.3 Kugelkoordinaten (r, ϑ, φ)

- Kugelkoordinaten (für $0 \leq \vartheta < \pi$ und $0 \leq \varphi < 2\pi$), dargestellt in kartesischen Koordinaten:

$$\begin{aligned} r &= \sqrt{x^2 + y^2 + z^2}, \\ \vartheta &= \arccos \frac{z}{\sqrt{x^2 + y^2 + z^2}} = \operatorname{arccot} \frac{z}{\sqrt{x^2 + y^2}}, \\ \varphi &= \begin{cases} \arctan\left(\frac{y}{x}\right) & , \text{ wenn } x > 0, \\ \arctan\left(\frac{y}{x}\right) + \pi & , \text{ wenn } x < 0 \wedge y \geq 0, \\ \arctan\left(\frac{y}{x}\right) - \pi & , \text{ wenn } x < 0 \wedge y < 0, \\ \arccos \frac{x}{r} & , \text{ wenn } y \geq 0, \\ 2\pi - \arccos \frac{x}{r} & , \text{ wenn } y < 0. \end{cases} \end{aligned}$$

- Ortsvektor $\vec{r}$:

$$\begin{aligned} \vec{r} &= \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} r \cdot \sin \vartheta \cos \varphi \\ r \cdot \sin \vartheta \sin \varphi \\ r \cdot \cos \vartheta \end{pmatrix} = r \sin \vartheta \cos \varphi \cdot \vec{e}_x + r \sin \vartheta \sin \varphi \cdot \vec{e}_y + r \cos \vartheta \cdot \vec{e}_z \\ &= r \cdot \vec{e}_r \Rightarrow \end{aligned}$$

$$\begin{aligned} \vec{g}_r &= \frac{\partial \vec{r}}{\partial r} = \begin{pmatrix} \sin \vartheta \cos \varphi \\ \sin \vartheta \sin \varphi \\ \cos \vartheta \end{pmatrix}, \quad \vec{g}_\vartheta = \frac{\partial \vec{r}}{\partial \vartheta} = \begin{pmatrix} r \cos \vartheta \cos \varphi \\ r \cos \vartheta \sin \varphi \\ -r \sin \vartheta \end{pmatrix}, \quad \vec{g}_\varphi = \frac{\partial \vec{r}}{\partial \varphi} = \begin{pmatrix} -r \sin \vartheta \sin \varphi \\ r \sin \vartheta \cos \varphi \\ 0 \end{pmatrix}, \\ |\vec{g}_r| &= \left| \frac{\partial \vec{r}}{\partial r} \right| = g_r = 1, \quad |\vec{g}_\vartheta| = \left| \frac{\partial \vec{r}}{\partial \vartheta} \right| = g_\vartheta = r, \quad |\vec{g}_\varphi| = \left| \frac{\partial \vec{r}}{\partial \varphi} \right| = g_\varphi = r \sin \vartheta. \end{aligned}$$

- Basiseinheitsvektoren:

$$\vec{e}_r = \frac{\frac{\partial \vec{r}}{\partial r}}{\left| \frac{\partial \vec{r}}{\partial r} \right|} = \begin{pmatrix} \sin \vartheta \cos \varphi \\ \sin \vartheta \sin \varphi \\ \cos \vartheta \end{pmatrix}, \quad \vec{e}_\vartheta = \frac{\frac{\partial \vec{r}}{\partial \vartheta}}{\left| \frac{\partial \vec{r}}{\partial \vartheta} \right|} = \begin{pmatrix} \cos \vartheta \cos \varphi \\ \cos \vartheta \sin \varphi \\ -\sin \vartheta \end{pmatrix}, \quad \vec{e}_\varphi = \frac{\frac{\partial \vec{r}}{\partial \varphi}}{\left| \frac{\partial \vec{r}}{\partial \varphi} \right|} = \begin{pmatrix} -\sin \varphi \\ \cos \varphi \\ 0 \end{pmatrix}. \quad (5.1)$$

- Jacobi-Matrix $J = (\vec{g}_r, \vec{g}_\vartheta, \vec{g}_\varphi)$:

$$\begin{aligned} J &= \begin{pmatrix} \frac{\partial(x, y, z)}{\partial(r, \vartheta, \varphi)} \end{pmatrix} = \begin{pmatrix} \frac{\partial x}{\partial r} & \frac{\partial x}{\partial \vartheta} & \frac{\partial x}{\partial \varphi} \\ \frac{\partial y}{\partial r} & \frac{\partial y}{\partial \vartheta} & \frac{\partial y}{\partial \varphi} \\ \frac{\partial z}{\partial r} & \frac{\partial z}{\partial \vartheta} & \frac{\partial z}{\partial \varphi} \end{pmatrix} = \begin{pmatrix} \sin \vartheta \cos \varphi & r \cos \vartheta \cos \varphi & -r \sin \vartheta \sin \varphi \\ \sin \vartheta \sin \varphi & r \cos \vartheta \sin \varphi & r \sin \vartheta \cos \varphi \\ \cos \vartheta & -r \sin \vartheta & 0 \end{pmatrix}, \\ J^{-1} &= \begin{pmatrix} \frac{\partial(r, \vartheta, \varphi)}{\partial(x, y, z)} \end{pmatrix} = \begin{pmatrix} \frac{\partial r}{\partial x} & \frac{\partial r}{\partial y} & \frac{\partial r}{\partial z} \\ \frac{\partial \vartheta}{\partial x} & \frac{\partial \vartheta}{\partial y} & \frac{\partial \vartheta}{\partial z} \\ \frac{\partial \varphi}{\partial x} & \frac{\partial \varphi}{\partial y} & \frac{\partial \varphi}{\partial z} \end{pmatrix} = \begin{pmatrix} \sin \vartheta \cos \varphi & \sin \vartheta \sin \varphi & \cos \vartheta \\ \frac{1}{r} \cos \vartheta \cos \varphi & \frac{1}{r} \cos \vartheta \sin \varphi & -\frac{1}{r} \sin \vartheta \\ -\frac{1}{r} \frac{\sin \varphi}{\sin \vartheta} & \frac{1}{r} \frac{\cos \varphi}{\sin \vartheta} & 0 \end{pmatrix} \\ &= \begin{pmatrix} \frac{x}{r} & \frac{y}{r} & \frac{z}{r} \\ \frac{xz}{r^2 \sqrt{x^2 + y^2}} & \frac{yz}{r^2 \sqrt{x^2 + y^2}} & \frac{-(x^2 + y^2)}{r^2 \sqrt{x^2 + y^2}} \\ \frac{-y}{x^2 + y^2} & \frac{x}{x^2 + y^2} & 0 \end{pmatrix}. \end{aligned}$$

- Drehmatrix (Rotationsmatrix) D :

$$J = D \cdot h = \underbrace{\begin{pmatrix} \sin \vartheta \cos \varphi & \cos \vartheta \cos \varphi & -\sin \varphi \\ \sin \vartheta \sin \varphi & \cos \vartheta \sin \varphi & \cos \varphi \\ \cos \vartheta & -\sin \vartheta & 0 \end{pmatrix}}_D \cdot \underbrace{\begin{pmatrix} 1 & 0 & 0 \\ 0 & r & 0 \\ 0 & 0 & r \sin \vartheta \end{pmatrix}}_h.$$

- Funktionaldeterminante $\det J = r^2 \sin \vartheta$.

- Metriktenor: $(g_{kl}) = J^T \cdot J = h^2 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & r^2 & 0 \\ 0 & 0 & r^2 \sin^2 \vartheta \end{pmatrix}$.

- Transformation der Basiseinheitsvektoren:

$$(\vec{e}_x, \vec{e}_y, \vec{e}_z) \cdot D = (\vec{e}_r, \vec{e}_\vartheta, \vec{e}_\varphi) \Leftrightarrow (\vec{e}_r, \vec{e}_\vartheta, \vec{e}_\varphi) \cdot D^T = (\vec{e}_x, \vec{e}_y, \vec{e}_z).$$

- Transformation von Differentialen:

$$J^{-1} \cdot \begin{pmatrix} dx \\ dy \\ dz \end{pmatrix} = \begin{pmatrix} dr \\ d\vartheta \\ d\varphi \end{pmatrix} \Leftrightarrow J \cdot \begin{pmatrix} dr \\ d\vartheta \\ d\varphi \end{pmatrix} = \begin{pmatrix} dx \\ dy \\ dz \end{pmatrix}.$$

- Transformation partieller Ableitungen:

$$\left(\frac{\partial}{\partial x}, \frac{\partial}{\partial y}, \frac{\partial}{\partial z} \right) \cdot J = \left(\frac{\partial}{\partial r}, \frac{\partial}{\partial \vartheta}, \frac{\partial}{\partial \varphi} \right) \Leftrightarrow \left(\frac{\partial}{\partial r}, \frac{\partial}{\partial \vartheta}, \frac{\partial}{\partial \varphi} \right) \cdot J^{-1} = \left(\frac{\partial}{\partial x}, \frac{\partial}{\partial y}, \frac{\partial}{\partial z} \right).$$

- Transformation eines Vektorfeldes $\vec{F}(\vec{r})$:

$$D^T \cdot \begin{pmatrix} F_x \\ F_y \\ F_z \end{pmatrix} = \begin{pmatrix} F_r \\ F_\vartheta \\ F_\varphi \end{pmatrix} \Leftrightarrow D \cdot \begin{pmatrix} F_r \\ F_\vartheta \\ F_\varphi \end{pmatrix} = \begin{pmatrix} F_x \\ F_y \\ F_z \end{pmatrix},$$

$$\vec{F} = F_x \vec{e}_x + F_y \vec{e}_y + F_z \vec{e}_z = F_r \vec{e}_r + F_\vartheta \vec{e}_\vartheta + F_\varphi \vec{e}_\varphi.$$

- Volumenelement: $dV = dS \cdot dr = |\det J| \cdot dr d\vartheta d\varphi = r^2 \sin \vartheta \cdot dr d\vartheta d\varphi$.

- Flächenelement auf der Kugeloberfläche mit Radius $r = \text{const}$:

$$dS = \frac{dV}{dr} = |\det J| \cdot d\vartheta d\varphi = r^2 \sin \vartheta \cdot d\vartheta d\varphi.$$

- Linienelement ds :

$$(d\vec{r})^2 = (ds)^2 = g_r^2 (dr)^2 + g_\vartheta^2 (d\vartheta)^2 + g_\varphi^2 (d\varphi)^2 = (dr)^2 + r^2 (d\vartheta)^2 + r^2 \sin^2 \vartheta (d\varphi)^2.$$

- Gradient: $\text{grad } f(r, \vartheta, \varphi) = \frac{\partial f}{\partial r} \vec{e}_r + \frac{1}{r} \frac{\partial f}{\partial \vartheta} \vec{e}_\vartheta + \frac{1}{r \sin \vartheta} \frac{\partial f}{\partial \varphi} \vec{e}_\varphi$.

- Divergenz: $\text{div } \vec{F}(r, \vartheta, \varphi) = \frac{1}{r^2} \frac{\partial}{\partial r} (r^2 \cdot F_r) + \frac{1}{r \sin \vartheta} \frac{\partial}{\partial \vartheta} (\sin \vartheta \cdot F_\vartheta) + \frac{1}{r \sin \vartheta} \frac{\partial F_\varphi}{\partial \varphi}$.

- Rotation: $\text{rot } \vec{F}(r, \vartheta, \varphi) = \left[\frac{1}{r \sin \vartheta} \left(\frac{\partial}{\partial \vartheta} (\sin \vartheta \cdot F_\varphi) - \frac{\partial F_\vartheta}{\partial \varphi} \right) \right] \vec{e}_r$
 $+ \left[\frac{1}{r \sin \vartheta} \frac{\partial F_r}{\partial \varphi} - \frac{1}{r} \frac{\partial}{\partial r} (r \cdot F_\varphi) \right] \vec{e}_\vartheta$
 $+ \left[\frac{1}{r} \frac{\partial}{\partial r} (r \cdot F_\vartheta) - \frac{1}{r} \frac{\partial F_r}{\partial \vartheta} \right] \vec{e}_\varphi.$

- Laplacian: $\Delta f(r, \vartheta, \varphi) = \underbrace{\frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \cdot \frac{\partial f}{\partial r} \right)}_{\text{Radialableitung}} + \frac{1}{r^2 \sin \vartheta} \frac{\partial}{\partial \vartheta} \left(\sin \vartheta \cdot \frac{\partial f}{\partial \vartheta} \right) + \frac{1}{r^2 \sin^2 \vartheta} \frac{\partial^2 f}{\partial \varphi^2},$

$$\frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \cdot \frac{\partial f}{\partial r} \right) = \frac{\partial^2 f}{\partial r^2} + \frac{2}{r} \frac{\partial f}{\partial r} = \frac{1}{r} \frac{\partial^2}{\partial r^2} (r \cdot f).$$

6 Linien-, Flächen-, Volumenelement

6.1 Linienelement

Weil Linien oder Kurven eindimensionale geometrische Objekte sind, benötigt man für ihre Beschreibung nur *einen* freien Parameter, den Kurvenparameter, auch wenn es sich beispielsweise um eine Raumkurve (im 3-dimensionalen Raum) handelt.

6.1.1 Linienelement in kartesischen Koordinaten

Im kartesischen Koordinatensystem kann folglich jeder Punkt einer Linie durch den zugehörigen parametrisierten Ortsvektor dargestellt werden, z. B.

$$\vec{r}(t) = \begin{pmatrix} x(t) \\ y(t) \\ z(t) \end{pmatrix} \text{ mit Parameter } t, \quad \vec{r}(x) = \begin{pmatrix} x \\ y(x) \\ z(x) \end{pmatrix} \text{ mit Parameter } x.$$

Der Parameter t kann z. B. die Zeit sein und der Parameter x ist die Ortskoordinate x selbst. Aus dem Differentialquotienten

$$\frac{d\vec{r}(t)}{dt} = \begin{pmatrix} dx(t)/dt \\ dy(t)/dt \\ dz(t)/dt \end{pmatrix}$$

nach dem Parameter t resultiert das

$$\text{vektorielle Linienelement } d\vec{r}(t) = \frac{d\vec{r}(t)}{dt} \cdot dt = \begin{pmatrix} dx(t)/dt \\ dy(t)/dt \\ dz(t)/dt \end{pmatrix} \cdot dt.$$

Die Arbeit W , die eine Kraft $\vec{F}(\vec{r})$ längs eines Weges $\vec{r}(t)$ verrichtet ist damit das **Arbeitsintegral**

$$\begin{aligned} \int \vec{F} \cdot d\vec{r} &= \int \begin{pmatrix} F_x(x(t), y(t), z(t)) \\ F_y(x(t), y(t), z(t)) \\ F_z(x(t), y(t), z(t)) \end{pmatrix} \cdot \begin{pmatrix} dx(t)/dt \\ dy(t)/dt \\ dz(t)/dt \end{pmatrix} \cdot dt \\ &= \int F_x(\vec{r}(t)) \frac{dx(t)}{dt} \cdot dt + \int F_y(\vec{r}(t)) \frac{dy(t)}{dt} \cdot dt + \int F_z(\vec{r}(t)) \frac{dz(t)}{dt} \cdot dt. \end{aligned}$$

Interessant und wichtig ist, dass $d\vec{r}(t)/dt$, d. h. die Ableitung der Kurvenfunktion nach dem Kurvenparameter, der Tangentenvektor an die Kurve im Punkt $\vec{r}(t)$ ist.¹

Das (**skalare**) **Linienelement** $ds = ds(t)$ resultiert aus der Ableitung der Kurvenlängenfunktion $s(t)$ nach dem Parameter t wie folgt:

$$\begin{aligned} \text{Kurvenstücklänge} &= s(t) \Big|_{t_1}^{t_2} \approx \sum_i \sqrt{(\Delta x_i(t))^2 + (\Delta y_i(t))^2 + (\Delta z_i(t))^2} \Big|_{t_1}^{t_2} \\ &= \lim_{\Delta t \rightarrow 0} \sum_i \sqrt{\left(\frac{\Delta x_i(t)}{\Delta t}\right)^2 + \left(\frac{\Delta y_i(t)}{\Delta t}\right)^2 + \left(\frac{\Delta z_i(t)}{\Delta t}\right)^2} \cdot \Delta t \Big|_{t_1}^{t_2} \\ &= \int_{t_1}^{t_2} \sqrt{\left(\frac{dx(t)}{dt}\right)^2 + \left(\frac{dy(t)}{dt}\right)^2 + \left(\frac{dz(t)}{dt}\right)^2} \cdot dt = \int_{t_1}^{t_2} ds(t). \end{aligned}$$

¹Siehe dazu: Lothar Papula, Mathematik für Ingenieure und Naturwissenschaftler Band 3, Vieweg-Verlag, Braunschweig, Wiesbaden, 4. Auflage, 2001, Abschnitt 1.2 *Differentiation eines Vektors nach einem Parameter*, Seite 4 bis Seite 20.

Mit $\sqrt{\left(\frac{dx}{dt}\right)^2 + \left(\frac{dy}{dt}\right)^2 + \left(\frac{dz}{dt}\right)^2} = \sqrt{\frac{d\vec{r}}{dt} \cdot \frac{d\vec{r}}{dt}} = \left|\frac{d\vec{r}}{dt}\right|$ folgt daraus

$$\boxed{ds = \left|\frac{d\vec{r}}{dt}\right| \cdot dt = |d\vec{r}| \Leftrightarrow \frac{ds(t)}{dt} = \left|\frac{d\vec{r}(t)}{dt}\right|}.$$

Das Linien- oder Kurvenintegral längs $\vec{r}(t)$ in einem skalaren Feld $f(x, y, z)$ ist damit

$$\int f \cdot ds = \int f(x(t), y(t), z(t)) \cdot \left|\frac{d\vec{r}(t)}{dt}\right| \cdot dt.$$

Allgemein, also beispielsweise für krummlinig-schiefwinklige Koordinatensysteme, wird das Linienelement ds wie folgt im Tensorkalkül dargestellt:²

$$(ds)^2 = g_{kl} du^k du^l.$$

Dabei ist g_{kl} der metrische Tensor oder Metriktensor, der sich aus dem verwendeten Koordinatensystem mit der Jakobi-Matrix J durch

$$J^T \cdot J = h^2 = (g_{kl}) \longleftrightarrow g_{kl}$$

ergibt. Wenn man die Matrix (g_{kl}) in einen Tensor umschreibt, erhält man den Metriktensor g_{kl} , der die Metrik des Raums beschreibt.

6.1.2 Linienelement in Polarkoordinaten

Im Folgenden verwenden wir gelegentlich die Laufindizes $i \in \{1, 2, 3\}$ und $j \in \{1, 2, 3\}$ am Koordinatensymbol u , sodass $u \equiv u_1$, $v \equiv u_2$, $w \equiv u_3$.

Jetzt leiten wir das Linienelement für krummlinig-orthogonale Koordinatensysteme her. Dies betrifft insbesondere die in der Praxis wichtigen **Polarkoordinaten**. Dazu gehören insbesondere

- ebene Polarkoordinaten oder kurz Polarkoordinaten,
- Zylinderkoordinaten,
- sphärische Polarkoordinaten oder Kugelkoordinaten.

Hierbei handelt es sich um spezielle (u, v, w) - bzw. (u_1, u_2, u_3) -Koordinatensysteme, deren Koordinatenlinien $(u, v, w) \equiv (u_1, u_2, u_3)$ und folglich auch ihre Basiseinheitsvektoren

$$\vec{e}_{u_1} \equiv \vec{e}_u = \frac{\frac{\partial \vec{r}}{\partial u}}{\left|\frac{\partial \vec{r}}{\partial u}\right|}, \quad \vec{e}_{u_2} \equiv \vec{e}_v = \frac{\frac{\partial \vec{r}}{\partial v}}{\left|\frac{\partial \vec{r}}{\partial v}\right|}, \quad \vec{e}_{u_3} \equiv \vec{e}_w = \frac{\frac{\partial \vec{r}}{\partial w}}{\left|\frac{\partial \vec{r}}{\partial w}\right|}$$

stets senkrecht aufeinander stehen, sodass gilt:

$$\vec{e}_{u_i} \cdot \vec{e}_{u_j} = \delta_{ij} = \begin{cases} 0 & \text{für } i \neq j, \\ 1 & \text{für } i = j. \end{cases}$$

Außerdem bilden die Basiseinheitsvektoren in der Reihenfolge u, v, w bzw. u_1, u_2, u_3 ein Rechtssystem, also insgesamt ein orthogonales Rechtssystem bzw. ein rechtwinkliges rechtshändiges Dreibein gemäß

$$\vec{e}_i \cdot (\vec{e}_{u_j} \times \vec{e}_{u_k}) = \varepsilon_{ijk}.$$

Aus

$$d\vec{r} = \frac{\partial \vec{r}}{\partial u} du + \frac{\partial \vec{r}}{\partial v} dv + \frac{\partial \vec{r}}{\partial w} dw = \vec{e}_u \left|\frac{\partial \vec{r}}{\partial u}\right| du + \vec{e}_v \left|\frac{\partial \vec{r}}{\partial v}\right| dv + \vec{e}_w \left|\frac{\partial \vec{r}}{\partial w}\right| dw$$

²Im Abschnitt 6.8 *Der metrische Tensor (Metriktensor)* meines Skripts *Grundlegendes zur Elektrodynamik und Quantenmechanik* ist das Linienelementquadrat mit dem Metriktensor vollständig ausgeschrieben. Um derartige „Monstren“ zu handhaben, ist der Tensorkalkül entwickelt worden. Es ist gut zu wissen, was sich hinter derartigen Tensoren verbirgt.

und mit den metrischen Koeffizienten

$$g_u = \left| \frac{\partial \vec{r}}{\partial u} \right|, \quad g_v = \left| \frac{\partial \vec{r}}{\partial v} \right|, \quad g_w = \left| \frac{\partial \vec{r}}{\partial w} \right|$$

folgt das **vektorielle Linienelement** (6.1)

$$d\vec{r} = g_u du \cdot \vec{e}_u + g_v dv \cdot \vec{e}_v + g_w dw \cdot \vec{e}_w \Rightarrow \quad (6.1)$$

$$(d\vec{r})^2 = (ds)^2 = g_u^2 du^2 + g_v^2 dv^2 + g_w^2 dw^2. \quad (6.2)$$

Durch Radizieren des Linienelementquadrates (6.2) resultiert das (**skalare**) **Linienelement**

$$ds = \sqrt{g_u^2 du^2 + g_v^2 dv^2 + g_w^2 dw^2}. \quad (6.3)$$

Die Darstellung einer Kurve erfolgt beispielsweise durch die Parametrisierung – hier mit dem Parameter t – des Ortsvektors gemäß

$$\vec{r}(u, v, w) \longrightarrow \vec{r}(u(t), v(t), w(t)).$$

Bei der Bildung von Kurvenintegralen längs der Kurve $\vec{r}(t)$ benötigen wir das parametrisierte Linienelement. Dabei ist zu beachten, dass hier die metrischen Koeffizienten $g_{u_i}(t)$ längs der zugehörigen Koordinaten $u_i(t)$ konstant sind und folglich hinsichtlich der Ableitung der Differentiale $du_i(t)$ nach dem Parameter t als konstante Faktoren betrachtet werden können. Zwar sind die Basiseinheitsvektoren bei krummlinig-orthogonalen Koordinatensystemen ortsabhängig und werden deshalb parametrisiert, doch ergeben sie im Skalarprodukt unter dem Integral $\vec{e}_{u_i}(t) \cdot \vec{e}_{u_j}(t) = 0$ für $i \neq j$ und $\vec{e}_{u_i}(t) \cdot \vec{e}_{u_i}(t) = 1$. Im Übrigen erfolgt die Parametrisierung völlig analog zu der in kartesischen Koordinaten:

$$d\vec{r}(t) = \left(g_u(t) \frac{du(t)}{dt} \cdot \vec{e}_u(t) + g_v(t) \frac{dv(t)}{dt} \cdot \vec{e}_v(t) + g_w(t) \frac{dw(t)}{dt} \cdot \vec{e}_w(t) \right) \cdot dt,$$

$$ds(t) = \sqrt{g_u^2(t) \left(\frac{du(t)}{dt} \right)^2 + g_v^2(t) \left(\frac{dv(t)}{dt} \right)^2 + g_w^2(t) \left(\frac{dw(t)}{dt} \right)^2} \cdot dt.$$

Beispiel:

Betrachten wir wieder das **Arbeitsintegral**, aber diesmal in Kugelkoordinaten. Ist das Kraftfeld $\vec{F}(\vec{r})$ in kartesischen Koordinaten gegeben, müssen wir es zunächst in Kugelkoordinaten transformieren. Dies zeigen wir am Kraftfeld

$$\vec{F}(x, y, z) = \begin{pmatrix} x - yz \\ y + xz \\ z \end{pmatrix}.$$

Mit

$$\begin{aligned} x &= r \cdot \sin \vartheta \cdot \cos \varphi, \\ y &= r \cdot \sin \vartheta \cdot \sin \varphi, \\ z &= r \cdot \cos \vartheta \end{aligned}$$

ersetzen wir die kartesischen Koordinaten in $\vec{F}(x, y, z)$ und erhalten die folgende Darstellung in Kugelkoordinaten:

$$\vec{F}(r, \vartheta, \varphi) = \begin{pmatrix} r \sin \vartheta \cos \varphi - r \sin \vartheta \sin \varphi \cdot r \cos \vartheta \\ r \sin \vartheta \sin \varphi + r \sin \vartheta \cos \varphi \cdot r \cos \vartheta \\ r \cos \vartheta \end{pmatrix} \quad (6.4)$$

$$= r \cdot \underbrace{\begin{pmatrix} \sin \vartheta \cos \varphi \\ \sin \vartheta \sin \varphi \\ \cos \vartheta \end{pmatrix}}_{=\vec{e}_r} + r^2 \sin \vartheta \cdot \underbrace{\begin{pmatrix} -\sin \varphi \\ \cos \varphi \\ 0 \end{pmatrix}}_{=\vec{e}_\varphi} \Rightarrow$$

$$\begin{aligned} \vec{F}(r, \vartheta, \varphi) &= (r \sin \vartheta \cos \varphi - r^2 \sin \vartheta \sin \varphi \cos \vartheta) \vec{e}_x + (r \sin \vartheta \sin \varphi + r^2 \sin \vartheta \cos \varphi \cos \vartheta) \vec{e}_y + r \cos \vartheta \vec{e}_z \\ &= r \cdot \vec{e}_r + r^2 \sin \vartheta \cdot \vec{e}_\varphi \\ &= F_r \cdot \vec{e}_r + F_\varphi \cdot \vec{e}_\varphi. \end{aligned} \quad (6.5)$$

An diesem Beispiel lässt sich leicht überprüfen, dass man die skalaren Vektorkomponenten F_r , F_ϑ , F_φ von $\vec{F}$ (in Kugelkoordinaten) entweder durch

$$D^T \cdot \begin{pmatrix} F_x \\ F_y \\ F_z \end{pmatrix} = \begin{pmatrix} \sin \vartheta \cos \varphi & \sin \vartheta \sin \varphi & \cos \vartheta \\ \cos \vartheta \cos \varphi & \cos \vartheta \sin \varphi & -\sin \vartheta \\ -\sin \varphi & \cos \varphi & 0 \end{pmatrix} \cdot \begin{pmatrix} r \sin \vartheta \cos \varphi - r \sin \vartheta \sin \varphi \cdot r \cos \vartheta \\ r \sin \vartheta \sin \varphi + r \sin \vartheta \cos \varphi \cdot r \cos \vartheta \\ r \cos \vartheta \end{pmatrix} = \begin{pmatrix} F_r \\ F_\vartheta \\ F_\varphi \end{pmatrix}$$

oder durch die Projektion des Vektors $\vec{F}(r, \vartheta, \varphi)$ in Darstellung (6.4) auf die Basiseinheitsvektoren $\vec{e}_r$, $\vec{e}_\vartheta$, $\vec{e}_\varphi$ in Darstellung (5.1) berechnen kann:

$$\vec{F}(r, \vartheta, \varphi) \cdot \vec{e}_r = F_r, \quad \vec{F}(r, \vartheta, \varphi) \cdot \vec{e}_\vartheta = F_\vartheta, \quad \vec{F}(r, \vartheta, \varphi) \cdot \vec{e}_\varphi = F_\varphi.$$

Mit dem (6.1) entsprechenden vektoriellen Linienelement für Kugelkoordinaten

$$\begin{aligned} d\vec{r}(u, v, w) &= g_u \cdot du \cdot \vec{e}_u + g_v \cdot dv \cdot \vec{e}_v + g_w \cdot dw \cdot \vec{e}_w \Rightarrow \\ d\vec{r}(r, \vartheta, \varphi) &= 1 \cdot dr \cdot \vec{e}_r + r \cdot d\vartheta \cdot \vec{e}_\vartheta + r \sin \vartheta \cdot d\varphi \cdot \vec{e}_\varphi \end{aligned}$$

können wir jetzt das Arbeitsintegral in Kugelkoordinaten bilden:

$$\begin{aligned} \int \vec{F} \cdot d\vec{r} &= \int (F_r \cdot \vec{e}_r + F_\vartheta \cdot \vec{e}_\vartheta + F_\varphi \cdot \vec{e}_\varphi) \cdot (1 \cdot dr \cdot \vec{e}_r + r \cdot d\vartheta \cdot \vec{e}_\vartheta + r \sin \vartheta \cdot d\varphi \cdot \vec{e}_\varphi) \\ &= \int (r \cdot \vec{e}_r + 0 \cdot \vec{e}_\vartheta + r^2 \sin \vartheta \cdot \vec{e}_\varphi) \cdot (1 \cdot dr \cdot \vec{e}_r + r \cdot d\vartheta \cdot \vec{e}_\vartheta + r \sin \vartheta \cdot d\varphi \cdot \vec{e}_\varphi). \end{aligned}$$

Die Parametrisierung mit dem Kurvenparameter t liefert daraus

$$\begin{aligned} \int \vec{F}(\vec{r}(t)) \cdot d\vec{r}(t) &= \int (r(t) \cdot \vec{e}_r(t) + 0 \cdot \vec{e}_\vartheta(t) + (r(t))^2 \sin \vartheta(t) \cdot \vec{e}_\varphi(t)) \\ &\quad \cdot \left(1 \cdot \frac{dr(t)}{dt} \cdot \vec{e}_r(t) + r(t) \cdot \frac{d\vartheta(t)}{dt} \cdot \vec{e}_\vartheta(t) + r(t) \sin \vartheta(t) \cdot \frac{d\varphi(t)}{dt} \cdot \vec{e}_\varphi(t) \right) \cdot dt. \end{aligned}$$

Weil im Skalarprodukt (unter dem Integral) stets $\vec{e}_{u_i}(t) \cdot \vec{e}_{u_i}(t) = 1$ gilt, ist das Arbeitsintegral in unserem Beispielkraftfeld für ein Kurvenstück von $\vec{r}(t_1)$ bis $\vec{r}(t_2)$

$$\int_{\vec{r}(t_1)}^{\vec{r}(t_2)} \vec{F} \cdot d\vec{r} = \int_{t_1}^{t_2} \left(r(t) \frac{dr(t)}{dt} + (r(t))^3 \sin^2 \vartheta(t) \frac{d\varphi(t)}{dt} \right) \cdot dt.$$

In Polarkoordinaten (u, v, w) ist das mit t parametrisierte Kurvenintegral in einem Vektorfeld $\vec{F}(u, v, w)$

$$\begin{aligned} \int \vec{F} \cdot d\vec{r} &= \int \vec{F}(u(t), v(t), w(t)) \cdot d\vec{r}(u(t), v(t), w(t)) \\ &= \left(F_u(t) \cdot g_u \frac{du(t)}{dt} + F_v(t) \cdot g_v \frac{dv(t)}{dt} + F_w(t) \cdot g_w \frac{dw(t)}{dt} \right) \cdot dt. \end{aligned}$$

In Polarkoordinaten ist das Linien- oder Kurvenintegral längs $\vec{r}(t)$ in einem skalaren Feld $f(u, v, w)$ mit dem skalaren Linienelement (6.3) nach Parametrisierung mit t

$$\begin{aligned} \int f \cdot ds &= \int f(u(t), v(t), w(t)) \cdot ds(u(t), v(t), w(t)) \\ &= \int f(u(t), v(t), w(t)) \cdot \sqrt{g_u^2 \left(\frac{du(t)}{dt} \right)^2 + g_v^2 \left(\frac{dv(t)}{dt} \right)^2 + g_w^2 \left(\frac{dw(t)}{dt} \right)^2} \cdot dt. \end{aligned}$$

6.2 Flächenelement

Flächen S werden durch 2 freie Parameter beschrieben, z. B. durch die Parameter bzw. Koordinaten u und v :

$$S := \vec{r}(u, v) .$$

Allgemein ist damit das **vektorielle Flächenelement**

$$d\vec{S}(u, v) = \underbrace{\left(\frac{\partial \vec{r}(u, v)}{\partial u} \times \frac{\partial \vec{r}(u, v)}{\partial v} \right)}_{\vec{n}} \cdot \underbrace{du dv}_{dA} = \vec{n} \cdot dA .$$

$\vec{n}$ ist der Normalenvektor auf S und es gilt $\vec{S} \uparrow \vec{n}$. Das (**skalare**) **Flächenelement** ist folglich

$$dS(u, v) = \left| \frac{\partial \vec{r}(u, v)}{\partial u} \times \frac{\partial \vec{r}(u, v)}{\partial v} \right| \cdot du dv = |\vec{n}| \cdot dA .$$

Weiterhin gilt mit dem Einheitsnormalenvektor $\vec{n}^0$

$$d\vec{S} = \frac{\vec{n}}{|\vec{n}|} \cdot |\vec{n}| dA = \vec{n}^0 \cdot dS$$

Flächenelement auf der Kugeloberfläche mit Radius R in kartesischen Koordinaten:

Als freie Parameter wählen wir die beiden kartesischen Koordinaten x und y :

$$\begin{aligned} \text{Kugel\ss{}che } S &:= \vec{r}(x, y) = \begin{pmatrix} x \\ y \\ \sqrt{R^2 - x^2 - y^2} \end{pmatrix} \Rightarrow \\ \vec{n} &= \frac{\partial \vec{r}(x, y)}{\partial x} \times \frac{\partial \vec{r}(x, y)}{\partial y} = \begin{pmatrix} x \cdot (R^2 - x^2 - y^2)^{-\frac{1}{2}} \\ y \cdot (R^2 - x^2 - y^2)^{-\frac{1}{2}} \\ 1 \end{pmatrix}, \quad |\vec{n}| = \frac{R}{\sqrt{R^2 - x^2 - y^2}}, \\ \vec{n}^0 &= \frac{\vec{n}}{|\vec{n}|} = \frac{1}{R} \cdot \begin{pmatrix} x \\ y \\ \sqrt{R^2 - x^2 - y^2} \end{pmatrix}. \end{aligned}$$

$$dS = |\vec{n}| \cdot dA = \frac{R}{\sqrt{R^2 - x^2 - y^2}} \cdot dx dy ,$$

$$d\vec{S} = \vec{n}^0 \cdot dS = \vec{n} \cdot dA = \begin{pmatrix} x \cdot (R^2 - x^2 - y^2)^{-\frac{1}{2}} \\ y \cdot (R^2 - x^2 - y^2)^{-\frac{1}{2}} \\ 1 \end{pmatrix} \cdot dx dy .$$

Flächenelement auf der Kugeloberfläche mit Radius R in Kugelkoordinaten:

Die freien Parameter der Kugeloberfläche sind die Kugelkoordinaten ϑ und φ :

$$\begin{aligned} \text{Kugel\ss{}che } S &:= \vec{r}(\vartheta, \varphi) = \begin{pmatrix} R \cdot \sin \vartheta \cos \varphi \\ R \cdot \sin \vartheta \sin \varphi \\ R \cdot \cos \vartheta \end{pmatrix} \Rightarrow \\ \vec{n} &= \frac{\partial \vec{r}(\vartheta, \varphi)}{\partial \vartheta} \times \frac{\partial \vec{r}(\vartheta, \varphi)}{\partial \varphi} = R^2 \begin{pmatrix} \sin^2 \vartheta \cos \varphi \\ \sin^2 \vartheta \sin \varphi \\ \sin \vartheta \cos \vartheta \end{pmatrix}, \quad |\vec{n}| = R^2 \sin \vartheta, \\ \vec{n}^0 &= \frac{\vec{n}}{|\vec{n}|} = \begin{pmatrix} \sin \vartheta \cos \varphi \\ \sin \vartheta \sin \varphi \\ \cos \vartheta \end{pmatrix}. \end{aligned}$$

$$dS = |\vec{n}| \cdot dA = R^2 \sin \vartheta \cdot d\vartheta d\varphi ,$$

$$d\vec{S} = \vec{n}^0 \cdot dS = \vec{n} \cdot dA = \begin{pmatrix} \sin \vartheta \cos \varphi \\ \sin \vartheta \sin \varphi \\ \cos \vartheta \end{pmatrix} R^2 \sin \vartheta \cdot d\vartheta d\varphi .$$

6.3 Volumenelement

Siehe auch:

Siegfried Großmann, Mathematischer Einführungskurs für die Physik, Teubner-Verlag, Stuttgart, Leipzig, 8. Auflage, 2000, Abschnitt 5.3.4. *Wechsel der Variablen*, Abschnitt 5.3.4.2. *Die Funktionaldeterminante* und Abschnitt 5.3.4.3. *Die Transformation von Flächenelementen*, Seite 203 bis Seite 207.

https://vhm.mathematik.uni-stuttgart.de/Vorlesungen/Mehrdimensionale_Integration/Folien_Mehrdimensionale_Integration.pdf

Volumina besitzen 3 freie Parameter. Es ist sofort klar, dass das Volumenelement

$$\text{in kartesischen Koordinaten } (x, y, z) : \quad dV = dx dy dz$$

ist. Und schaut man sich die Abbildungen zum Volumenelement im „Großmann“ oder im „Weltner“³ an, so ist ebenfalls intuitiv klar, dass gilt

$$\text{in Kugelkoordinaten } (r, \vartheta, \varphi) : \quad dV = dS \cdot dr = r^2 \sin \vartheta \cdot dr d\vartheta d\varphi .$$

Allgemein kann man sich das Volumenelement als infinitesimales Parallelepiped, auch Spat genannt, vorstellen, dessen Kanten gebildet werden von den infinitesimalen Tangentenvektoren an die Koordinatenlinien des betrachteten Koordinatensystems. Im (geradlinigen, global orthogonalen) kartesischen (x, y, z) -Koordinatensystem fallen die infinitesimalen Tangentenvektoren mit den Koordinatenlinien zusammen, sodass das Volumenelement dV ein Kubus mit dem Volumen $dV = dx \cdot dy \cdot dz$ ist.

Die Tangentenvektoren an die Koordinatenlinien nichtkartesischer (u, v, w) -Koordinatensysteme erhalten wir, falls die Koordinatentransformation $(x, y, z) = g(u, v, w)$ bzw.

$$\vec{r} = \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} x(u, v, w) \\ y(u, v, w) \\ z(u, v, w) \end{pmatrix}$$

bekannt ist, mit der Jacobi-Matrix J wie folgt:

$$J = \begin{pmatrix} \frac{\partial x}{\partial u} & \frac{\partial x}{\partial v} & \frac{\partial x}{\partial w} \\ \frac{\partial y}{\partial u} & \frac{\partial y}{\partial v} & \frac{\partial y}{\partial w} \\ \frac{\partial z}{\partial u} & \frac{\partial z}{\partial v} & \frac{\partial z}{\partial w} \end{pmatrix} = \left(\frac{\partial \vec{r}(u, v, w)}{\partial u}, \frac{\partial \vec{r}(u, v, w)}{\partial v}, \frac{\partial \vec{r}(u, v, w)}{\partial w} \right),$$

$$J = (\vec{g}_u, \vec{g}_v, \vec{g}_w) .$$

Die Vektoren $\vec{g}_u, \vec{g}_v, \vec{g}_w$ sind die Tangentenvektoren an die Koordinatenlinien u, v, w und bilden die Spalten in der Jacobi-Matrix. Der Betrag des Spatprodukts, also

$$|(\vec{g}_u \times \vec{g}_v) \cdot \vec{g}_w| = |\det J| ,$$

aus diesen Tangentenvektoren ergibt das Volumen des Parallelepipeds im Verhältnis zum Volumen des kartesischen Kubus' $(\partial \vec{r} / \partial x \times \partial \vec{r} / \partial y) \cdot \partial \vec{r} / \partial z$ mit $\vec{r}(x, y, z) = (x, y, z)$:

$$\left| \det \left(\frac{\partial \vec{r}(x, y, z)}{\partial x}, \frac{\partial \vec{r}(x, y, z)}{\partial y}, \frac{\partial \vec{r}(x, y, z)}{\partial z} \right) \right| = \left| \det \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \right| = 1 .$$

Der Betrag der Determinante der Jacobi-Matrix, d. h. **der Betrag der Funktionaldeterminante ist der Skalierungsfaktor für das Volumenelement** gemäß

$$dV = 1 \cdot dx dy dz = dx dy dz = |\det J| \cdot du dv dw ,$$

³Klaus Weltner, Springerlehrbuch, Mathematik für Physiker 2, Springer-Verlag, Berlin, Heidelberg, New York, 12. Auflage, 2001, Abschnitt 15.4.3 *Kugelkoordinaten*, Seite 51.

oder ausführlich geschrieben:

$$\begin{aligned} dV &= \left| \left(\frac{\partial \vec{r}}{\partial x} dx \times \frac{\partial \vec{r}}{\partial y} dy \right) \cdot \frac{\partial \vec{r}}{\partial z} dz \right| = 1 \cdot dx dy dz \\ &= \left| \left(\frac{\partial \vec{r}}{\partial u} du \times \frac{\partial \vec{r}}{\partial v} dv \right) \cdot \frac{\partial \vec{r}}{\partial w} dw \right| = |\det J| \cdot du dv dw . \end{aligned}$$

Zu beachten ist, dass in den Volumenelementen (und auch in den Flächenelementen) der **Betrag der Funktionaldeterminante** verwendet werden muss, um ein orientierungsabhängiges Vorzeichen zu vermeiden.

Eine vereinfachende Besonderheit gibt es bei lokal orthogonalen Koordinatensystemen. In diesem Fall sind die Spalten der Jacobi-Matrix orthogonal zueinander, sodass für den Betrag der Funktionaldeterminante

$$|\det J| = |\vec{g}_u| \cdot |\vec{g}_v| \cdot |\vec{g}_w| = g_u \cdot g_v \cdot g_w$$

gilt. Angewandt auf Kugelkoordinaten ergibt dies mit

$$g_r = \left| \frac{\partial}{\partial r} \begin{pmatrix} r \sin \vartheta \cos \varphi \\ r \sin \vartheta \sin \varphi \\ r \cos \vartheta \end{pmatrix} \right| = 1, \quad g_\vartheta = \left| \frac{\partial}{\partial \vartheta} \begin{pmatrix} r \sin \vartheta \cos \varphi \\ r \sin \vartheta \sin \varphi \\ r \cos \vartheta \end{pmatrix} \right| = r, \quad g_\varphi = \left| \frac{\partial}{\partial \varphi} \begin{pmatrix} r \sin \vartheta \cos \varphi \\ r \sin \vartheta \sin \varphi \\ r \cos \vartheta \end{pmatrix} \right| = r \sin \vartheta$$

das Volumenelement

$$dV = |\vec{g}_r| |\vec{g}_\vartheta| |\vec{g}_\varphi| \cdot dr d\vartheta d\varphi = 1 \cdot r \cdot r \sin \vartheta \cdot dr d\vartheta d\varphi = r^2 \sin \vartheta \cdot dr d\vartheta d\varphi .$$

6.4 Gelegentlich hilfreiche Ergänzungen

- **Zur Funktionaldeterminante bei Flächenelementen im $\mathbb{R}^2$:**

Flächen sind 2-parametrische geometrische Objekte, besitzen also zwei freie Parameter. Die Analogie zum Volumenelement im 3-dimensionalen Raum ist folglich das Flächenelement im 2-dimensionalen Raum bzw. in der Ebene mit der quadratischen **Jacobi-Matrix**

$$J = \begin{pmatrix} \frac{\partial x(u,v)}{\partial u} & \frac{\partial x(u,v)}{\partial v} \\ \frac{\partial y(u,v)}{\partial u} & \frac{\partial y(u,v)}{\partial v} \end{pmatrix} = \begin{pmatrix} x_u & x_v \\ y_u & y_v \end{pmatrix} = \begin{pmatrix} \frac{\partial(x,y)}{\partial(u,v)} \end{pmatrix}$$

bzw. der zugehörigen **Inversen**

$$J^{-1} = \begin{pmatrix} \frac{\partial u(x,y)}{\partial x} & \frac{\partial u(x,y)}{\partial y} \\ \frac{\partial v(x,y)}{\partial x} & \frac{\partial v(x,y)}{\partial y} \end{pmatrix} = \begin{pmatrix} u_x & u_y \\ v_x & v_y \end{pmatrix} = \begin{pmatrix} \frac{\partial(u,v)}{\partial(x,y)} \end{pmatrix}.$$

Das Flächenelement im $\mathbb{R}^2$ ist dann

$$dS(u,v) = \left| \frac{\partial \vec{r}(u,v)}{\partial u} \times \frac{\partial \vec{r}(u,v)}{\partial v} \right| \cdot du dv = |\det J| \cdot du dv = \left| \det \begin{pmatrix} \frac{\partial(x,y)}{\partial(u,v)} \end{pmatrix} \right| \cdot du dv,$$

was man sofort für den Fall ebener Polarkoordinaten verifizieren kann (siehe Abschnitt 5.1).

- **Beziehung zwischen Funktionaldeterminante und ihrer Inversen:**

Es gilt

$$dV(x,y,z) = dx dy dz = \left| \det \begin{pmatrix} \frac{\partial(x,y,z)}{\partial(u,v,w)} \end{pmatrix} \right| \cdot du dv dw = dV(u,v,w). \quad (6.6)$$

Und mit

$$d\tilde{V}(u,v,w) = du dv dw = \left| \det \begin{pmatrix} \frac{\partial(u,v,w)}{\partial(x,y,z)} \end{pmatrix} \right| \cdot dx dy dz$$

resultiert aus (6.6)

$$\begin{aligned} dx dy dz &= \underbrace{\left| \det \begin{pmatrix} \frac{\partial(x,y,z)}{\partial(u,v,w)} \end{pmatrix} \right| \cdot \left| \det \begin{pmatrix} \frac{\partial(u,v,w)}{\partial(x,y,z)} \end{pmatrix} \right|}_{=1} \cdot dx dy dz = dx dy dz \\ \Rightarrow \quad \left| \det \begin{pmatrix} \frac{\partial(u,v,w)}{\partial(x,y,z)} \end{pmatrix} \right| &= \frac{1}{\left| \det \begin{pmatrix} \frac{\partial(x,y,z)}{\partial(u,v,w)} \end{pmatrix} \right|} \end{aligned}$$

in Übereinstimmung mit

$$A \text{ invertierbar} \Rightarrow \det A^{-1} = (\det A)^{-1} = \frac{1}{\det A}.$$

Angewandt auf die Funktionaldeterminante $\det J$ hat dann diese Beziehung in vereinfachter Notation die folgende Gestalt:

$$\det J^{-1} = \begin{vmatrix} u_x & u_y & u_z \\ v_x & v_y & v_z \\ w_x & w_y & w_z \end{vmatrix} = \frac{1}{\begin{vmatrix} x_u & x_v & x_w \\ y_u & y_v & y_w \\ z_u & z_v & z_w \end{vmatrix}} = \frac{1}{\det J}.$$

Außerdem sieht man, dass

$$d\tilde{V}(u,v,w) = du dv dw \neq dx dy dz = |\det J| \cdot du dv dw = dV(u,v,w).$$

- **Zum Skalierungsfaktor bei Flächenelementen im $\mathbb{R}^3$:**

Weil Flächen nur zwei freie Parameter besitzen, gibt es für sie im $\mathbb{R}^3$ keine Funktionaldeterminante in der bereits diskutierten Form.⁴ Man kann aber auch für Flächen im $\mathbb{R}^3$ eine

⁴Siehe Siegfried Großmann, Mathematischer Einführungskurs für die Physik, Teubner-Verlag, Stuttgart, Leipzig, 8. Auflage, 2000, Abschnitt 5.3.4.2. *Die Funktionaldeterminante*, Seite 205.

Determinante konstruieren, die den richtigen Skalierungsfaktor liefert, wenn wir davon ausgehen, dass der Skalierungsfaktor der Betrag eines Spatprodukts

$$(\vec{a} \times \vec{b}) \cdot \vec{c} = (\vec{b} \times \vec{c}) \cdot \vec{a} = (\vec{c} \times \vec{a}) \cdot \vec{b}$$

sein soll, und wenn wir berücksichtigen, dass man ein Spatprodukt als Determinante 3. Ordnung schreiben kann:

$$\det \begin{pmatrix} a_x & b_x & c_x \\ a_y & b_y & c_y \\ a_z & b_z & c_z \end{pmatrix} = \det(\vec{a}, \vec{b}, \vec{c}) = (\vec{a} \times \vec{b}) \cdot \vec{c}.$$

Die Elemente der Determinante sind hier skalare Vektorkomponenten und dürfen nicht mit partiellen Ableitungen nach den Koordinaten x, y, z verwechselt werden. Wenn jetzt $\frac{\partial \vec{r}}{\partial u} = \vec{a}$ und $\frac{\partial \vec{r}}{\partial v} = \vec{b}$ die Tangentenvektoren an die Koordinatenlinien der Fläche sind und wenn $\vec{n}^0$ der nach außen gerichtete Einheitsnormalenvektor auf dieser Fläche ist gemäß

$$\vec{a} \times \vec{b} = \vec{n} \quad \Rightarrow \quad \vec{n}^0 = \frac{\vec{n}}{|\vec{n}|} = \frac{\vec{a} \times \vec{b}}{|\vec{a} \times \vec{b}|}, \quad |\vec{n}^0| = 1,$$

dann können wir damit ein Spatprodukt konstruieren, das uns den Skalierungsfaktor für das Flächenelement liefert:

$$|(\vec{a} \times \vec{b}) \cdot \vec{n}^0| = |\vec{n} \cdot \vec{n}^0| = |\vec{n}| \cdot |\vec{n}^0| = |\vec{n}| = |\vec{a} \times \vec{b}| \quad \Rightarrow$$

$$\left| \frac{\partial \vec{r}(u, v)}{\partial u} \times \frac{\partial \vec{r}(u, v)}{\partial v} \right| = \left| \left(\frac{\partial \vec{r}(u, v)}{\partial u} \times \frac{\partial \vec{r}(u, v)}{\partial v} \right) \cdot \vec{n}^0 \right| = \left| \det \left(\frac{\partial \vec{r}}{\partial u}, \frac{\partial \vec{r}}{\partial v}, \vec{n}^0 \right) \right|.$$

Sehr bequem lässt sich diese Determinante für die Kugeloberfläche ermitteln, denn Kugelkoordinaten sind lokal orthogonal. Deshalb steht der Vektor $\vec{g}_r$ stets senkrecht auf den Tangentenvektoren $\vec{g}_u$ und $\vec{g}_v$ an die Koordinatenlinien u und v der Kugeloberfläche:

$$\vec{g}_r = \begin{pmatrix} \sin \vartheta \cos \varphi \\ \sin \vartheta \sin \varphi \\ \cos \vartheta \end{pmatrix} = \vec{n}^0, \quad |\vec{g}_r| = g_r = 1.$$

Der Skalierungsfaktor für das Flächenelement der Kugeloberfläche mit $r = \text{const} = R$ ist damit

$$\begin{aligned} \left| \frac{\partial \vec{r}(\vartheta, \varphi)}{\partial \vartheta} \times \frac{\partial \vec{r}(\vartheta, \varphi)}{\partial \varphi} \right| &= \left| \det(\vec{g}_\vartheta, \vec{g}_\varphi, \vec{g}_r) \right| \\ &= \left| \det \begin{pmatrix} R \cos \vartheta \cos \varphi & -R \sin \vartheta \sin \varphi & \sin \vartheta \cos \varphi \\ R \cos \vartheta \sin \varphi & R \sin \vartheta \cos \varphi & \sin \vartheta \sin \varphi \\ -R \sin \vartheta & 0 & \cos \vartheta \end{pmatrix} \right| = R^2 \sin \vartheta. \end{aligned}$$

- Das **Raumwinkelement** $d\Omega$ ist

$$d\Omega := \sin \vartheta \, d\vartheta \, d\varphi,$$

sodass das Integral über den gesamten Raum $\mathbb{R}^3$ den **vollen Raumwinkel**

$$\Omega_{\mathbb{R}^3} = \int_{\mathbb{R}^3} d\Omega = \int_{\vartheta=0}^{\pi} \int_{\varphi=0}^{2\pi} \sin \vartheta \, d\vartheta \, d\varphi = 4\pi$$

ergibt. Der Raumwinkel Ω ist nämlich über die Gesamtoberfläche A_{OK} einer Kugel mit dem Radius R und über deren Teiloberflächenstücke A wie folgt definiert:

$$\Omega_{\mathbb{R}^3} = \frac{A_{\text{OK}}}{R^2} = \frac{4\pi R^2}{R^2} = 4\pi \quad \Rightarrow \quad \Omega := \frac{A}{R^2}.$$

7 Differentialoperatoren in krummlinig-orthogonalen Koordinatensystemen

(Nach dem Lehrbuch aus der Teubner Studienbücherei Physik: Siegfried Großmann, Mathematischer Einführungskurs für die Physik, Teubner-Verlag, Stuttgart, Leipzig, 8. Auflage, 2000, Abschnitt 7.2. *Differentialoperatoren in krummlinig-orthogonalen Koordinaten*, Seite 256 bis Seite 258.)

In diesem Kapitel verallgemeinern wir die Differentialoperatoren des geradlinig-orthogonalen (kartesischen) (x, y, z) -Koordinatensystems auf die Differentialoperatoren in krummlinig-orthogonalen (u, v, w) -Koordinatensystemen.

7.1 Gradient

In kartesischen Koordinaten ist der Gradient einer skalaren Funktion bzw. eines skalaren Feldes $f = f(\vec{r}) = f(x, y, z)$ definiert durch

$$\text{grad } f = \frac{\partial f}{\partial x} \cdot \vec{e}_x + \frac{\partial f}{\partial y} \cdot \vec{e}_y + \frac{\partial f}{\partial z} \cdot \vec{e}_z =: \frac{df}{d\vec{r}}. \quad (7.1)$$

Bei $\frac{df}{d\vec{r}}$ handelt es sich um eine symbolische Schreibweise, denn sie ist inkonsistent mit dem Vektorkalkül. Das Skalarprodukt

$$\text{grad } f \cdot d\vec{r} = df$$

beschreibt die Feldänderung *bei* Verschiebung $d\vec{r}$.

Der Gradient ist also ein Vektor, der die Änderung der Feldgröße f *pro* Verschiebung (Länge) beschreibt und in die Richtung der maximalen Änderung von f zeigt. Das aber bedeutet, dass die u -Komponente grad_u die Projektion des Gradienten(vektors) auf den Basiseinheitsvektor $\vec{e}_u$ ist:

$$\text{grad}_u = \text{grad } f \cdot \vec{e}_u. \quad (7.2)$$

Mit

$$\vec{e}_u = \frac{\frac{\partial \vec{r}}{\partial u}}{\left| \frac{\partial \vec{r}}{\partial u} \right|} = \frac{1}{g_u} \frac{\partial \vec{r}}{\partial u} = \frac{1}{g_u} \left(\frac{\partial x}{\partial u} \vec{e}_x + \frac{\partial y}{\partial u} \vec{e}_y + \frac{\partial z}{\partial u} \vec{e}_z \right)$$

resultiert aus (7.2) das Skalarprodukt bzw. die skalare Vektorkomponente des Gradienten von f in Richtung der u -Koordinate:

$$\begin{aligned} \text{grad}_u &= \left(\frac{\partial f}{\partial x} \cdot \vec{e}_x + \frac{\partial f}{\partial y} \cdot \vec{e}_y + \frac{\partial f}{\partial z} \cdot \vec{e}_z \right) \cdot \frac{1}{g_u} \left(\frac{\partial x}{\partial u} \vec{e}_x + \frac{\partial y}{\partial u} \vec{e}_y + \frac{\partial z}{\partial u} \vec{e}_z \right) \\ &= \frac{1}{g_u} \cdot \underbrace{\left(\frac{\partial f}{\partial x} \frac{\partial x}{\partial u} + \frac{\partial f}{\partial y} \frac{\partial y}{\partial u} + \frac{\partial f}{\partial z} \frac{\partial z}{\partial u} \right)}_{= \frac{\partial f}{\partial u}}. \end{aligned}$$

Der unterklammerte Term ergibt nicht $3 \frac{\partial f}{\partial u}$ sondern $\frac{\partial f}{\partial u}$ wie beim totalen Differential gemäß der **Kettenregel**, sodass wir für die skalare Vektorkomponente in u -Richtung

$$\text{grad}_u = \frac{1}{g_u} \cdot \frac{\partial f}{\partial u}$$

erhalten. Die skalaren Gradientenkomponenten in v - und in w -Richtung ergeben sich völlig analog. Mit diesen drei Komponenten ist dann der Gradient in den orthogonalen Koordinaten u, v, w

$$\text{grad } f(\vec{r}(u, v, w)) = \frac{1}{g_u} \frac{\partial f}{\partial u} \vec{e}_u + \frac{1}{g_v} \frac{\partial f}{\partial v} \vec{e}_v + \frac{1}{g_w} \frac{\partial f}{\partial w} \vec{e}_w.$$

Wählen wir für die allgemeinen orthogonalen Koordinaten u, v, w speziell die kartesischen Koordinaten x, y, z mit $g_x = g_y = g_z = 1$, so resultiert erwartungsgemäß (7.1).

7.2 Divergenz

Die folgende Herleitung ist analog zu der Herleitung der Divergenz für kartesische Koordinaten im Kapitel 15 meines *Basisskripts zur Infinitesimalrechnung*. Im Folgenden verwenden wir gelegentlich den Laufindex $i \in \{1, 2, 3\}$ am Koordinatensymbol u , sodass $u \equiv u_1$, $v \equiv u_2$, $w \equiv u_3$.

Wenn der infinitesimale Abstandsvektor in orthogonalen Koordinaten

$$d\vec{r} = \frac{\partial \vec{r}}{\partial u} du + \frac{\partial \vec{r}}{\partial v} dv + \frac{\partial \vec{r}}{\partial w} dw$$

ist und für die entsprechenden orthogonalen Basiseinheitsvektoren

$$\vec{e}_{u_i} = \frac{1}{g_{u_i}} \frac{d\vec{r}}{du_i} \Leftrightarrow \frac{d\vec{r}}{du_i} = \vec{e}_{u_i} g_{u_i}$$

gilt, dann können wir für den infinitesimalen Abstandsvektor unter Verwendung der Basiseinheitsvektoren

$$d\vec{r} = \vec{e}_u g_u du + \vec{e}_v g_v dv + \vec{e}_w g_w dw$$

schreiben und das Volumenelement ist

$$dV = g_u du \cdot g_v dv \cdot g_w dw = g_u g_v g_w \cdot du dv dw.$$

Und für das Volumen bei kleinen Abständen gilt dann

$$\Delta V \approx g_u g_v g_w \cdot \Delta u \Delta v \Delta w.$$

Für den Fluss des Vektors $\vec{F}(u, v, w) = (F_u(u, v, w), F_v(u, v, w), F_w(u, v, w))$ durch die einhüllende Fläche $S = \partial V$ eines quaderanalogen Bereichs V , dessen Kanten gebildet werden von den krummlinig-orthogonalen (u, v, w) -Koordinatenlinien, gilt

$$\begin{aligned} \oint_S \vec{F} \cdot \vec{n}^0 dS &= \int_{v_0}^{v_0+\Delta v} \int_{w_0}^{w_0+\Delta w} [F_u(u_0 + \Delta u, v, w) - F_u(u_0, v, w)] g_v dv g_w dw \\ &+ \int_{w_0}^{w_0+\Delta w} \int_{u_0}^{u_0+\Delta u} [F_v(u, v_0 + \Delta v, w) - F_v(u, v_0, w)] g_w dw g_u du \\ &+ \int_{u_0}^{u_0+\Delta u} \int_{v_0}^{v_0+\Delta v} [F_w(u, v, w_0 + \Delta w) - F_w(u, v, w_0)] g_u du g_v dv \\ &= \int_{v_0}^{v_0+\Delta v} \int_{w_0}^{w_0+\Delta w} [g_v g_w F_u(u_0 + \Delta u, v, w) - g_v g_w F_u(u_0, v, w)] dv dw \quad (7.3) \\ &+ \int_{w_0}^{w_0+\Delta w} \int_{u_0}^{u_0+\Delta u} [g_w g_u F_v(u, v_0 + \Delta v, w) - g_w g_u F_v(u, v_0, w)] dw du \\ &+ \int_{u_0}^{u_0+\Delta u} \int_{v_0}^{v_0+\Delta v} [g_u g_v F_w(u, v, w_0 + \Delta w) - g_u g_v F_w(u, v, w_0)] du dv. \end{aligned}$$

Die Taylor-Entwicklung beispielsweise des Integranden in (7.3) an der Stelle u_0 liefert

$$\begin{aligned} &[g_v g_w F_u(u_0 + \Delta u, v, w) - g_v g_w F_u(u_0, v, w)] \\ &= g_v g_w F_u(u_0, v, w) + \frac{\partial}{\partial u} [g_v g_w F_u(u_0, v, w)] \Delta u + \frac{1}{2} \frac{\partial^2}{\partial u^2} [g_v g_w F_u(u_0, v, w)] \Delta u^2 + \dots \\ &\quad - g_v g_w F_u(u_0, v, w), \end{aligned}$$

$$\left[g_v g_w F_u(u_0 + \Delta u, v, w) - g_v g_w F_u(u_0, v, w) \right] = \frac{\partial}{\partial u} \left[g_v g_w F_u(u_0, v, w) \right] \Delta u \dots$$

Die Glieder höherer Ordnung haben wir deshalb nicht mehr aufgeführt, weil alle Glieder der Ordnung > 1 bei der später erfolgenden Grenzwertbildung mit $\Delta u_i \rightarrow 0$ verschwinden. Diese Taylor-Entwicklung, die analog auch für die beiden übrigen Teilintegranden gilt, setzen wir in das Flussintegral ein:

$$\begin{aligned} \oint_S \vec{F} \cdot \vec{n}^0 dS &= \int_{v_0}^{v_0+\Delta v} \int_{w_0}^{w_0+\Delta w} \left\{ \frac{\partial}{\partial u} \left[g_v g_w F_u(u_0, v, w) \right] \Delta u \dots \right\} dv dw \\ &+ \int_{w_0}^{w_0+\Delta w} \int_{u_0}^{u_0+\Delta u} \left\{ \frac{\partial}{\partial v} \left[g_w g_u F_v(u, v_0, w) \right] \Delta v \dots \right\} dw du \\ &+ \int_{u_0}^{u_0+\Delta u} \int_{v_0}^{v_0+\Delta v} \left\{ \frac{\partial}{\partial w} \left[g_u g_v F_w(u, v, w_0) \right] \Delta w \dots \right\} du dv . \end{aligned}$$

Darauf wenden wir den **Mittelwertsatz der Integralrechnung**

$$\int_{x_0}^{x_1} f(x) dx = (x_1 - x_0) f(\tilde{x}) = f(\tilde{x}) \Delta x$$

an und erhalten

$$\begin{aligned} \oint_S \vec{F} \cdot \vec{n}^0 dS &= \frac{\partial}{\partial u} \left[g_v g_w F_u(u_0, \tilde{v}, \tilde{w}) \right] \Delta u \Delta v \Delta w \dots \\ &+ \frac{\partial}{\partial v} \left[g_w g_u F_v(\tilde{u}, v_0, \tilde{w}) \right] \Delta u \Delta v \Delta w \dots \\ &+ \frac{\partial}{\partial w} \left[g_u g_v F_w(\tilde{u}, \tilde{v}, w_0) \right] \Delta u \Delta v \Delta w \dots . \end{aligned}$$

Die mittlere Quelldichte von $\vec{F}$ für den quaderanalogen Bereich V ist folglich

$$\begin{aligned} \frac{1}{\Delta V} \oint_S \vec{F} \cdot \vec{n}^0 dS &\approx \frac{1}{g_u g_v g_w \cdot \Delta u \Delta v \Delta w} \oint_S \vec{F} \cdot \vec{n}^0 dS \approx \frac{1}{g_u g_v g_w} \frac{\partial}{\partial u} \left[g_v g_w F_u(u_0, \tilde{v}, \tilde{w}) \right] \dots \\ &+ \frac{1}{g_u g_v g_w} \frac{\partial}{\partial v} \left[g_w g_u F_v(\tilde{u}, v_0, \tilde{w}) \right] \dots \\ &+ \frac{1}{g_u g_v g_w} \frac{\partial}{\partial w} \left[g_u g_v F_w(\tilde{u}, \tilde{v}, w_0) \right] \dots . \end{aligned}$$

Wenn wir jetzt das Volumen ΔV gegen Null gehen lassen gemäß

$$\Delta V \rightarrow 0 \Rightarrow \begin{cases} \Delta u \rightarrow 0 & \tilde{u} \rightarrow u_0 \\ \Delta v \rightarrow 0 & \tilde{v} \rightarrow v_0 \\ \Delta w \rightarrow 0 & \tilde{w} \rightarrow w_0 \end{cases} ,$$

erhalten wir **exakt** die (lokale) Quelldichte bzw. die Divergenz von $\vec{F}$ für den Raumpunkt $\vec{r}_0 = (u_0, v_0, w_0)$:

$$\begin{aligned} \lim_{\Delta V \rightarrow 0} \frac{1}{\Delta V} \oint_S \vec{F} \cdot \vec{n}^0 dS &= \text{div } \vec{F}(\vec{r}_0) \\ &= \frac{1}{g_u g_v g_w} \left\{ \frac{\partial}{\partial u} \left[g_v g_w F_u(u_0, v_0, w_0) \right] + \frac{\partial}{\partial v} \left[g_w g_u F_v(u_0, v_0, w_0) \right] + \frac{\partial}{\partial w} \left[g_u g_v F_w(u_0, v_0, w_0) \right] \right\} . \end{aligned}$$

Weil aber diese Divergenzformel für jeden Raumpunkt $\vec{r}$ im Vektorfeld $\vec{F}(\vec{r})$ gelten soll, können wir $\vec{r}_0$ durch $\vec{r}$ ersetzen und erhalten schließlich

$$\text{div } \vec{F}(\vec{r}(u, v, w)) = \frac{1}{g_u g_v g_w} \left[\frac{\partial}{\partial u} (g_v g_w F_u) + \frac{\partial}{\partial v} (g_w g_u F_v) + \frac{\partial}{\partial w} (g_u g_v F_w) \right] .$$

7.3 Laplace-Operator

In kartesischen Koordinaten x, y, z ist der Laplace-Operator die Anwendung des Divergenz-Operators auf den Nabla-Operator („Gradient-Operator“) gemäß¹

$$\Delta := \operatorname{div} \operatorname{grad} = \vec{\nabla}^2 = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2}.$$

Dabei darf man das Laplace-Operator-Symbol Δ nicht mit dem Differenz-Delta verwechseln. Auf den Vektorpfeil über dem Nabla-Symbol $\vec{\nabla}$ für den Gradienten wird meistens verzichtet, weil von vornherein klar ist, dass es sich beim Gradienten um einen Vektor handelt. Der Laplace-Operator beinhaltet also zuerst die Bildung des Gradienten $\vec{F}(\vec{r})$ einer Funktion bzw. eines skalaren Feldes² $f(\vec{r})$ und daran anschließend die Bildung der Divergenz des Gradienten bzw. Vektorfeldes $\vec{F}(\vec{r})$. Das gilt in dieser Weise natürlich auch für krummlinig-orthogonale Koordinaten:

$$\operatorname{grad} f(u, v, w) = (F_u, F_v, F_w) = \vec{F}(u, v, w) \longrightarrow \operatorname{div} \vec{F}(u, v, w) = \Delta f(u, v, w).$$

Um den Laplace-Operator in krummlinig-orthogonalen Koordinaten zu erhalten, brauchen wir also nur die Ergebnisse der Abschnitte (7.1) und (7.2) zu kombinieren:

$$\begin{aligned} \operatorname{grad} f &= \frac{1}{g_u} \frac{\partial f}{\partial u} \vec{e}_u + \frac{1}{g_v} \frac{\partial f}{\partial v} \vec{e}_v + \frac{1}{g_w} \frac{\partial f}{\partial w} \vec{e}_w \\ &= F_u \vec{e}_u + F_v \vec{e}_v + F_w \vec{e}_w, \\ \operatorname{div} \vec{F} &= \frac{1}{g_u g_v g_w} \left[\frac{\partial}{\partial u} (g_v g_w F_u) + \frac{\partial}{\partial v} (g_w g_u F_v) + \frac{\partial}{\partial w} (g_u g_v F_w) \right] \Rightarrow \end{aligned}$$

$$\Delta f = \operatorname{div} \operatorname{grad} f(\vec{r}(u, v, w)) = \frac{1}{g_u g_v g_w} \left[\frac{\partial}{\partial u} \left(\frac{g_v g_w}{g_u} \frac{\partial f}{\partial u} \right) + \frac{\partial}{\partial v} \left(\frac{g_w g_u}{g_v} \frac{\partial f}{\partial v} \right) + \frac{\partial}{\partial w} \left(\frac{g_u g_v}{g_w} \frac{\partial f}{\partial w} \right) \right].$$

¹ $\vec{\nabla}^2 = \vec{\nabla} \cdot \vec{\nabla} = \left(\frac{\partial}{\partial x} \vec{e}_x + \frac{\partial}{\partial y} \vec{e}_y + \frac{\partial}{\partial z} \vec{e}_z \right) \cdot \left(\frac{\partial}{\partial x} \vec{e}_x + \frac{\partial}{\partial y} \vec{e}_y + \frac{\partial}{\partial z} \vec{e}_z \right) = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2}.$ □

²Es gibt auch den Vektorgradienten, d. h. den Gradienten eines Vektorfeldes. Das schon an dieser Stelle zu diskutieren, würde zu weit führen.

7.4 Rotation

Die folgende Herleitung ist analog zu der Herleitung der Rotation für kartesische Koordinaten im Kapitel 16 meines *Basisskripts zur Infinitesimalrechnung*. Im Folgenden verwenden wir gelegentlich den Laufindex $i \in \{1, 2, 3\}$ am Koordinatensymbol u , sodass $u \equiv u_1$, $v \equiv u_2$, $w \equiv u_3$.

Die Rotation ist ein Vektor bzw. ein Vektorfeld. In krummlinig-orthogonalen Koordinaten u, v, w besitzt die Rotation $\text{rot } \vec{F}$ des Vektorfeldes $\vec{F}(\vec{r}(u, v, w))$ die

$$\text{Vektorkomponenten } \vec{e}_{u_i} \cdot \text{rot}_{u_i} \vec{F}.$$

Das bedeutet, dass die

$$\text{rot}_{u_i} \vec{F} = (\text{rot } \vec{F}) \cdot \vec{e}_{u_i}$$

die skalaren Vektorkomponenten sind. Weil die Herleitung der einzelnen skalaren Vektorkomponenten der Rotation völlig analog zueinander ist, werden wir nur die w -Komponente der Rotation herleiten und die beiden übrigen Komponenten schlussfolgern.

Mit dem vektoriellen Linienelement

$$d\vec{r} = \vec{e}_u \cdot g_u du + \vec{e}_v \cdot g_v dv + \vec{e}_w \cdot g_w dw$$

gilt in krummlinig-orthogonalen Koordinaten für das vektorielle Linienintegral in einem Vektorfeld $\vec{F}(\vec{r})$ (siehe Arbeitsintegral im Abschnitt 6.1.2)

$$\int \vec{F} \cdot d\vec{r} = \int (F_u \cdot g_u du + F_v \cdot g_v dv + F_w \cdot g_w dw).$$

Betrachten wir in einem krummlinig-orthogonalen (u, v, w) -Koordinatensystem ein rechteckanalogen Flächenstück S_w . Dieses werde bei $w_0 = \text{const}$ im mathematisch positiven Umlaufsinn umrandet von der geschlossenen Kurve $\partial S_w = C_{uv}$, die wiederum aus den folgenden vier Teilkurven besteht:

von (u_0, v_0, w_0)	längs u bis	$(u_0 + \Delta u, v_0, w_0)$,
von $(u_0 + \Delta u, v_0, w_0)$	längs v bis	$(u_0 + \Delta u, v_0 + \Delta v, w_0)$,
von $(u_0 + \Delta u, v_0 + \Delta v, w_0)$	längs u bis	$(u_0, v_0 + \Delta v, w_0)$,
von $(u_0, v_0 + \Delta v, w_0)$	längs v bis	(u_0, v_0, w_0) .

Damit bilden wir das Linienintegral oder genauer gesagt das Umlaufintegral bzw. die Zirkulation von $\vec{F}$ längs C_{uv} :

$$\begin{aligned} \oint_{C_{uv}} \vec{F} \cdot d\vec{r} &= \int_{u_0}^{u_0 + \Delta u} F_u(u, v_0, w_0) \cdot g_u du + \int_{v_0}^{v_0 + \Delta v} F_v(u_0 + \Delta u, v, w_0) \cdot g_v dv \\ &\quad + \int_{u_0 + \Delta u}^{u_0} F_u(u, v_0 + \Delta v, w_0) \cdot g_u du + \int_{v_0 + \Delta v}^{v_0} F_v(u_0, v, w_0) \cdot g_v dv, \\ \oint_{C_{uv}} \vec{F} \cdot d\vec{r} &= \int_{v_0}^{v_0 + \Delta v} \left[g_v F_v(u_0 + \Delta u, v, w_0) - g_v F_v(u_0, v, w_0) \right] dv \\ &\quad - \int_{u_0}^{u_0 + \Delta u} \left[g_u F_u(u, v_0 + \Delta v, w_0) - g_u F_u(u, v_0, w_0) \right] du. \end{aligned}$$

Analog zur Herleitung der Divergenz liefert die Taylor-Entwicklung der Integranden

$$\oint_{C_{uv}} \vec{F} \cdot d\vec{r} = \int_{v_0}^{v_0 + \Delta v} \left\{ \frac{\partial}{\partial u} [g_v F_v(u_0, v, w_0)] \Delta u + \dots \right\} dv - \int_{u_0}^{u_0 + \Delta u} \left\{ \frac{\partial}{\partial v} [g_u F_u(u, v_0, w_0)] \Delta v + \dots \right\} du. \quad (7.4)$$

Auch bei dieser Taylor-Entwicklung brauchen wir die Glieder der Ordnung > 1 nicht zu berücksichtigen, weil diese am Ende bei der Grenzwertbildung $\Delta S \rightarrow 0$ verschwinden. Mit dem Mittelwertsatz der Integralrechnung schreiben wir für (7.4)

$$\oint_{C_{uv}} \vec{F} \cdot d\vec{r} = \left\{ \frac{\partial}{\partial u} [g_v F_v(u_0, \tilde{v}, w_0)] \Delta u + \dots \right\} \Delta v - \left\{ \frac{\partial}{\partial v} [g_u F_u(\tilde{u}, v_0, w_0)] \Delta v + \dots \right\} \Delta u .$$

Die Division der Zirkulation durch den Flächeninhalt³

$$\Delta S_w \approx g_u g_v \cdot \Delta u \Delta v$$

des rechteckanalogen Flächenstücks S_w ergibt die mittlere Zirkulationsflächendichte von $\vec{F}$ in S_w :

$$\frac{1}{\Delta S_w} \oint_{C_{uv}} \vec{F} \cdot d\vec{r} \approx \frac{1}{g_u g_v} \left\{ \frac{\partial}{\partial u} [g_v F_v(u_0, \tilde{v}, w_0)] + \dots - \frac{\partial}{\partial v} [g_u F_u(\tilde{u}, v_0, w_0)] + \dots \right\} .$$

Wenn wir schließlich den Flächeninhalt ΔS_w gegen Null gehen lassen bzw. die geschlossene Kurve C_{uv} auf den Punkt $\vec{r}_0 = (u_0, v_0, w_0)$ zusammenziehen gemäß

$$\left. \begin{array}{l} \Delta u \\ \Delta v \end{array} \right\} \rightarrow 0 , \quad \begin{array}{l} \tilde{u} \rightarrow u_0 \\ \tilde{v} \rightarrow v_0 \end{array} ,$$

erhalten wir **exakt** die (lokale) Zirkulationsflächendichte in S_w bzw. die skalare w -Komponente der Rotation von $\vec{F}$ im Raumpunkt $(u_0, v_0, w_0) = \vec{r}_0$:

$$\lim_{\Delta S_w \rightarrow 0} \frac{1}{\Delta S_w} \oint_{C_{uv}} \vec{F} \cdot d\vec{r} = \frac{1}{g_u g_v} \left\{ \frac{\partial}{\partial u} [g_v F_v(u_0, v_0, w_0)] - \frac{\partial}{\partial v} [g_u F_u(u_0, v_0, w_0)] \right\} = \text{rot}_w \vec{F}(\vec{r}_0) .$$

Weil die Rotation für das Vektorfeld $\vec{F}$ in der gesamten von C_{uv} umrandeten Fläche gelten soll, ersetzen wir wieder $\vec{r}_0$ durch $\vec{r}$. Durch zyklische Vertauschung von u, v, w erhalten wir schließlich auch die u - und v -Komponente der Rotation und somit das Rotationsvektorfeld

$$\begin{aligned} \text{rot } \vec{F}(\vec{r}(u, v, w)) &= \frac{1}{g_v g_w} \left[\frac{\partial}{\partial v} (g_w \cdot F_w) - \frac{\partial}{\partial w} (g_v \cdot F_v) \right] \vec{e}_u \\ &+ \frac{1}{g_w g_u} \left[\frac{\partial}{\partial w} (g_u \cdot F_u) - \frac{\partial}{\partial u} (g_w \cdot F_w) \right] \vec{e}_v \\ &+ \frac{1}{g_u g_v} \left[\frac{\partial}{\partial u} (g_v \cdot F_v) - \frac{\partial}{\partial v} (g_u \cdot F_u) \right] \vec{e}_w . \end{aligned}$$

Dies lässt sich auch kompakt als Determinante schreiben:

$$\text{rot } \vec{F} = \begin{vmatrix} \vec{e}_u & \vec{e}_v & \vec{e}_w \\ \frac{\partial}{g_v g_w} & \frac{\partial}{g_w g_u} & \frac{\partial}{g_u g_v} \\ \frac{\partial}{\partial u} & \frac{\partial}{\partial v} & \frac{\partial}{\partial w} \\ g_u \cdot F_u & g_v \cdot F_v & g_w \cdot F_w \end{vmatrix} .$$

³Weil wir den Mittelwertsatz wie üblich mit dem „Differenz- Δ “ in ΔS_w anwenden, haben wir von der Benutzung des Flächenelements $dS_w = g_u g_v \cdot du dv$ Abstand genommen.

8 Green'sche Identitäten

Die Green'schen Identitäten werden auch Green'sche Theoreme genannt. Der Vollständigkeit halber geben wir zu Beginn den

$$\textbf{Satz von Stokes} \quad \int_S \text{rot} \vec{F} \cdot d\vec{S} = \oint_{\partial S} \vec{F} \cdot d\vec{r}$$

und den

$$\textbf{Gauß'schen Integralsatz} \quad \int_V \text{div} \vec{F} dV = \oint_{\partial V} \vec{F} \cdot d\vec{S}$$

an.

Die Funktionen Φ und Ψ beschreiben Skalarfelder. Φ sei mindestens ein mal und Ψ mindestens zwei mal stetig differenzierbar. Dann ergibt sich die 1. Green'sche Identität aus dem Gauß'schen Integralsatz wie folgt:

$$\oint_{\partial V} (\Phi \nabla \Psi) \cdot d\vec{S} = \int_V \nabla \cdot (\Phi \nabla \Psi) dV = \int_V (\Phi \nabla^2 \Psi + \nabla \Phi \cdot \nabla \Psi) dV \Rightarrow$$

$$\textbf{1. Green'sche Identität} \quad \oint_{\partial V} (\Phi \nabla \Psi) \cdot d\vec{S} = \int_V (\Phi \Delta \Psi + \nabla \Phi \cdot \nabla \Psi) dV .$$

Jetzt seien sowohl Φ als auch Ψ mindestens zwei mal stetig differenzierbar. Dann ergibt sich die 2. Green'sche Identität aus der 1. Green'schen Identität wie folgt:

$$\oint_{\partial V} (\Phi \nabla \Psi) \cdot d\vec{S} = \int_V (\Phi \Delta \Psi + \nabla \Phi \cdot \nabla \Psi) dV , \quad (8.1)$$

$$\oint_{\partial V} (\Psi \nabla \Phi) \cdot d\vec{S} = \int_V (\Psi \Delta \Phi + \nabla \Psi \cdot \nabla \Phi) dV \Rightarrow \quad (8.2)$$

Wir subtrahieren (8.2) von (8.1) und erhalten mit $\nabla \Phi \cdot \nabla \Psi = \nabla \Psi \cdot \nabla \Phi$ die

$$\textbf{2. Green'sche Identität} \quad \oint_{\partial V} (\Phi \nabla \Psi - \Psi \nabla \Phi) \cdot d\vec{S} = \int_V (\Phi \Delta \Psi - \Psi \Delta \Phi) dV .$$

9 Taylor-Entwicklung eines skalaren Feldes

Siehe auch:

Wolfgang Nolting, *Grundkurs Theoretische Physik 3, Elektrodynamik*, 7. Auflage, Springer, Berlin, Heidelberg, New York, 2004, Abschnitt 1.2 *Taylor-Entwicklung*, Seite 8 bis Seite 13,

Siegfried Großmann, *Mathematischer Einführungskurs für die Physik*, 8. Auflage, Teubner-Verlag, Stuttgart, Leipzig, 2000, Abschnitt 3.3.5. *Taylorentwicklung für Felder*, Seite 113 bis Seite 116.

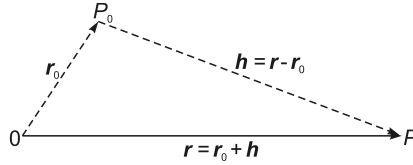


Abb. 9.1 Das skalare Feld Φ wird entwickelt um den Punkt P_0 , d. h. an der Stelle $\vec{r}_0$, für den Aufpunkt P , d. h. für den Ort $\vec{r} = \vec{r}_0 + \vec{h}$.

In Analogie zur Taylor-Entwicklung

$$f(x) = f(x_0 + h) = \sum_{n=0}^{\infty} \frac{1}{n!} f^{(n)}(x) \Big|_{x_0} \cdot h^n, \quad h = x - x_0$$

an der Stelle x_0 , also für $h \rightarrow 0$, entwickeln wir die skalare Feldfunktion

$$\Phi(\vec{r}) = \Phi(\vec{r}_0 + \vec{h}) = \Phi(x_{10} + h_1, x_{20} + h_2, x_{30} + h_3)$$

gemäß Abbildung 9.1 an der Stelle $\vec{r}_0$ mit Hilfe eines Tricks, indem wir

$$\tilde{\Phi}(t) = \tilde{\Phi}(\vec{r}_0 + \vec{h} \cdot t) = \tilde{\Phi}(x_{10} + h_1 t, x_{20} + h_2 t, x_{30} + h_3 t)$$

definieren, sodass

$$\tilde{\Phi}(t=1) \equiv \Phi(\vec{r}_0 + \vec{h}) = \Phi(\vec{r}), \quad \tilde{\Phi}(t=0) = \tilde{\Phi}(0) \equiv \Phi(\vec{r}_0). \quad (9.1)$$

Die Entwicklung von $\tilde{\Phi}(t)$ an der Stelle $t=0$ ist dann

$$\tilde{\Phi}(t) = \sum_{n=0}^{\infty} \frac{1}{n!} \tilde{\Phi}^{(n)} \Big|_{t=0} \cdot t^n.$$

Darin enthalten sind:¹

$$\boxed{\tilde{\Phi}^{(0)}(t) \Big|_{t=0} = \tilde{\Phi}(0) = \Phi(\vec{r}_0)},$$

¹Wegen $\partial x_{i0} = \partial x_i \Big|_{\vec{r}=\vec{r}_0}$ dürfen wir schreiben:

$$\frac{\partial \tilde{\Phi}(\vec{r}_0 + \vec{h} t)}{\partial (x_{i0} + h_i t)} \Big|_{t=0} = \frac{\partial \tilde{\Phi}(\vec{r}_0)}{\partial x_{i0}} = \frac{\partial \Phi(\vec{r})}{\partial x_i} \Big|_{\vec{r}_0}.$$

$$\begin{aligned}
\tilde{\Phi}^{(1)}(t)\Big|_{t=0} &= \left[\frac{\partial \tilde{\Phi}(\vec{r}_0 + \vec{h}t)}{\partial(x_{10} + h_1 t)} \frac{d}{dt}(x_{10} + h_1 t) \right]_{t=0} \\
&+ \left[\frac{\partial \tilde{\Phi}(\vec{r}_0 + \vec{h}t)}{\partial(x_{20} + h_2 t)} \frac{d}{dt}(x_{20} + h_2 t) \right]_{t=0} \\
&+ \left[\frac{\partial \tilde{\Phi}(\vec{r}_0 + \vec{h}t)}{\partial(x_{30} + h_3 t)} \frac{d}{dt}(x_{30} + h_3 t) \right]_{t=0} \\
&= \frac{\partial \Phi(\vec{r})}{\partial x_1} \Big|_{\vec{r}_0} \cdot h_1 + \frac{\partial \Phi(\vec{r})}{\partial x_2} \Big|_{\vec{r}_0} \cdot h_2 + \frac{\partial \Phi(\vec{r})}{\partial x_3} \Big|_{\vec{r}_0} \cdot h_3 ,
\end{aligned}$$

$$\boxed{\tilde{\Phi}^{(1)}(t)\Big|_{t=0} = \sum_{i=1}^3 \frac{\partial \Phi(\vec{r})}{\partial x_i} \Big|_{\vec{r}_0} \cdot h_i} ,$$

$$\begin{aligned}
\tilde{\Phi}^{(2)}(t)\Big|_{t=0} &= \frac{d}{dt} \left[\frac{\partial \tilde{\Phi}(\vec{r}_0 + \vec{h}t)}{\partial(x_{10} + h_1 t)} \frac{d}{dt}(x_{10} + h_1 t) \right]_{t=0} \\
&+ \frac{d}{dt} \left[\frac{\partial \tilde{\Phi}(\vec{r}_0 + \vec{h}t)}{\partial(x_{20} + h_2 t)} \frac{d}{dt}(x_{20} + h_2 t) \right]_{t=0} \\
&+ \frac{d}{dt} \left[\frac{\partial \tilde{\Phi}(\vec{r}_0 + \vec{h}t)}{\partial(x_{30} + h_3 t)} \frac{d}{dt}(x_{30} + h_3 t) \right]_{t=0} \\
&= h_1 \frac{\partial^2 \Phi(\vec{r})}{\partial x_1 \partial x_1} \Big|_{\vec{r}_0} \cdot h_1 + h_1 \frac{\partial^2 \Phi(\vec{r})}{\partial x_1 \partial x_2} \Big|_{\vec{r}_0} \cdot h_2 + h_1 \frac{\partial^2 \Phi(\vec{r})}{\partial x_1 \partial x_3} \Big|_{\vec{r}_0} \cdot h_3 \\
&+ h_2 \frac{\partial^2 \Phi(\vec{r})}{\partial x_2 \partial x_1} \Big|_{\vec{r}_0} \cdot h_1 + h_2 \frac{\partial^2 \Phi(\vec{r})}{\partial x_2 \partial x_2} \Big|_{\vec{r}_0} \cdot h_2 + h_2 \frac{\partial^2 \Phi(\vec{r})}{\partial x_2 \partial x_3} \Big|_{\vec{r}_0} \cdot h_3 \\
&+ h_3 \frac{\partial^2 \Phi(\vec{r})}{\partial x_3 \partial x_1} \Big|_{\vec{r}_0} \cdot h_1 + h_3 \frac{\partial^2 \Phi(\vec{r})}{\partial x_3 \partial x_2} \Big|_{\vec{r}_0} \cdot h_2 + h_3 \frac{\partial^2 \Phi(\vec{r})}{\partial x_3 \partial x_3} \Big|_{\vec{r}_0} \cdot h_3
\end{aligned}$$

$$\begin{aligned}
\tilde{\Phi}^{(2)}(t)\Big|_{t=0} &= \sum_{i=1}^3 \sum_{j=1}^3 \frac{\partial^2 \Phi(\vec{r})}{\partial x_i \partial x_j} \Big|_{\vec{r}_0} \cdot h_i h_j = \left(\sum_{i=1}^3 h_i \frac{\partial}{\partial x_i} \right)^2 \Phi(\vec{r}) \Big|_{\vec{r}_0} , \\
&\vdots
\end{aligned}$$

$$\boxed{\tilde{\Phi}^{(n)}(t)\Big|_{t=0} = \left(\sum_{i=1}^3 h_i \frac{\partial}{\partial x_i} \right)^n \Phi(\vec{r}) \Big|_{\vec{r}_0}} .$$

Folglich ist die Taylor-Entwicklung von $\tilde{\Phi}(t)$ an der Stelle $t = 0$

$$\tilde{\Phi}(t) = \sum_{n=0}^{\infty} \frac{1}{n!} \left(\sum_{i=1}^3 h_i \frac{\partial}{\partial x_i} \right)^n \Phi(\vec{r}) \Big|_{\vec{r}_0} \cdot t^n . \tag{9.2}$$

Schließlich erhalten wir aus der Entwicklung (9.2) und unter Berücksichtigung von (9.1) für $t = 1$ die Taylor-Entwicklung der skalaren Feldfunktion $\Phi(\vec{r}) = \Phi(\vec{r}_0 + \vec{h})$ an der Stelle $\vec{r}_0$ wie folgt:

$$\tilde{\Phi}(t = 1) = \Phi(\vec{r}_0 + \vec{h}) = \Phi(\vec{r}) = \sum_{n=0}^{\infty} \frac{1}{n!} \left(\sum_{i=1}^3 h_i \frac{\partial}{\partial x_i} \right)^n \Phi(\vec{r}) \Big|_{\vec{r}_0}. \quad (9.3)$$

Schreibt man vereinfachend $\vec{r}$ für $\vec{r}_0$, erhält die Taylor-Entwicklung (9.3) schließlich die Form

$$\Phi(\vec{r} + \vec{h}) = \sum_{n=0}^{\infty} \frac{1}{n!} \left(\sum_{i=1}^3 h_i \frac{\partial}{\partial x_i} \right)^n \Phi(\vec{r}).$$

Als Beispiel betrachten wir das Potential φ einer Punktladung q :

$$\varphi = \varphi(\vec{r}) = \frac{\alpha}{|\vec{r} - \vec{r}_0|} = \frac{\alpha}{h} = \varphi(h), \quad \alpha = \frac{q}{4\pi\epsilon_0}.$$

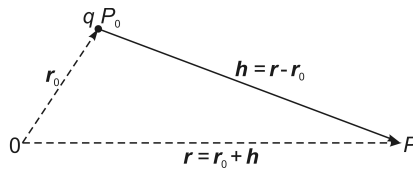


Abb. 9.2 Entwicklung des Coulomb-Potentialfeldes an der Stelle $\vec{r} = \vec{0}$ für den Aufpunkt P .

Konventionsgemäß und wie in Abbildung 9.2 dargestellt, ist q im Punkt $P_0 := \vec{r}_0$ lokalisiert, sodass der Aufpunkt (Messpunkt) $P := \vec{r}$ für das Potential den Abstand $h = |\vec{h}| = |\vec{r} - \vec{r}_0|$ von q besitzt. Tatsächlich also ist φ direkt abhängig vom Betrag des Abstandsvektors $\vec{h} = \vec{r} - \vec{r}_0$ und nicht von $\vec{r}$. Die Entwicklung an der Stelle $P_0 := \vec{r}_0$ mit $\vec{h} = \vec{0} \Rightarrow \vec{r} = \vec{r}_0$ ist wegen

$$\frac{1}{|\vec{r} - \vec{r}_0|} \Rightarrow \frac{1}{|\vec{r}_0 - \vec{r}_0|}$$

nicht definiert. Wir entwickeln deshalb nicht an der Stelle $\vec{h} = \vec{0}$ sondern betrachten $\vec{r}$ statt $\vec{h}$ als die unabhängige Variable und entwickeln an der Stelle $\vec{r} = \vec{0}$ bzw. $x_i = 0$. Die Rolle von $\vec{h}$ in (9.3) wird folglich in diesem Fall übernommen von $\vec{r}$. Weiterhin verwenden wir die geläufige Notation $|\vec{r} - \vec{r}_0| = |\vec{r}_0| = r_0$.

Taylor-Entwicklung von $\varphi(\vec{r}) = \frac{\alpha}{|\vec{r} - \vec{r}_0|}$ bis zur zweiten Ordnung:

$$n = 0 : \\ \Rightarrow \text{Glied 0. Ordnung: } \varphi(\vec{r}) \Big|_{\vec{r}=\vec{0}} = \frac{\alpha}{|\vec{r} - \vec{r}_0|} \Big|_{\vec{r}=\vec{0}} = \frac{\alpha}{|\vec{r}_0|} = \frac{\alpha}{r_0}.$$

$n = 1 :$

$$\begin{aligned} \frac{\partial \varphi(\vec{r})}{\partial x_i} \Big|_{\vec{r}=\vec{0}} &= \frac{\partial}{\partial x_i} \frac{\alpha}{|\vec{r} - \vec{r}_0|} \Big|_{\vec{r}=\vec{0}} \\ &= \alpha \frac{\partial}{\partial x_i} [(x_1 - x_{10})^2 + (x_2 - x_{20})^2 + (x_3 - x_{30})^2]^{-\frac{1}{2}} \Big|_{\vec{r}=\vec{0}} \\ &= \alpha \left(-\frac{1}{2} \right) \frac{2(x_i - x_{i0})}{|\vec{r} - \vec{r}_0|^3} \Big|_{\vec{r}=\vec{0}} = \alpha \frac{x_{i0}}{r_0^3} \end{aligned}$$

$$\Rightarrow \text{Glied 1. Ordnung: } \frac{1}{1!} \sum_{i=1}^3 \frac{\partial \varphi(\vec{r})}{\partial x_i} \Big|_{\vec{r}=\vec{0}} \cdot x_i = \alpha \sum_{i=1}^3 \frac{x_{i0}}{r_0^3} \cdot x_i = \alpha \frac{\vec{r} \cdot \vec{r}_0}{r_0^3} .$$

$n = 2 :$

$$\begin{aligned} \frac{\partial^2 \varphi(\vec{r})}{\partial x_i \partial x_j} \Big|_{\vec{r}=\vec{0}} &= -\alpha \frac{\partial}{\partial x_j} (x_i - x_{i0}) \cdot [(x_1 - x_{10})^2 + (x_2 - x_{20})^2 + (x_3 - x_{30})^2]^{-\frac{3}{2}} \Big|_{\vec{r}=\vec{0}} \\ &= -\alpha \left[\frac{\delta_{ij}}{|\vec{r} - \vec{r}_0|^3} - \frac{3(x_i - x_{i0})(x_j - x_{j0})}{|\vec{r} - \vec{r}_0|^5} \right] \Big|_{\vec{r}=\vec{0}} \\ &= \alpha \left(\frac{3 x_{i0} x_{j0}}{r_0^5} - \frac{\delta_{ij}}{r_0^3} \right) \end{aligned}$$

$\Rightarrow$ Glied 2. Ordnung:

$$\begin{aligned} \frac{1}{2!} \sum_{i=1}^3 \sum_{j=1}^3 \frac{\partial^2 \varphi(\vec{r})}{\partial x_i \partial x_j} \Big|_{\vec{r}=\vec{0}} \cdot x_i x_j &= \alpha \frac{1}{2} \sum_{i=1}^3 \sum_{j=1}^3 \left(\frac{3 x_{i0} x_{j0}}{r_0^5} - \frac{\delta_{ij}}{r_0^3} \right) \cdot x_i x_j \\ &= \alpha \frac{1}{2} \sum_{i=1}^3 \sum_{j=1}^3 \frac{3 x_{i0} x_{j0} \cdot x_i x_j - r_0^2 \delta_{ij} \cdot x_i x_j}{r_0^5} = \alpha \frac{1}{2} \cdot \frac{3(\vec{r} \cdot \vec{r}_0)^2 - r^2 r_0^2}{r_0^5} . \end{aligned}$$

Die Taylor-Entwicklung des Potentialfeldes einer Punktladung ist somit

$$\begin{aligned} \varphi(\vec{r}) &= \frac{\alpha}{|\vec{r} - \vec{r}_0|} = \sum_{n=0}^{\infty} \frac{1}{n!} \left(\sum_{i=1}^3 x_i \frac{\partial}{\partial x_i} \right)^n \varphi(\vec{r} = \vec{0}) \\ &= \alpha \left[\frac{1}{r_0} + \frac{\vec{r} \cdot \vec{r}_0}{r_0^3} + \frac{1}{2} \cdot \frac{3(\vec{r} \cdot \vec{r}_0)^2 - r^2 r_0^2}{r_0^5} + \dots \right] . \end{aligned} \tag{9.4}$$

Wir haben in (9.4) $\varphi(\vec{r} = \vec{0})$ geschrieben, weil die Ableitungen von $\varphi(\vec{r})$ an der Stelle $\vec{r} = \vec{0}$ genommen werden sollen.

10 Dirac'sche delta-Funktion

Siehe auch:

Christian B. Lang und Norbert Pucker, *Mathematische Methoden in der Physik*, Spektrum, Heidelberg, Berlin, 1998, Abschnitt 13.4 Die Diracsche Deltafunktion, Seite 403 bis Seite 408,

Wolfgang Nolting, *Grundkurs Theoretische Physik 3, Elektrodynamik*, 7. Auflage, Springer, Berlin, Heidelberg, New York, 2004, Abschnitt 1.1 Diracsche δ -Funktion, Seite 3 bis Seite 8,

Torsten Fließbach, *Elektrodynamik, Lehrbuch zur Theoretischen Physik II*, 4. Auflage, Elsevier-Spektrum, München, 2005, Kapitel 3 Distributionen, Seite 21 bis Seite 24,

Gero Hillebrandt und Matthias Köhler, *Die Dirac'sche δ -Funktion*, 20. Oktober 2013, www.physik.uni-halle.de/~tpobx/deltafkt.pdf

10.1 Definition der δ -Funktion

Die Dirac'sche delta-Funktion (δ -Funktion) ist eine „uneigentliche“ Funktion, gehört zu den Distributionen und sollte deshalb korrekt als Dirac'sches δ -Funktional bezeichnet werden. Wie im physikalischen Sprachgebrauch üblich, werden wir sie aber dennoch kurz δ -Funktion nennen. Distribution bedeutet Verteilung. Schauen wir uns also zunächst eine kontinuierliche Verteilung an, die Gauß'sche Normalverteilung, auch Gauß'sche Glockenkurve genannt:

$$g_{\sigma}(x) = \frac{1}{\sqrt{2\pi}\sigma} \cdot e^{-\frac{1}{2}\left(\frac{x-x_0}{\sigma}\right)^2} \quad \Rightarrow \quad \int_{-\infty}^{\infty} g_{\sigma}(x) dx = 1, \quad \sigma > 0.$$

Der Parameter σ ist die Standardabweichung und bestimmt Breite und Höhe der Glockenkurve. Je größer σ ist, desto höher liegt das Maximum und desto steiler verläuft die Kurve. Der Parameter x_0 ist die Stelle des Maximums. Die Gauß'sche Normalverteilung nennt man im Fall $x_0 = 0$ und $\sigma = 1$ Standardnormalverteilung. Wie man sieht, ist die *Gauß'sche Normalverteilung* auf 1 normiert. deshalb kann man $g_{\sigma}(x)$ als *Wahrscheinlichkeitsdichte* interpretieren.

Analog zur Gauß'schen Normalverteilung gelte für eine Verteilungsdichtefunktion $d_{\sigma}(x - x_0)$:

- Aus $y = x - x_0 \Rightarrow \frac{dy}{dx} = 1 \Rightarrow dy = dx$ und mit $f(y) \equiv d_{\sigma}(x - x_0)$ resultiert

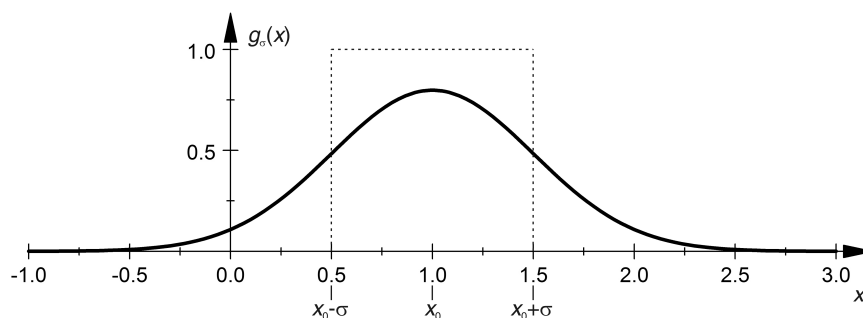


Abb. 10.1 Gauß'sche Normalverteilung $g_{\sigma}(x)$ für $\sigma = 0,5$ und $x_0 = 1$. Der Flächeninhalt des von den punktierten Linien und der x -Achse eingeschlossenen Quadrats ist gleich 1 und somit gleich dem Flächeninhalt zwischen dem Graphen $g_{\sigma}(x)$ und der x -Achse.

$$\int f(y) dy = \int d_\sigma(x - x_0) dx .$$

$d_\sigma(x - x_0)$ ist gegenüber $d_\sigma(x)$ um x_0 längs der x -Achse verschoben, und zwar bei $x_0 > 0$ nach rechts und bei $x_0 < 0$ nach links.

- $d_\sigma(x - x_0)$ sei stets positiv, also $d_\sigma(x - x_0) \stackrel{!}{>} 0$.
(Die Herleitung mit einer stets negativen Verteilungsdichtefunktion ist ebenfalls möglich.)
- $d_\sigma(x - x_0)$ sei eine gerade Funktion und symmetrisch zu x_0 .
- $d_\sigma(x - x_0)$ sei auf 1 normiert, d. h.

$$\int_{-\infty}^{\infty} d_\sigma(x - x_0) dx = 1 . \quad (10.1)$$

Dafür muss $d_\sigma(x - x_0)$ zu beiden Seiten von x_0 hinreichend schnell abfallen, also hinreichend schnell gegen Null gehen.

Eine solche Funktion ist beispielsweise auch die sehr einfache, in einem Bereich der Größe σ um $x = x_0$ lokalisierte, diskrete Verteilungsdichtefunktion (Rechteckfunktion)

$$d_\sigma(x - x_0) = \begin{cases} \frac{1}{\sigma} & \text{für } x_0 - \frac{\sigma}{2} \leq x \leq x_0 + \frac{\sigma}{2} , \\ 0 & \text{sonst .} \end{cases} \quad (10.2)$$

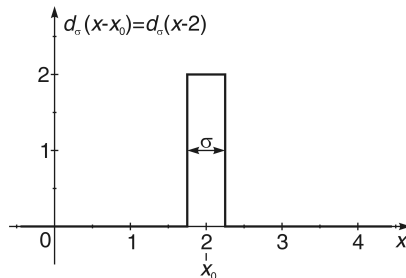


Abb. 10.2 Die Rechteckfunktion $d_\sigma(x - x_0)$ für $\sigma = 0,5$ und $x_0 = 2$. Der vom Rechteck eingeschlossene Flächeninhalt ist gleich 1.

Wie gefordert, ergibt diese Verteilungsdichtefunktion

$$\int_{-\infty}^{\infty} d_\sigma(x) dx = \int_{x_0 - \frac{\sigma}{2}}^{x_0 + \frac{\sigma}{2}} \frac{1}{\sigma} dx = \frac{1}{\sigma} \cdot x \Big|_{x_0 - \frac{\sigma}{2}}^{x_0 + \frac{\sigma}{2}} = \frac{1}{\sigma} \cdot \left(\frac{\sigma}{2} + \frac{\sigma}{2} \right) = 1 .$$

Aus der Verteilungsdichtefunktion $d_\sigma(x - x_0)$ erhalten wir die δ -Funktion durch Grenzwertbildung:

$$\lim_{\sigma \rightarrow 0} d_\sigma(x - x_0) = \delta(x - x_0) .$$

Auch im Limes soll gelten:

$$\boxed{\lim_{\sigma \rightarrow 0} \int_{-\infty}^{\infty} d_\sigma(x - x_0) dx = \int_{-\infty}^{\infty} \delta(x - x_0) dx = 1} . \quad (10.3)$$

Die δ -Funktion ist also u. a. dadurch definiert, dass sie unter dem Integral den Wert 1 liefert. Wie man sieht, ist dabei zuerst das Integral und danach der Grenzwert zu bilden.

Anschaulich wird die Verteilungsdichtefunktion $d_\sigma(x - x_0)$ bei der Grenzwertbildung für $\sigma \rightarrow 0$ immer schmaler und höher bzw. ihre Breite geht gegen Null und ihre Höhe gegen Unendlich, wobei der Flächeninhalt unter dem Graphen von $d_\sigma(x - x_0)$ stets gleich 1 ist, also erhalten bleibt. Das Ergebnis dieser Grenzwertbildung ist die δ -Funktion, eine Funktion mit verschwindender Breite, unendlicher Höhe und dem Flächeninhalt 1 unter ihrem Graphen. Mit anderen Worten, die δ -Funktion liefert nur einen Beitrag an der Stelle $x = x_0$ und ist sonst gleich Null:

$$\delta(x - x_0) \begin{cases} \rightarrow \infty, & x = x_0, \\ = 0, & x \neq x_0. \end{cases}$$

bzw.

$$\int_{x_1}^{x_2} \delta(x - x_0) dx = \begin{cases} 1, & x_0 \in]x_1, x_2[, \\ 0, & x_0 \notin]x_1, x_2[. \end{cases}$$

Für $x_0 = 0$ erhält die δ -Funktion die einfache Form

$$x_0 = 0 \quad \Rightarrow \quad \delta(x - x_0) = \delta(x).$$

10.2 Faltungsintegral mit der δ -Funktion

$$(f * g)(t) := \int_{-\infty}^{+\infty} f(\tau) \cdot g(t - \tau) d\tau \quad (10.4)$$

ist die Definition für die Faltung $(f * g)(t)$ der beiden absolut integrierbaren Funktionen $f(\tau)$ und $g(\tau)$ sowie des zugehörigen Faltungsintegrals. Die δ -Funktion ist gerade und symmetrisch zu x_0 . Deshalb gilt

$$\delta(x - x_0) = \delta(x_0 - x),$$

sodass wir analog zu (10.4) für die Faltung der δ -Funktion mit einer Funktion $f(x)$ folgendes schreiben können:

$$\begin{aligned} (f * \delta)(x_0) &:= \int_{-\infty}^{+\infty} f(x) \cdot \delta(x_0 - x) dx = \int_{-\infty}^{+\infty} f(x) \cdot \delta(x - x_0) dx \\ (f * \delta)(x_0) &:= \int_{-\infty}^{+\infty} \delta(x - x_0) \cdot f(x) dx. \end{aligned}$$

Die Funktion $f(x)$ wird in diesem Zusammenhang auch als Testfunktion bezeichnet.

Bei der Grenzwertbildung für $\sigma \rightarrow 0$ in (10.3) wird der Bereich $\Delta x = x_2 - x_1 = (x_2 - x_0) - (x_1 - x_0)$, der zum Integral beiträgt, immer kleiner und verschwindet schließlich. Insofern sind hierbei die Schreibweisen $\sigma \rightarrow 0$, $\Delta x \rightarrow 0$ und $x_\pm \rightarrow x_0$ gleichbedeutend. Diese Argumentation werden wir im Folgenden verwenden.

Wir zeigen jetzt, wie sich die Grenzwertbildung $\sigma \rightarrow 0$ auf die Faltung unserer als Beispiel verwendeten Verteilungsdichtefunktion (10.2) mit einer beliebigen absolut integrierbaren Testfunktion $f(x)$ auswirkt. Dabei ist zu beachten, dass wir zuerst das Faltungsintegral und danach den Grenzwert bilden:

$$\int_{-\infty}^{\infty} d_{\sigma}(x - x_0) \cdot f(x) \, dx = \int_{x_0 - \frac{\sigma}{2}}^{x_0 + \frac{\sigma}{2}} \frac{1}{\sigma} \cdot f(x) \, dx = \frac{1}{\sigma} \int_{x_0 - \frac{\sigma}{2}}^{x_0 + \frac{\sigma}{2}} f(x) \, dx . \quad (10.5)$$

Mit dem Mittelwertsatz der Integralrechnung

$$\begin{aligned} \int_{x_1}^{x_2} f(x) \, dx &= (x_2 - x_1) \cdot f(\bar{x}) = \\ \Delta F &= \Delta x \cdot f(\bar{x}) \\ \Leftrightarrow f(\bar{x}) &= \frac{\Delta F}{\Delta x} \end{aligned}$$

erhalten wir aus (10.5)

$$\begin{aligned} \int_{-\infty}^{\infty} d_{\sigma}(x - x_0) \cdot f(x) \, dx &= \frac{1}{\sigma} \cdot \left[\left(x_0 + \frac{\sigma}{2} \right) - \left(x_0 - \frac{\sigma}{2} \right) \right] \cdot f(\bar{x}) \\ &= \frac{1}{\sigma} \cdot \sigma \cdot f(\bar{x}) \\ &= \frac{\Delta F}{\Delta x} . \end{aligned}$$

Abschließend bilden wir den Grenzwert für $\sigma \rightarrow 0$ bzw. $\Delta x \rightarrow 0$ und $x_{\pm} \rightarrow x_0$:

$$\lim_{\sigma \rightarrow 0} \int_{-\infty}^{\infty} d_{\sigma}(x - x_0) \cdot f(x) \, dx = \lim_{\substack{\Delta x \rightarrow 0 \\ x_{\pm} \rightarrow x_0}} \frac{\Delta F}{\Delta x} = \left. \frac{dF}{dx} \right|_{x_0} = f(x_0) . \quad (10.6)$$

Für (10.6) kann man in Analogie zu (10.3) mit der δ -Funktion abgekürzt und verallgemeinernd schreiben:

$$\begin{aligned} \lim_{\sigma \rightarrow 0} \int_{-\infty}^{\infty} d_{\sigma}(x - x_0) \cdot f(x) \, dx &= \\ \boxed{\int_{-\infty}^{\infty} \delta(x - x_0) \cdot f(x) \, dx = f(x_0)} &. \end{aligned} \quad (10.7)$$

Dies ist aber gleichbedeutend mit

$$\int_{-\infty}^{\infty} \delta(x - x_0) \cdot f(x_0) \, dx = f(x_0) \int_{-\infty}^{\infty} \delta(x - x_0) \, dx = f(x_0) ,$$

sodass

$$\delta(x - x_0) \cdot f(x) = \delta(x - x_0) \cdot f(x_0) .$$

Wie man sieht, liefert das Faltungsintegral der δ -Funktion $\delta(x - x_0)$ mit einer absolut integrierbaren Funktion $f(x)$ den Funktionswert dieser Funktion an der Stelle x_0 , also $f(x_0)$. Diese Eigenschaft der δ -Funktion kann man im Umgang mit Punktgrößen wie z. B. der Punktladung nutzen.

10.3 Eigenschaften der δ -Funktion – Rechenregeln

- Die δ -Funktion ist symmetrisch zu x_0 und gerade:

$$\delta(x - x_0) = \delta[-(x - x_0)] = \delta(x_0 - x) .$$

- Verschiebung der δ -Funktion:

$$\delta(x - x_0) \tag{10.8}$$

ist gegenüber $\delta(x)$ um $|x_0|$ längs der x -Achse nach rechts verschoben.

$$\delta(x + x_0) \tag{10.9}$$

ist gegenüber $\delta(x)$ um $|x_0|$ längs der x -Achse nach links verschoben. Diese Verschiebung nach links ergibt sich, wenn wir in $\delta(x - x_0)$ ein $x_0 < 0$ einsetzen.

- Die δ -Funktion ist definitionsgemäß auf 1 normiert:

$$\int_{-\infty}^{\infty} \delta(x - x_0) \, dx = 1 .$$

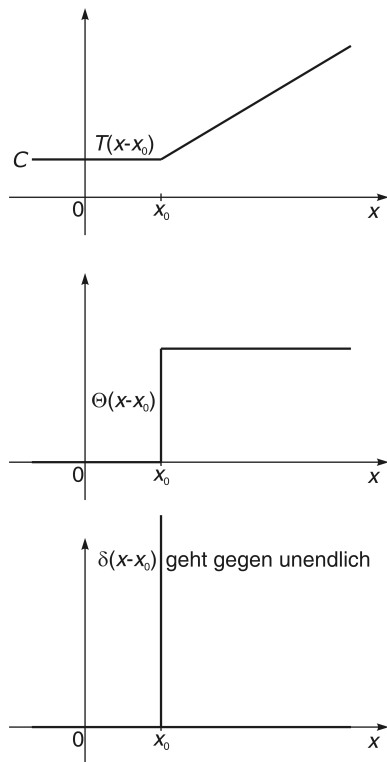


Abb. 10.3 Die Ableitung der oben dargestellten Funktion $T(x - x_0)$ ist die Θ -Funktion (mittig) und die Ableitung der Θ -Funktion ist die δ -Funktion (unten).

- Die δ -Funktion ist die Ableitung der Einheitssprungfunktion¹:
Um dies zu zeigen, gehen wir aus von der stetigen Funktion

$$T(x - x_0) = \begin{cases} C, & x \leq x_0, \\ x - x_0 + C, & x \geq x_0. \end{cases}$$

Wie wir in der Abbildung 10.3 erkennen können, ist die Ableitung der Funktion $T(x - x_0)$ nach x die Θ -Funktion

$$T'(x - x_0) = \Theta(x - x_0) = \int_{-\infty}^{\infty} \delta(x - x_0) \, dx = \begin{cases} 0, & x < x_0, \\ 1, & x > x_0 \end{cases}$$

und die Ableitung der Θ -Funktion nach x ist die δ -Funktion

$$T''(x - x_0) = \Theta'(x - x_0) = \delta(x - x_0).$$

Dass die Funktion $\delta(x - x_0)$ die Ableitung der Einheitssprungfunktion

$$\Theta(x - x_0) = \begin{cases} 0, & x < x_0, \\ 1, & x > x_0 \end{cases}$$

¹Die Einheitssprungfunktion wird auch als Stufenfunktion, Θ -Funktion oder Heaviside-Funktion bezeichnet.

ist, lässt sich auch folgendermaßen zeigen:

$$\begin{aligned}
& \int_{-\infty}^{\infty} \Theta'(x - x_0) f(x) \, dx = \\
& \stackrel{\text{p.I.}}{=} \Theta'(x - x_0) \cdot f(x) \Big|_{-\infty}^{\infty} - \int_{-\infty}^{\infty} \Theta(x - x_0) f'(x) \, dx \\
& = 1 \cdot f(\infty) - 0 \cdot f(-\infty) - \int_{x_0}^{\infty} f'(x) \, dx \\
& = f(\infty) - f(\infty) + f(x_0) \\
& = f(x_0) = \int_{-\infty}^{\infty} \delta(x - x_0) f(x) \, dx \\
& \Rightarrow \Theta'(x - x_0) = \delta(x - x_0) . \quad \square
\end{aligned}$$

Die Θ -Funktion ist ungerade, die δ -Funktion ist gerade und folglich ist die Ableitung der δ -Funktion wieder eine ungerade Funktion gemäß

$$\delta'(x - x_0) = -\delta'(x_0 - x) .$$

- Deshalb verschwindet das Integral über die Ableitung der δ -Funktion:

$$\int_{-\infty}^{\infty} \delta'(x - x_0) \, dx = 0 .$$

Dies lässt sich durch die Faltung von $\delta'(x - x_0)$ mit einer Testfunktion $f(x)$ zeigen. Dabei berücksichtigen wir in der partiellen Integration (p.I.), dass $\delta(x - x_0) = 0$ für $x \rightarrow \pm\infty$.

$$\begin{aligned}
& \int_{-\infty}^{\infty} \delta'(x - x_0) \cdot f(x) \, dx \\
& \stackrel{\text{p.I.}}{=} \underbrace{\delta(x - x_0) \cdot f(x) \Big|_{-\infty}^{\infty}}_{=0} - \int_{-\infty}^{\infty} \delta(x - x_0) \cdot f'(x) \, dx = -f'(x_0) .
\end{aligned}$$

Als Testfunktion verwenden wir jetzt $f(x) \equiv 1 \Rightarrow f'(x) \equiv 0$ und erhalten

$$\int_{-\infty}^{\infty} \delta'(x - x_0) \cdot f(x) \, dx = \int_{-\infty}^{\infty} \delta'(x - x_0) \, dx = -f'(x_0) = 0 . \quad \square$$

- Ist die δ -Funktion die Funktion eines Vektors, z. B. des Ortsvektors $\vec{r} = (x, y, z)$, so erhält sie die Form²

$$\delta(\vec{r} - \vec{r}_0) = \delta(x - x_0) \cdot \delta(y - y_0) \cdot \delta(z - z_0) \quad \Rightarrow$$

²Werden die Integrationsgrenzen nicht angegeben, so ist üblicherweise die Integration von $-\infty$ bis ∞ oder über den ganzen (unendlich ausgedehnten) Raum gemeint.

$$\int \delta(\vec{r} - \vec{r}_0) d^3r = \underbrace{\int \delta(x - x_0) dx}_{=1} \cdot \underbrace{\int \delta(y - y_0) dy}_{=1} \cdot \underbrace{\int \delta(z - z_0) dz}_{=1} = 1 .$$

Allgemein kann man dafür schreiben

$$\delta(\vec{x} - \vec{x}_0) = \prod_i \delta(x_i - x_{i0}) \Rightarrow \prod_i \int \delta(x_i - x_{i0}) dx_i = 1 .$$

- Das auf 1 normierte Integral über die δ -Funktion ist dimensionslos. Deshalb besitzt die δ -Funktion die Maßeinheit des Kehrwerts der Größe, von der sie abhängt. Ist die δ -Funktion z.B. vom Ortsvektor abhängig, so gilt

$$\begin{aligned} [\delta(\vec{r} - \vec{r}_0)] &= [\delta(x - x_0) \cdot \delta(y - y_0) \cdot \delta(z - z_0)] \\ &= \frac{1}{\text{m}} \cdot \frac{1}{\text{m}} \cdot \frac{1}{\text{m}} = \frac{1}{\text{m}^3} = \left[\frac{1}{V} \right] . \end{aligned}$$

- Beschreibung von Punktgrößen mit Hilfe der δ -Funktion am Beispiel der Punktmasse:

$$\varrho(\vec{r}) = m \delta(\vec{r} - \vec{r}_0)$$

sei die Massendichte eines Punktteilchens im Raumpunkt $\vec{r}_0$.³ Das Integral über einen Raumbereich ΔV liefert dann die Punktmasse wie folgt:

$$\int_{\Delta V} \varrho(\vec{r}) d^3r = \int_{\Delta V} m \delta(\vec{r} - \vec{r}_0) d^3r = \begin{cases} m & \text{bei } \vec{r}_0 \text{ wenn } \vec{r}_0 \in \Delta V , \\ 0 & \text{sonst .} \end{cases}$$

- Herleitung von

$$\boxed{\delta(ax) = \frac{1}{|a|} \delta(x)} :$$

Dies gilt nur für $x_0 = 0$. Wir machen eine Fallunterscheidung.

Für $a > 0$ verwenden wir die Substitution

$$ax = |a|x = y \Leftrightarrow x = \frac{y}{|a|} , \quad dx = \frac{1}{|a|} dy ,$$

die Integrationsgrenzen

$$x \rightarrow \pm \infty \Rightarrow y \rightarrow \pm \infty ,$$

die Beziehung

$$f\left(\frac{y}{|a|}\right)\Big|_{y=0} = f(x)\Big|_{x=0} = f(0)$$

³Man beachte, dass $m \cdot \delta(\vec{r} - \vec{r}_0)$ tatsächlich die Maßeinheit (physikalische Dimension) $\text{g} \cdot \frac{1}{\text{m}^3} = \text{g m}^{-3}$ der Massendichte besitzt.

und bilden damit das entsprechende Faltungsintegral:

$$\begin{aligned}
\int_{x \rightarrow -\infty}^{x \rightarrow +\infty} \delta(ax) f(x) \, dx &= \int_{y \rightarrow -\infty}^{y \rightarrow +\infty} \delta(y) f\left(\frac{y}{|a|}\right) \cdot \frac{1}{|a|} \, dy \\
&= \frac{1}{|a|} \int \delta(y) f\left(\frac{y}{|a|}\right) \cdot dy \\
&= \frac{1}{|a|} f(0) = \frac{1}{|a|} \int \delta(x) f(x) \, dx \\
&= \int \frac{1}{|a|} \delta(x) f(x) \, dx \\
\Rightarrow \delta(ax) &= \frac{1}{|a|} \delta(x) .
\end{aligned}$$

Für $a < 0$ verwenden wir die Substitution

$$ax = -|a|x = y \Leftrightarrow x = -\frac{y}{|a|}, \quad dx = -\frac{1}{|a|} dy ,$$

die Integrationsgrenzen

$$x \rightarrow \pm \infty \quad \Rightarrow \quad y \rightarrow \mp \infty ,$$

die Beziehung

$$f\left(-\frac{y}{|a|}\right)\Big|_{y=0} = f(x)\Big|_{x=0} = f(0)$$

und bilden damit das entsprechende Faltungsintegral:

$$\begin{aligned}
\int_{x \rightarrow -\infty}^{x \rightarrow +\infty} \delta(ax) f(x) \, dx &= \int_{y \rightarrow +\infty}^{y \rightarrow -\infty} \delta(y) f\left(-\frac{y}{|a|}\right) \cdot \left(-\frac{1}{|a|}\right) \, dy \\
&= -\frac{1}{|a|} \int_{y \rightarrow +\infty}^{y \rightarrow -\infty} \delta(y) f\left(-\frac{y}{|a|}\right) \cdot dy \\
&= \frac{1}{|a|} \int_{y \rightarrow -\infty}^{y \rightarrow +\infty} \delta(y) f\left(-\frac{y}{|a|}\right) \cdot dy \\
&= \frac{1}{|a|} f(0) = \frac{1}{|a|} \int \delta(x) f(x) \, dx \\
&= \int \frac{1}{|a|} \delta(x) f(x) \, dx \\
\Rightarrow \delta(ax) &= \frac{1}{|a|} \delta(x) . \quad \square
\end{aligned}$$

- Herleitung von

$$\delta(x^2 - x_0^2) = \frac{\delta(x + x_0) + \delta(x - x_0)}{|2x_0|}$$

 : (10.10)

Die Funktion $x^2 - x_0^2 = (x + x_0)(x - x_0) = y(x)$ hat die Nullstellen bei $x = |x_0|$ entsprechend $x_0 > 0$ und bei $x = -|x_0|$ entsprechend $x_0 < 0$. An diesen Nullstellen befinden sich aber auch die Peaks der δ -Funktion (10.15). Jeder dieser beiden Peaks bzw. jede dieser beiden δ -Funktionen liefert im Faltungsintegral mit einer Funktion $f(x)$ seinen eigenständigen Beitrag wie folgt:

$$\begin{aligned} & \int_{-\infty}^{\infty} \delta(x^2 - x_0^2) \cdot f(x) \, dx \\ &= \underbrace{\int_{-\infty}^0 \delta(x^2 - x_0^2) \cdot f(x) \, dx}_{x_0 < 0} + \underbrace{\int_0^{\infty} \delta(x^2 - x_0^2) \cdot f(x) \, dx}_{x_0 > 0} . \end{aligned} \quad (10.11)$$

Mit der Substitution

$$x^2 - x_0^2 = y(x) \Rightarrow x = \begin{cases} -\sqrt{y + x_0^2} & \Leftrightarrow x < 0 \Rightarrow x_0 < 0 , \\ +\sqrt{y + x_0^2} & \Leftrightarrow x > 0 \Rightarrow x_0 > 0 , \end{cases}$$

$$\frac{dy}{dx} = 2x = \begin{cases} -2\sqrt{y + x_0^2} \Rightarrow dx = \frac{1}{-2\sqrt{y + x_0^2}} dy & \text{für } x_0 < 0 , \\ +2\sqrt{y + x_0^2} \Rightarrow dx = \frac{1}{+2\sqrt{y + x_0^2}} dy & \text{für } x_0 > 0 \end{cases}$$

und den Grenzen

$$\begin{aligned} x \rightarrow -\infty & \Rightarrow y \rightarrow \infty , \\ x = 0 & \Rightarrow y = -x_0^2 , \\ x \rightarrow +\infty & \Rightarrow y \rightarrow \infty , \\ -\infty < x < \infty & \Rightarrow -x_0^2 < y < \infty \end{aligned}$$

erhalten wir aus (10.11)

$$\begin{aligned} & \int_{-\infty}^{\infty} \delta(x^2 - x_0^2) \cdot f(x) \, dx \\ &= \int_{-\infty}^{-x_0^2} \delta(y) \cdot f\left(-\sqrt{y + x_0^2}\right) \cdot \frac{1}{-2\sqrt{y + x_0^2}} \, dy \\ & \quad + \int_{-x_0^2}^{\infty} \delta(y) \cdot f\left(\sqrt{y + x_0^2}\right) \cdot \frac{1}{2\sqrt{y + x_0^2}} \, dy \\ &= \int_{-x_0^2}^{\infty} \delta(y) \cdot f\left(-\sqrt{y + x_0^2}\right) \cdot \frac{1}{2\sqrt{y + x_0^2}} \, dy \\ & \quad + \int_{-x_0^2}^{\infty} \delta(y) \cdot f\left(\sqrt{y + x_0^2}\right) \cdot \frac{1}{2\sqrt{y + x_0^2}} \, dy . \end{aligned} \quad (10.12)$$

Weil der Integrationsbereich $-\infty < y \leq -x_0^2$ keinen Beitrag liefert, dürfen wir die untere Integrationsgrenze auf $-\infty$ herabsetzen. Außerdem verwenden wir in (10.12)

$$\frac{1}{2\sqrt{y+x_0^2}} \cdot f\left(-\sqrt{y+x_0^2}\right) = \tilde{\Phi}(y) \quad \text{für } x_0 < 0 ,$$

$$\frac{1}{2\sqrt{y+x_0^2}} \cdot f\left(\sqrt{y+x_0^2}\right) = \Phi(y) \quad \text{für } x_0 > 0 ,$$

sodass

$$\begin{aligned} & \int_{-\infty}^{\infty} \delta(x^2 - x_0^2) \cdot f(x) \, dx \\ &= \int_{-\infty}^{\infty} \delta(y) \cdot \tilde{\Phi}(y) \, dy + \int_{-\infty}^{\infty} \delta(y) \cdot \Phi(y) \, dy \\ &= \tilde{\Phi}(0) + \Phi(0) = \frac{1}{2\sqrt{x_0^2}} \cdot f\left(-\sqrt{x_0^2}\right) + \frac{1}{2\sqrt{x_0^2}} \cdot f\left(\sqrt{x_0^2}\right) \\ &= \frac{1}{|2x_0|} \cdot f(-|x_0|) + \frac{1}{|2x_0|} \cdot f(|x_0|) \\ &= \frac{1}{|2x_0|} \cdot \left\{ \underbrace{f(-|x_0|)}_{x_0 < 0} + \underbrace{f(|x_0|)}_{x_0 > 0} \right\} . \end{aligned} \tag{10.13}$$

Unter Berücksichtigung von (10.7), (10.8) und (10.9) gilt aber

$$\begin{aligned} x_0 < 0 : \quad f(-|x_0|) &= \int_{-\infty}^{\infty} \delta(x+x_0) \cdot f(x) \, dx , \\ x_0 > 0 : \quad f(+|x_0|) &= \int_{-\infty}^{\infty} \delta(x-x_0) \cdot f(x) \, dx . \end{aligned}$$

Wenn wir dies in (10.13) einsetzen, resultiert schließlich

$$\begin{aligned} & \int_{-\infty}^{\infty} \delta(x^2 - x_0^2) \cdot f(x) \, dx \\ &= \frac{1}{|2x_0|} \cdot \left\{ \int_{-\infty}^{\infty} \delta(x+x_0) \cdot f(x) \, dx + \int_{-\infty}^{\infty} \delta(x-x_0) \cdot f(x) \, dx \right\} , \\ & \int_{-\infty}^{\infty} \delta(x^2 - x_0^2) \cdot f(x) \, dx = \int_{-\infty}^{\infty} \frac{\delta(x+x_0) + \delta(x-x_0)}{|2x_0|} \cdot f(x) \, dx . \end{aligned} \tag{10.14}$$

Der Vergleich der beiden Seiten von Gleichung (10.14) zeigt

$$\delta(x^2 - x_0^2) = \frac{\delta(x+x_0) + \delta(x-x_0)}{|2x_0|} \quad \text{mit } |x_0| = |y'(x_0)| . \quad \square$$

- Herleitung der allgemeinen Formel⁴

$$\boxed{\delta(g(x)) = \sum_n \frac{\delta(x - x_n)}{|g'(x_n)|}} : \quad (10.15)$$

x_n seien die Nullstellen n der Funktion $g(x)$, sodass gilt

$$\delta(g(x)) = 0 \quad \text{für } g(x) \neq 0 \text{ bzw. für } x \neq x_n$$

und folglich

$$\sum_n \frac{\delta(x - x_n)}{|g'(x_n)|} = 0 \quad \forall x \neq x_n \text{ bzw. für } g(x) \neq 0 .$$

Mit anderen Worten, $\delta(g(x))$ liefert nur einen Beitrag (Peak) an den Nullstellen x_n der Funktion $g(x)$.

Wir gehen aus vom Faltungsintegral

$$\mathcal{I} \equiv \int_{\alpha}^{\beta} dx \delta(g(x)) \cdot f(x) = \sum_n^{\alpha < x_n < \beta} \int_{x_n - \epsilon}^{x_n + \epsilon} dx \delta(g(x)) \cdot f(x) , \quad (10.16)$$

wobei in den jeweiligen Integrationsbereichen von $x_n - \epsilon$ bis $x_n + \epsilon$ immer nur eine Nullstelle, und zwar die Nullstelle n liegen soll.

Die Taylor-Entwicklung von $g(x)$ an den Nullstellen x_n ist

$$g(x) = \underbrace{g(x_n)}_{=0} + g'(x_n) \cdot (x - x_n) + \frac{1}{2} g''(x_n) \cdot (x - x_n)^2 + \dots .$$

Wenn wir davon ausgehen, dass die Summanden in der Taylor-Entwicklung voneinander unabhängig sind, müssen an den Nullstellen x_n alle Summanden verschwinden, sodass

$$g(x_n) = 0 = \underbrace{g(x_n)}_{=0} + \underbrace{g'(x_n) \cdot (x - x_n)}_{=0 \text{ für } x=x_n} + \underbrace{\frac{1}{2} g''(x_n) \cdot (x - x_n)^2 + 0}_{=0 \text{ für } x=x_n} ,$$

$$g(x_n) = 0 = g'(x_n) \cdot (x - x_n) .$$

Nach Voraussetzung können $\delta(g(x))$ und folglich auch $\mathcal{I}$ nur an den Stellen x_n , also für $g(x_n)$, einen Beitrag liefern. Diesen Sachverhalt berücksichtigen wir jetzt bei der Ersetzung von $g(x)$ in der δ -Funktion $\delta(g(x))$ von (10.16). Es gilt dann dort nämlich exakt

$$g(x) = g'(x_n) \cdot (x - x_n) ,$$

⁴Diese Herleitung ist teilweise zitiert aus: Wolfgang Nolting, *Grundkurs Theoretische Physik 3, Elektrodynamik*, 7. Auflage, Springer, Berlin, Heidelberg, New York, 2004, *Lösung zu Aufgabe 1.7.3*, Seite 371 und Seite 372.

sodass

$$\begin{aligned}
\mathcal{I} &= \sum_n \int_{x_n-\epsilon}^{x_n+\epsilon} dx \, \delta(g'(x_n) \cdot (x - x_n)) \cdot f(x) \\
\mathcal{I} &= \int_{\alpha}^{\beta} dx \sum_n \delta(g'(x_n) \cdot (x - x_n)) \cdot f(x) .
\end{aligned} \tag{10.17}$$

Mit der Substitution

$$g'(x_n) x_n = z_n , \quad g'(x_n) x = z \Leftrightarrow x = \frac{1}{g'(x_n)} z \Rightarrow dx = \frac{1}{g'(x_n)} dz$$

erhalten wir aus (10.17) für

$$g'(x_n) > 0 :$$

$$\begin{aligned}
\mathcal{I} &= \sum_n \int_{g'(x_n)\alpha}^{g'(x_n)\beta} dz \, \frac{1}{g'(x_n)} \delta(z - z_n) \cdot f\left(\frac{z}{g'(x_n)}\right) \\
&= \sum_n \frac{1}{g'(x_n)} \cdot f\left(\frac{z_n}{g'(x_n)}\right) \\
&= \sum_n \frac{1}{g'(x_n)} \cdot f(x_n) \\
\mathcal{I} &= \int_{\alpha}^{\beta} dx \sum_n \frac{1}{g'(x_n)} \delta(x - x_n) \cdot f(x) .
\end{aligned}$$

Der Vergleich mit (10.17) zeigt

$$\begin{aligned}
\delta(g'(x_n) \cdot (x - x_n)) &= \delta(g(x)) = \sum_n \frac{1}{g'(x_n)} \delta(x - x_n) \\
&= \sum_n \frac{1}{|g'(x_n)|} \delta(x - x_n) .
\end{aligned}$$

Weiterhin erhalten wir aus (10.17) für

$$g'(x_n) < 0 :$$

$$\begin{aligned}
\mathcal{I} &= - \sum_n \int_{-|g'(x_n)|^\alpha}^{-|g'(x_n)|^\beta} dz \frac{1}{|g'(x_n)|} \delta(z - z_n) \cdot f\left(\frac{z}{g'(x_n)}\right) \\
&= + \sum_n \int_{-|g'(x_n)|^\beta}^{-|g'(x_n)|^\alpha} dz \frac{1}{|g'(x_n)|} \delta(z - z_n) \cdot f\left(\frac{z}{g'(x_n)}\right) \\
&= \sum_n^{\substack{-|g'|^\alpha > z_n > -|g'|^\beta}} \frac{1}{|g'(x_n)|} \cdot f\left(\frac{z_n}{g'(x_n)}\right) \\
&= \sum_n^{\alpha < x_n < \beta} \frac{1}{|g'(x_n)|} \cdot f(x_n) \\
\mathcal{I} &= \int_\alpha^\beta dx \sum_n \frac{1}{|g'(x_n)|} \delta(x - x_n) \cdot f(x) .
\end{aligned}$$

Der Vergleich mit (10.17) zeigt wieder

$$\delta(g'(x_n) \cdot (x - x_n)) = \delta(g(x)) = \sum_n \frac{1}{|g'(x_n)|} \delta(x - x_n) . \quad \square$$

10.4 Fourier-Transformation und δ -Funktion

Die Fourier-Transformation bezeichnen wir mit $\mathcal{F}$ und ihre Umkehrung, die inverse Fourier-Transformation, mit $\mathcal{F}^{-1}$.

Konvention⁵:

$$\begin{aligned}
\mathcal{F}\{f(x)\} &:= \int_{-\infty}^{\infty} f(x) e^{-i2\pi\xi x} dx = \tilde{f}(\xi) , \\
\mathcal{F}^{-1}\{\tilde{f}(\xi)\} &:= \frac{1}{2\pi} \int_{-\infty}^{\infty} \tilde{f}(\xi) e^{i2\pi\xi x} d\xi = f(x)
\end{aligned}$$

mit dem Normierungsfaktor $\frac{1}{2\pi}$. Die Fourier-Transformierte $\tilde{f}(\xi)$ heißt auch Spektralfunktion der Funktion $f(x)$. $2\pi\xi$ kann z. B. die

$$\text{Ortsfrequenz } k = \frac{2\pi}{\lambda} = 2\pi\xi \quad \Leftrightarrow \quad \xi = \frac{1}{\lambda}$$

oder auch die

$$\text{Zeitfrequenz } \omega = \frac{2\pi}{T} = 2\pi\xi \quad \Leftrightarrow \quad \xi = \frac{1}{T} = \nu$$

sein. An die Stelle der Funktion $f(x)$ setzen wir jetzt die δ -Funktion, verwenden also

$$f(x) \equiv \delta(x - x_0) .$$

⁵Die Zuordnung und eventuelle Aufteilung des Normierungsfaktors $\frac{1}{2\pi}$ sowie die Zuordnung des Vorzeichens im Exponenten der e -Funktion sind Konvention.

Die Fourier-Transformation von $\delta(x - x_0)$ ist dann

$$\begin{aligned}\mathcal{F}\{\delta(x - x_0)\} &= \int_{-\infty}^{\infty} \delta(x - x_0) e^{-i2\pi\xi x} dx = \tilde{f}(\xi) \\ &= e^{-i2\pi\xi x_0} .\end{aligned}$$

Für die Fourier-Transformation von $\delta(x)$ gemäß $x_0 = 0$ gilt also

$$\begin{aligned}\mathcal{F}\{\delta(x)\} &= \int_{-\infty}^{\infty} \delta(x) e^{-i2\pi\xi x} dx = \tilde{f}(\xi) \\ &= e^0 = 1 .\end{aligned}$$

Durch die Umkehroperation, d.h. durch die inverse Fourier-Transformation von $\tilde{f}(\xi) = e^{-i2\pi\xi x_0}$ erhalten wir die Darstellung der δ -Funktion als Fourier-Transformierte von $\tilde{f}(\xi)$ wie folgt:

$$\mathcal{F}^{-1}\{\tilde{f}(\xi)\} = \mathcal{F}^{-1}\{e^{-i2\pi\xi x_0}\} = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{-i2\pi\xi x_0} \cdot e^{i2\pi\xi x} d\xi =$$

$$\delta(x - x_0) = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{i2\pi\xi (x-x_0)} d\xi$$

Die Darstellung der δ -Funktion mit Hilfe der Fourier-Transformation lässt sich auch auf andere Weise zeigen:

$$f(x_0) = \frac{1}{2\pi} \int_{k=-\infty}^{\infty} dk \tilde{f}(k) e^{ikx_0} , \quad \tilde{f}(k) = \int_{x=-\infty}^{\infty} dx f(x) e^{-ikx} .$$

$f(x)$ sei stetig und quadratintegrabel. Wir setzen $\tilde{f}(k)$ in $f(x_0)$ ein:

$$f(x_0) = \frac{1}{2\pi} \int_{k=-\infty}^{\infty} dk \left\{ \int_{x=-\infty}^{\infty} dx f(x) e^{-ikx} \right\} e^{ikx_0} .$$

Die Vertauschung der Integrationsreihenfolge liefert

$$f(x_0) = \int_{x=-\infty}^{\infty} dx f(x) \cdot \underbrace{\frac{1}{2\pi} \int_{k=-\infty}^{\infty} dk e^{ik(x_0-x)}}_{=\delta(x_0-x)=\delta(x-x_0)} \quad (10.18)$$

$$f(x_0) = \int_{x=-\infty}^{\infty} dx f(x) \cdot \delta(x - x_0) . \quad (10.19)$$

Der Vergleich von (10.19) mit (10.18) zeigt

$$\delta(x - x_0) = \frac{1}{2\pi} \int_{k=-\infty}^{\infty} dk e^{ik(x-x_0)} . \quad (10.20)$$

10.5 Die δ -Funktion in Kugelkoordinaten

Wir ersetzen die kartesischen Koordinaten in

$$\int \delta(\vec{r} - \vec{r}') d^3\vec{r} = \int \delta(x - x') \delta(y - y') \delta(z - z') dx dy dz$$

durch Kugelkoordinaten und verwenden dabei

$$d^3\vec{r} = r^2 \sin \vartheta dr d\vartheta d\varphi = r^2 dr d\cos \vartheta d\varphi ,$$

weil $\sin \vartheta d\vartheta = -d\cos \vartheta$ ist, aber hinsichtlich des Vorzeichens bei Berücksichtigung der Integrationsgrenzen unter dem Integral

$$\int_0^\pi \sin \vartheta d\vartheta = \int_{\cos(0)=1}^{\cos(\pi)=-1} -d\cos \vartheta = -\cos \vartheta \Big|_{+1}^{-1} = \int_{-1}^1 d\cos \vartheta = \cos \vartheta \Big|_{-1}^{+1}$$

gilt. Mit

$$\begin{aligned} \int \frac{1}{r^2} \delta(r - r') r^2 dr &= \int \delta(r - r') dr = 1 , \\ \int \delta(\cos \vartheta - \cos \vartheta') d\cos \vartheta &= 1 , \\ \int \delta(\varphi - \varphi') d\varphi &= 1 \end{aligned}$$

erhalten wir dann

$$\int \delta(\vec{r} - \vec{r}') d^3\vec{r} = \int \frac{1}{r^2} \delta(r - r') \delta(\cos \vartheta - \cos \vartheta') \delta(\varphi - \varphi') \underbrace{r^2 dr d\cos \vartheta d\varphi}_{= d^3\vec{r}} = 1 ,$$

sodass

$$\delta(\vec{r} - \vec{r}') = \frac{1}{r^2} \delta(r - r') \delta(\cos \vartheta - \cos \vartheta') \delta(\varphi - \varphi') .$$

10.6 Verallgemeinerung für das Ersetzen der kartesischen Koordinaten in der δ -Funktion (Koordinaten-Transformation)

Wir ersetzen die kartesischen Koordinaten x, y, z durch die Koordinaten u, v, w . Für das Volumenelement gilt dann

$$d\vec{r} = dx dy dz = \left| \frac{\partial(x, y, z)}{\partial(u, v, w)} \right| du dv dw$$

mit der *Funktionaldeterminante*

$$\frac{\partial(x, y, z)}{\partial(u, v, w)} = \det \begin{pmatrix} \frac{\partial x}{\partial u} & \frac{\partial y}{\partial u} & \frac{\partial z}{\partial u} \\ \frac{\partial x}{\partial v} & \frac{\partial y}{\partial v} & \frac{\partial z}{\partial v} \\ \frac{\partial x}{\partial w} & \frac{\partial y}{\partial w} & \frac{\partial z}{\partial w} \end{pmatrix}$$

und die δ -Funktion erhält dann die Form

$$\begin{aligned} \delta(\vec{r} - \vec{r}_0) &= \delta(x - x_0) \cdot \delta(y - y_0) \cdot \delta(z - z_0) \\ &= \gamma(u, v, w) \cdot \delta(u - u_0) \cdot \delta(v - v_0) \cdot \delta(w - w_0), \end{aligned}$$

wobei γ ein noch zu bestimmender Faktor ist. Weil aber

$$\begin{aligned} \int \delta(\vec{r} - \vec{r}_0) \cdot d^3\vec{r} &= \\ &= \int \gamma(u, v, w) \delta(u - u_0) \delta(v - v_0) \delta(w - w_0) \cdot \left| \frac{\partial(x, y, z)}{\partial(u, v, w)} \right| du dv dw \\ &= \int \delta(u - u_0) \delta(v - v_0) \delta(w - w_0) \cdot du dv dw = 1 \end{aligned}$$

ist, muss

$$\gamma(u, v, w) \cdot \left| \frac{\partial(x, y, z)}{\partial(u, v, w)} \right| = 1 \quad \Leftrightarrow \quad \gamma(u, v, w) = \frac{1}{\left| \frac{\partial(x, y, z)}{\partial(u, v, w)} \right|}$$

und folglich

$$\delta(\vec{r} - \vec{r}_0) = \frac{1}{\left| \frac{\partial(x, y, z)}{\partial(u, v, w)} \right|} \cdot \delta(u - u_0) \delta(v - v_0) \delta(w - w_0)$$

gelten. Beispielsweise ist damit die δ -Funktion in Zylinderkoordinaten

$$\delta(\vec{r} - \vec{r}_0) = \frac{1}{\varrho} \delta(\varrho - \varrho_0) \delta(\varphi - \varphi_0) \delta(z - z_0),$$

weil im Volumenelement $\varrho d\varrho d\varphi dz$ in Zylinderkoordinaten der **Betrag der Funktionaldeterminante** gleich ϱ ist.

10.7 Herleitung von $\Delta \frac{1}{|\vec{r} - \vec{r}'|} = -4\pi \delta(\vec{r} - \vec{r}')$

Vergleiche: Wolfgang Nolting, *Grundkurs Theoretische Physik 3, Elektrodynamik*, 7. Auflage, Springer, Berlin, Heidelberg, New York, 2004, Seite 34 und Seite 35.

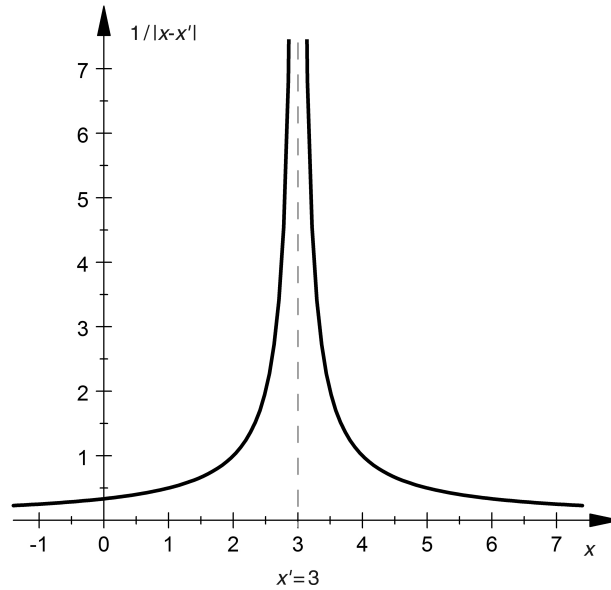


Abb. 10.4 Die Funktion $1/|x - x'| = 1/|x - 3|$, also für $x' = 3$. Zur Vereinfachung von $1/|\vec{r} - \vec{r}'|$ wurden die y - und die z -Achse unterdrückt.

In dieser Herleitung werden wir die folgende Notation verwenden:

$$\begin{aligned}\vec{r} - \vec{r}' &= \vec{\tilde{r}}, \\ x - x' &= \tilde{x}, \\ |\vec{r} - \vec{r}'| &= |\vec{\tilde{r}}| = \tilde{r}.\end{aligned}$$

Die gerade und stets positive Funktion $1/|\vec{r} - \vec{r}'|$ besteht aus zwei Ästen, die symmetrisch liegen zu ihrer Polstelle bei $\vec{r} = \vec{r}'$ (vgl. Abb. 10.4). Die Funktion $1/|\vec{r} - \vec{r}'|$ ist folglich an der Stelle $\vec{r} - \vec{r}' = 0$ nicht definiert und nicht differenzierbar. Die feste Größe $\vec{r}'$ in $1/|\vec{r} - \vec{r}'|$ bewirkt nur eine Verschiebung längs der Achsen der variablen Größe. Wenn wir die y - und die z -Achse unterdrücken und nur die x -Achse betrachten, ergeben sich durch *Fallunterscheidung*, wie in Abbildung 10.4 beispielsweise für $x' = 3$ dargestellt, die beiden Äste

$$\frac{1}{|x - x'|} = \begin{cases} -\frac{1}{x - x'} & \text{für } x < x', \\ +\frac{1}{x - x'} & \text{für } x > x'. \end{cases} \quad (10.21)$$

Weiterhin stellen wir fest, dass x' eine Verschiebung der Funktionsäste längs der x -Achse bewirkt, und zwar

$$\begin{aligned}\text{bei } x' > 0 &\text{ nach rechts,} \\ \text{bei } x' < 0 &\text{ nach links.}\end{aligned}$$

10.7.1 Betrachtungen für $\vec{r} \neq \vec{r}'$

Betrachten wir den Gradienten von $1/|\vec{r} - \vec{r}'|$ für $\vec{r} \neq \vec{r}'$:

$$\begin{aligned}
 \nabla_{\vec{r}} \frac{1}{|\vec{r} - \vec{r}'|} &= \nabla_{\vec{r}} \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{1}{2}} \\
 &= -\frac{1}{2} \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{3}{2}} \cdot \begin{pmatrix} 2(x - x') \\ 2(y - y') \\ 2(z - z') \end{pmatrix}, \\
 &= - \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{3}{2}} \cdot \begin{pmatrix} (x - x') \\ (y - y') \\ (z - z') \end{pmatrix},
 \end{aligned} \tag{10.22}$$

$$\boxed{\nabla_{\vec{r}} \frac{1}{|\vec{r} - \vec{r}'|} = - \frac{\vec{r} - \vec{r}'}{|\vec{r} - \vec{r}'|^3} = \frac{-\vec{e}_{\vec{r}}}{\widetilde{r}^2} \quad \forall \vec{r} \neq \vec{r}'}, \tag{10.23}$$

$$\begin{aligned}
 \nabla_{\vec{r}'} \frac{1}{|\vec{r} - \vec{r}'|} &= -\frac{1}{2} \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{3}{2}} \cdot \begin{pmatrix} -2(x - x') \\ -2(y - y') \\ -2(z - z') \end{pmatrix} \\
 &= + \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{3}{2}} \cdot \begin{pmatrix} (x - x') \\ (y - y') \\ (z - z') \end{pmatrix},
 \end{aligned} \tag{10.24}$$

$$\boxed{\nabla_{\vec{r}'} \frac{1}{|\vec{r} - \vec{r}'|} = + \frac{\vec{r} - \vec{r}'}{|\vec{r} - \vec{r}'|^3} = \frac{+\vec{e}_{\vec{r}}}{\widetilde{r}^2} \quad \forall \vec{r} \neq \vec{r}'}.$$

Jetzt wenden wir den Laplace-Operator $\Delta = \nabla \cdot \nabla = \text{div grad}$ auf die Funktion $1/|\vec{r} - \vec{r}'|$ für $\vec{r} \neq \vec{r}'$ an und gehen dabei von (10.22) aus. Zur Vereinfachung leiten wir zunächst nur die x -Komponente

$$\frac{\partial}{\partial x} \frac{1}{|\vec{r} - \vec{r}'|} = -(x - x') \cdot \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{3}{2}}$$

von (10.22) nach x ab:

$$\begin{aligned}
& \frac{\partial^2}{\partial x^2} \frac{1}{|\vec{r} - \vec{r}'|} = \\
& = - \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{3}{2}} + \\
& \quad \frac{3}{2} \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{5}{2}} \cdot 2(x - x')^2 \quad (10.25) \\
& = - \frac{1}{|\vec{r} - \vec{r}'|^3} + \frac{3(x - x')^2}{|\vec{r} - \vec{r}'|^5} \\
& = \frac{3(x - x')^2 - \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]}{|\vec{r} - \vec{r}'|^5}.
\end{aligned}$$

Analog erhält man die Ableitung der y -Komponente von (10.22) nach y und die Ableitung der z -Komponente von (10.22) nach z . Die Summe dieser drei Ableitungen liefert schließlich

$$\begin{aligned}
& \left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} \right) \frac{1}{|\vec{r} - \vec{r}'|} = \Delta_{\vec{r}} \frac{1}{|\vec{r} - \vec{r}'|} = \\
& = \frac{3(x - x')^2 + 3(y - y')^2 + 3(z - z')^2 - 3 \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]}{|\vec{r} - \vec{r}'|^5} \\
& = 0. \quad (10.26)
\end{aligned}$$

Aus (10.24) bzw. aus

$$\frac{\partial}{\partial x'} \frac{1}{|\vec{r} - \vec{r}'|} = + (x - x') \cdot \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{3}{2}}$$

erhält man auf die gleiche Weise $\Delta_{\vec{r}'} \frac{1}{|\vec{r} - \vec{r}'|} = 0$, denn wie in (10.25) resultiert

$$\begin{aligned}
& \frac{\partial^2}{\partial x'^2} \frac{1}{|\vec{r} - \vec{r}'|} = \\
& = - \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{3}{2}} + \\
& \quad \frac{3}{2} \left[(x - x')^2 + (y - y')^2 + (z - z')^2 \right]^{-\frac{5}{2}} \cdot 2(x - x')^2, \quad (10.27)
\end{aligned}$$

sodass schließlich gilt:

$$\boxed{\Delta_{\vec{r}} \frac{1}{|\vec{r} - \vec{r}'|} = \Delta_{\vec{r}'} \frac{1}{|\vec{r} - \vec{r}'|} = 0 \quad \forall \vec{r} \neq \vec{r}'} \quad (10.28)$$

Unter Berücksichtigung von (10.28) verallgemeinern wir diese Ergebnisse für den $\mathbb{R}^3$ und stellen fest

$$\boxed{\Delta_{\vec{r}} \frac{1}{|\vec{r} - \vec{r}'|} = \Delta_{\vec{r}'} \frac{1}{|\vec{r} - \vec{r}'|} = \Delta \frac{1}{|\vec{r} - \vec{r}'|} = \Delta \frac{1}{\tilde{r}}},$$

sodass wir die Indizierung des Laplace-Operators bei seiner Anwendung auf die Funktion $1/|\vec{r} - \vec{r}'| = 1/\tilde{r}$ vernachlässigen können und nur noch das Symbol Δ verwenden.

Dass $\Delta \frac{1}{|\vec{r} - \vec{r}'|}$ für $\vec{r} \neq \vec{r}'$ verschwindet, liegt an der 3-Dimensionalität des betrachteten Raums. Dies erkennt man daran, wie die Exponenten bei der Differentiation der Potenzfunktionen $[\sum_{i=1}^n (x_i - x'_i)^2]^{-1/2}$ für verschiedene Dimensionalitäten n in die Rechnung eingehen und schließlich durch die Faktoren 3 den Zähler von (10.26) verschwinden lassen. Insofern ist (10.28) eine Eigenschaft des 3-dimensionalen Ortsraums.

Zur Veranschaulichung der Eigenschaften der Funktion $1/|\vec{r} - \vec{r}'|$ reduzieren wir diese auf den eindimensionalen Fall mit den beiden Funktionen

$$f(x) = \frac{1}{|x - x'|} = f(x, x_0) ,$$

$$g(x') = \frac{1}{|x - x'|} = g(x_0, x') .$$

mit

$$f(x, x_0) \equiv g(x_0, x') \quad \text{für alle } x_0 = x' .$$

Bilden wir den Gradienten dieser beiden Funktionen:

$$\nabla_x \frac{1}{|x - x'|} \Rightarrow \frac{d}{dx} \frac{1}{|x - x'|} = -\frac{x - x'}{|x - x'|^3} = -\frac{\tilde{x}}{|\tilde{x}|^3} = -\frac{1}{|\tilde{x}|^2} \cdot \frac{\tilde{x}}{|\tilde{x}|} , \quad (10.29)$$

$$\nabla_{x'} \frac{1}{|x - x'|} \Rightarrow \frac{d}{dx'} \frac{1}{|x - x'|} = +\frac{x - x'}{|x - x'|^3} = +\frac{\tilde{x}}{|\tilde{x}|^3} = +\frac{1}{|\tilde{x}|^2} \cdot \frac{\tilde{x}}{|\tilde{x}|} \quad (10.30)$$

$$= -\frac{x' - x}{|x - x'|^3} = -\frac{-\tilde{x}}{|\tilde{x}|^3} = -\frac{1}{|\tilde{x}|^2} \cdot \frac{-\tilde{x}}{|\tilde{x}|} . \quad (10.31)$$

Wegen $x - x' = \tilde{x} \Leftrightarrow x' - x = -\tilde{x}$ schlussfolgern wir daraus

$$\frac{\tilde{x}}{|\tilde{x}|} \hat{=} \operatorname{sgn} \left(\frac{x - x'}{|x - x'|} \right) \cdot \vec{e}_x \quad \text{mit} \quad \operatorname{sgn} \left(\frac{x - x'}{|x - x'|} \right) = 1 > 0 \text{ für } x > x' ,$$

$$\frac{-\tilde{x}}{|\tilde{x}|} \hat{=} \operatorname{sgn} \left(\frac{x' - x}{|x - x'|} \right) \cdot \vec{e}_{x'} \quad \text{mit} \quad \operatorname{sgn} \left(\frac{x' - x}{|x - x'|} \right) = 1 > 0 \text{ für } x' > x .$$

Wie wir sehen ist der Gradient nach x' entgegengesetzt gerichtet zum Gradienten nach x entsprechend $\vec{e}_{x'} = -\vec{e}_x$. Anders gesagt, die x' -Achse ist entgegengesetzt gerichtet zur x -Achse, zeigt also in die negative x -Richtung. Diese Tatsache spielt aber keine Rolle für unsere Betrachtung, weil der Laplace-Operator $\Delta = \nabla^2$ unabhängig vom richtungsbedingtem Vorzeichen des Gradienten ist. Dafür liefern aber gemäß (10.29) und (10.31) die ersten Ableitungen

$$\frac{df(x, x_0)}{dx} = \frac{d}{dx} \frac{1}{|x - x'_0|} = -\frac{x - x'_0}{|x - x'_0|^3} ,$$

$$\frac{dg(x_0, x')}{dx'} = \frac{d}{dx'} \frac{1}{|x_0 - x'|} = -\frac{x' - x_0}{|x_0 - x'|^3}$$

identische Funktionsgraphen im Fall $x_0 = x'_0$, wie wir auch an den Abbildungen 10.5 und 10.6 sofort erkennen.

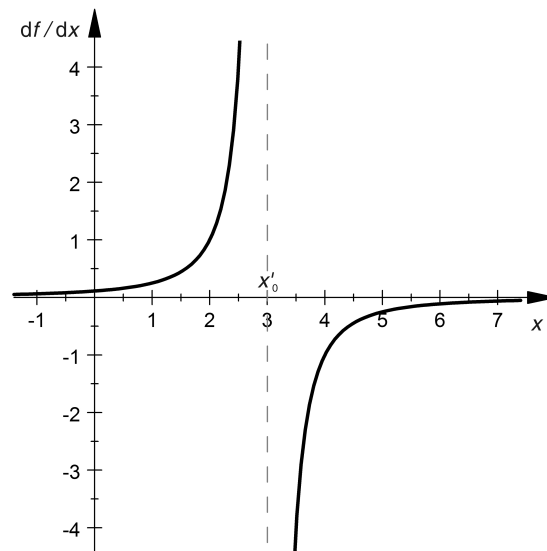


Abb. 10.5 Graphische Darstellung der ersten Ableitung $\frac{df(x, x_0')}{dx} = \frac{d}{dx} \frac{1}{|x - x_0'|} = -\frac{x - x_0'}{|x - x_0'|^3}$ für $x_0' = 3$ entsprechend (10.29). Die Steigung beider Funktionsäste ist stets positiv.

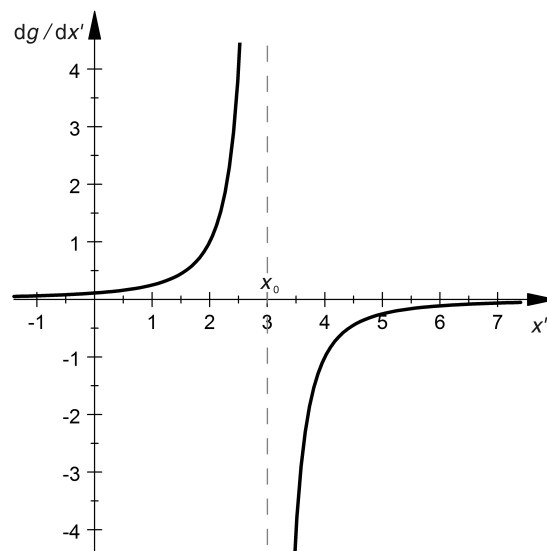


Abb. 10.6 Graphische Darstellung der ersten Ableitung $\frac{dg(x_0, x')}{dx'} = \frac{d}{dx'} \frac{1}{|x_0 - x'|} = -\frac{x' - x_0}{|x_0 - x'|^3}$ für $x_0 = 3$ entsprechend (10.31). Die x' -Achse zeigt konventionsgemäß in die positive Richtung und die Steigung beider Funktionsäste ist stets positiv.

Schließlich bilden wir die zweite Ableitung von $f(x) = 1/|x - x'|$ bzw. $g(x') = 1/|x - x'|$, d. h. wir leiten (10.29) nach x und (10.31) nach x' ab:

$$\left. \begin{aligned} \frac{\partial^2 f}{\partial x^2} &= \frac{d^2}{dx^2} \frac{1}{|x - x'|} \\ \frac{\partial^2 g}{\partial x'^2} &= \frac{d^2}{dx'^2} \frac{1}{|x - x'|} \end{aligned} \right\} = \frac{2}{|x - x'|^3} = \frac{2}{|\tilde{x}|^3} > 0. \quad (10.32)$$

Während die ersten Ableitungen von $1/|x - x'|$ nach x und auch nach x' ungerade Funktionen sind, ist die zweite Ableitung $2/|x - x'|^3$ eine stets positive gerade Funktion mit einem Pol bei $x = x'$ (Abb. 10.7).

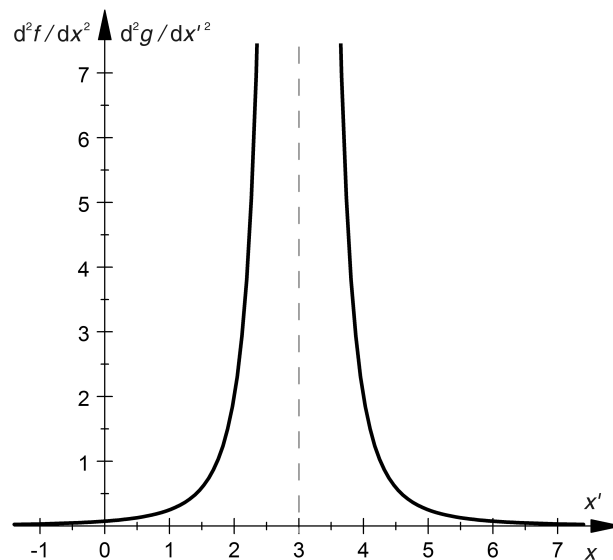


Abb. 10.7 Graphische Darstellung der zweiten Ableitung der Funktion $f(x, x'_0)$ für $x'_0 = 3$ bzw. der Funktion $g(x_0, x')$ für $x_0 = 3$ entsprechend (10.32). Wie man sieht, sind die Graphen der beiden Funktionen im Fall $x_0 = x'_0$ identisch.

Dass die zweite Ableitung von $1/|x - x'|$ nach x gleich der zweiten Ableitung nach x' ist, konnten wir bereits aus den Abbildungen 10.5 und 10.6 ersehen, denn identische Funktionsgraphen besitzen identische Steigungen (Ableitungen).

Achtung!

Die in der Abbildung 10.7 dargestellte Funktion (10.32) ist nicht das Ergebnis der Anwendung des Laplace-Operators auf die Funktion $1/|\vec{r} - \vec{r}'|$ mit anschließender Nullsetzung von y und z . Anders gesagt, (10.32) ist nicht das nur längs der x -Achse betrachtete Ergebnis der Anwendung des (dreidimensionalen) Laplace-Operators auf die Funktion $1/|\vec{r} - \vec{r}'|$. Diese Darstellung soll nur als Analogie dienen, um das Verhalten des Laplace-Operators hinsichtlich des Vorzeichens zu veranschaulichen.

Die Anwendung des Laplace-Operators auf die Funktion $1/|\vec{r} - \vec{r}'|$ ergibt längs der x -Achse, also für den Fall

$$\vec{r} - \vec{r}' = \begin{pmatrix} x - x' \\ 0 \\ 0 \end{pmatrix} \Rightarrow |\vec{r} - \vec{r}'| = |x - x'|$$

gemäß (10.26) und (10.27) :

$$\left. \begin{aligned} & \left[\Delta_{\vec{r}} \frac{1}{|\vec{r} - \vec{r}'|} \right]_{y'=z'=0} \\ & \left[\Delta_{\vec{r}'} \frac{1}{|\vec{r} - \vec{r}'|} \right]_{y=z=0} \end{aligned} \right\} = \frac{3(x - x')^2 - 3(x - x')^2}{|x - x'|^5} = 0 \quad \text{für} \quad |x - x'| > 0 .$$

Als **skalares Feld** im $\mathbb{R}^3$ ist die Funktion $1/|\vec{r} - \vec{r}'|$ nicht eine nur von x bzw. x' abhängige Funktion sondern eine Funktion von $\vec{r}$ bzw. $\vec{r}'$. Deshalb müssen wir zuerst den Laplace-Operator $\Delta_{\vec{r}}$ bzw. $\Delta_{\vec{r}'}$ auf die Funktion $1/|\vec{r} - \vec{r}'|$ **im dreidimensionalen Raum** anwenden und erst danach dürfen wir die Koordinatenwerte (hier z. B. $y = y' = z = z' = 0$) einsetzen.

Der Vollständigkeit halber und wegen der praktischen Bedeutung zeigen wir noch

$$\Delta \frac{1}{|\vec{r} - \vec{r}'|} = \Delta_r \frac{1}{r} \quad \text{für} \quad |\vec{r} - \vec{r}'| = r > 0$$

in **Kugelkoordinaten**. Je nach Bedarf sind für die Radialableitung des Laplace-Operators in Kugelkoordinaten bezüglich einer Funktion $f = f(r)$ drei Schreibweisen gebräuchlich:

$$\Delta_r f(r) = \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \cdot \frac{\partial f}{\partial r} \right) = \frac{2}{r} \frac{\partial f}{\partial r} + \frac{\partial^2 f}{\partial r^2} = \frac{1}{r} \frac{\partial^2}{\partial r^2} (r \cdot f) .$$

Setzen wir also $f(r) = 1/r$ ein:

$$\Delta_r \frac{1}{r} = \frac{1}{r} \frac{\partial^2}{\partial r^2} \left(r \cdot \frac{1}{r} \right) = \frac{1}{r} \frac{\partial^2}{\partial r^2} (1) = 0 . \quad \square$$

In Kugelkoordinaten benötigen wir hier nur die Radialableitung des Laplace-Operators, weil die beiden anderen Ableitungen sofort verschwinden und weil sich die Radialableitung wegen der Kugelsymmetrie über den ganzen $\mathbb{R}^3$ erstreckt.

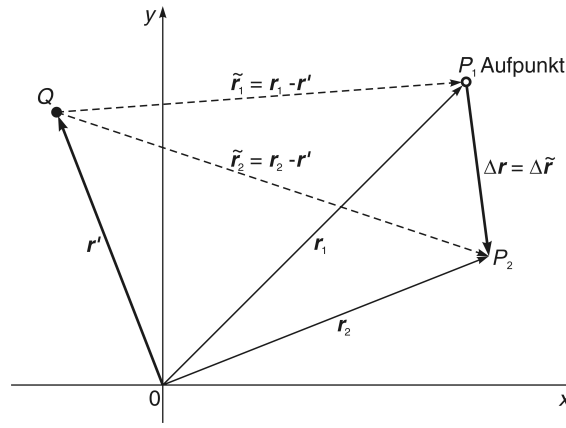


Abb. 10.8 Veranschaulichung des Vektors $\Delta \vec{r} = \vec{r}_2 - \vec{r}_1 = (\vec{r}_2 - \vec{r}') - (\vec{r}_1 - \vec{r}') = \vec{r}_2 - \vec{r}_1 = \Delta \vec{r}$ (zur Vereinfachung in zwei Dimensionen).

10.7.2 Betrachtungen unter Einschluss von $\vec{r} = \vec{r}'$

Wir bilden jetzt das Volumenintegral über die Funktion $\Delta \frac{1}{|\vec{r} - \vec{r}'|}$ mit den folgenden Substitutionen:

- $\vec{r}'$ bzw. x' sei eine gegebene Konstante und $\vec{r}$ bzw. x die Variable. Dann gilt für die Komponente x (stellvertretend für die Komponenten y und z):

$$\tilde{x} = x - x' \Rightarrow \frac{d\tilde{x}}{dx'} = -1, \quad \frac{d\tilde{x}}{dx} = 1 \Rightarrow dx = d\tilde{x}$$

$$\text{bzw. } \Delta x = \Delta \tilde{x} \text{ (s. Abb. 10.8)} \Rightarrow \frac{d}{dx} = \frac{d}{d\tilde{x}} \Rightarrow \frac{\partial^2}{\partial x^2} = \frac{\partial^2}{\partial \tilde{x}^2}.$$

- Wegen der Analogie zwischen den drei Raumkoordinaten gilt für den $\mathbb{R}^3$

$$\Delta \vec{r} = \Delta \tilde{\vec{r}} \quad \text{bzw.} \quad d\vec{r} = d\tilde{\vec{r}}, \quad dx dy dz = d^3r = d^3\tilde{r} = d\tilde{x} d\tilde{y} d\tilde{z} \quad \text{und}$$

$$\Delta = \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} = \frac{\partial^2}{\partial \tilde{x}^2} + \frac{\partial^2}{\partial \tilde{y}^2} + \frac{\partial^2}{\partial \tilde{z}^2}.$$

- Für den Integrationsbereich V des Volumenintegrals schreiben wir jetzt folglich $\tilde{V}$.

Wir erhalten somit

$$\int_V \Delta \frac{1}{|\vec{r} - \vec{r}'|} d^3r = \int_{\tilde{V}} \Delta \frac{1}{\tilde{r}} d^3\tilde{r}.$$

Für $\vec{r} \neq \vec{r}' \Rightarrow \tilde{r} \neq 0$ bzw. wenn $\vec{r}'$ nicht im Integrationsbereich $\tilde{V}$ liegt, gilt

$$\int_{\tilde{V}} \Delta \frac{1}{\tilde{r}} d^3\tilde{r} = 0, \quad \vec{r}' \notin \tilde{V}.$$

Alle Volumina $\tilde{V}$, die den Nullpunkt $\tilde{r} = 0$ nicht enthalten, ergeben also das Integral Null. Das aber bedeutet:

Wenn das Integral über ein Volumen $\tilde{V}$ existiert, welches den Nullpunkt $\tilde{r} = 0$ beinhaltet, hängt der Wert des Integrals nicht von der Form des Volumens ab, sondern sein Wert resultiert nur aus dem Punkt Q gemäß $\vec{r} = \vec{r}'$ bzw. $\tilde{r} = \vec{0}$. $\tilde{V}$ darf also auch die Form einer Kugel mit dem Mittelpunkt $\tilde{r} = \vec{0}$ entsprechend $\tilde{r} = 0$ und dem Radius $\tilde{R}$ haben. Der Punkt Q in der Abbildung 10.8 entspräche dann dem Mittelpunkt dieser Kugel mit dem Volumen $\tilde{V} = V_K$.

Diese Tatsache ermöglicht uns, das Integral in **Kugelkoordinaten** zu berechnen, wobei wir auch über $\Delta(1/\tilde{r})$ mit der Polstelle bei $\tilde{r} = 0$ integrieren. Dafür benötigen wir den Gradienten in Kugelkoordinaten:

$$\text{grad } \Phi(r, \vartheta, \varphi) = \frac{\partial \Phi}{\partial r} \vec{e}_r + \frac{1}{r} \frac{\partial \Phi}{\partial \vartheta} \vec{e}_\vartheta + \frac{1}{r \sin \vartheta} \frac{\partial \Phi}{\partial \varphi} \vec{e}_\varphi ,$$

das Volumenelement in Kugelkoordinaten:

$$dV = d^3r = dA \cdot dr = r^2 \sin \vartheta \, dr \, d\vartheta \, d\varphi ,$$

das Flächenelement auf der Oberfläche S einer Kugel mit dem Radius R , dem Kugelvolumen V_K und dem Mittelpunkt im Koordinatenursprung:

$$dA = R^2 \sin \vartheta \, d\vartheta \, d\varphi \quad \Rightarrow \quad d\vec{A} = R^2 \sin \vartheta \, d\vartheta \, d\varphi \cdot \vec{e}_r$$

sowie den Gauß'schen Satz

$$\int_V \text{div } \vec{v} \, d^3r = \oint_{S(V)} \vec{v} \, d\vec{A} .$$

Unter den genannten Bedingungen für das Volumen $\tilde{V}$ bzw. das Kugelvolumen V_K und mit (10.23) gilt

$$\nabla_{\vec{r}} \frac{1}{|\vec{r} - \vec{r}'|} = \text{grad}_{\tilde{r}} \frac{1}{\tilde{r}} = -\frac{1}{\tilde{r}^2} \vec{e}_{\tilde{r}} ,$$

sodass

$$\begin{aligned} \int_{\tilde{V}} \Delta \frac{1}{\tilde{r}} \, d^3\tilde{r} &= \int_{V_K} \text{div} \left(\text{grad} \frac{1}{\tilde{r}} \right) \, d^3\tilde{r} = \underbrace{\oint_{S(V_K)} \left(-\frac{1}{\tilde{R}^2} \vec{e}_{\tilde{r}} \right) \cdot d\vec{A}}_{\text{Hüllintegral über die Oberfläche } S \text{ der Kugel mit Radius } \tilde{R}} \\ &= \int_{\varphi=0}^{2\pi} \int_{\vartheta=0}^{\pi} \left(-\frac{1}{\tilde{R}^2} \cdot \vec{e}_{\tilde{r}} \right) \cdot \tilde{R}^2 \sin \vartheta \, d\vartheta \, d\varphi \cdot \vec{e}_{\tilde{r}} \\ &= - \int_{\varphi=0}^{2\pi} d\varphi \int_{\vartheta=0}^{\pi} \sin \vartheta \, d\vartheta = -2\pi \cdot \left[-\cos \vartheta \right]_{\vartheta=0}^{\pi} \\ &= -2\pi \cdot \left[1 - (-1) \right] \\ \int_{\tilde{V}} \Delta \frac{1}{\tilde{r}} \, d^3\tilde{r} &= -4\pi = \int_V \Delta \frac{1}{|\vec{r} - \vec{r}'|} \, d^3r \quad \text{für } \vec{r}' \in V . \end{aligned}$$

Zusammenfassend stellen wir fest:

$$\frac{1}{-4\pi} \int_V \Delta \frac{1}{|\vec{r} - \vec{r}'|} d^3r = \begin{cases} 0 & \text{falls } \vec{r}' \notin V \\ 1 & \text{falls } \vec{r}' \in V \end{cases} = \int_V \delta(\vec{r} - \vec{r}') d^3r .$$

Der Vergleich der Integranden liefert

$$\Delta \frac{1}{|\vec{r} - \vec{r}'|} = -4\pi \delta(\vec{r} - \vec{r}') .$$

Analog zu

$$\Delta \frac{1}{|\vec{r} - \vec{r}'|} = \Delta \frac{1}{\tilde{r}} = -4\pi \delta(\vec{r} - \vec{r}') = -4\pi \delta(\tilde{r})$$

findet man für $\vec{r}' = \vec{0}$

$$\Delta \frac{1}{r} = -4\pi \delta(\vec{r}) .$$

10.8 Herleitung von $(\Delta + k^2) \frac{\exp(\pm ikr)}{r} = -4\pi \delta(\vec{r} - \vec{r}')$

Siehe: Torsten Fließbach, *Elektrodynamik, Lehrbuch zur Theoretischen Physik II*, 4.Auflage, Elsevier, München, 2009, Seite 26.

In **Kugelkoordinaten** ist

$$\Delta_r \Phi(r) = \frac{1}{r} \cdot \frac{\partial^2}{\partial r^2} [r \cdot \Phi(r)] .$$

Damit erhalten wir im Fall $r \neq 0$

$$\begin{aligned} (\Delta + k^2) \frac{e^{\pm ikr}}{r} &= \Delta_r \frac{e^{\pm ikr}}{r} + k^2 \frac{e^{\pm ikr}}{r} \\ &= \frac{1}{r} \cdot (-k^2) e^{\pm ikr} + \frac{1}{r} \cdot k^2 e^{\pm ikr} = 0 . \quad \square \end{aligned} \quad (10.33)$$

Jetzt schließen wir den Fall $r = 0$ in unsere Überlegungen ein. Mit der Taylor-Entwicklung

$$e^{\pm ikr} = 1 \pm ikr - \frac{1}{2} k^2 r^2 \pm \frac{1}{6} ik^3 r^3 - \frac{1}{24} k^4 r^4 \pm \dots$$

an der Stelle $r = 0$ können wir

$$\begin{aligned} (\Delta + k^2) \frac{e^{\pm ikr}}{r} &= (\Delta + k^2) \left(\frac{1}{r} \pm ik - \frac{1}{2} k^2 r \pm \frac{1}{6} ik^3 r^2 - \frac{1}{24} k^4 r^3 \pm \dots \right) \\ &= \Delta \frac{1}{r} \pm \Delta(ik) - \Delta \left(\frac{1}{2} k^2 r \right) \pm \Delta \left(\frac{1}{6} ik^3 r^2 \right) - \Delta \left(\frac{1}{24} k^4 r^3 \right) \pm \dots \\ &\quad + \frac{k^2}{r} \pm ik^3 - \frac{1}{2} k^4 r \pm \frac{1}{6} ik^5 r^2 - \frac{1}{24} k^6 r^3 \pm \dots \end{aligned}$$

schreiben. Nach Anwendung des Laplace-Operators gemäß

$$\Delta \frac{1}{r} = -4\pi \delta(\vec{r}) , \quad \Delta r = \frac{2}{r} , \quad \Delta r^2 = 6 , \quad \Delta r^3 = 12r$$

resultiert dann

$$\begin{aligned} (\Delta + k^2) \frac{e^{\pm ikr}}{r} &= -4\pi \delta(\vec{r}) \pm 0 - \frac{k^2}{r} \pm ik^3 - \frac{1}{2} k^4 r \pm \dots \\ &\quad + \frac{k^2}{r} \pm ik^3 - \frac{1}{2} k^4 r \pm \frac{1}{6} ik^5 r^2 - \frac{1}{24} k^6 r^3 \pm \dots , \\ (\Delta + k^2) \frac{e^{\pm ikr}}{r} &= -4\pi \delta(\vec{r}) \pm 2ik^3 - k^4 r \pm \dots . \end{aligned} \quad (10.34)$$

Das Integral von (10.34) über den ganzen Raum mit dem Volumenelement $d^3r = r^2 \sin \vartheta \, dr \, d\vartheta \, d\varphi$ in Kugelkoordinaten kann nur an der Stelle $r = 0$ einen Betrag liefern, da der Integrand für $r \neq 0$ gemäß (10.33) verschwindet. Wir integrieren deshalb über den ganzen Raum einschließlich $r = 0$, also von $r = 0$ bis $r = \varepsilon > 0$, nachdem wir (10.34) mit einer beliebigen Testfunktion⁶ $f(\vec{r})$ multipliziert haben:

$$\begin{aligned} \int_{r \leq \varepsilon} d^3r \, f(\vec{r}) \cdot (\Delta + k^2) \frac{e^{\pm ikr}}{r} &= \int_{r \leq \varepsilon} d^3r \, f(\vec{r}) \cdot \left[-4\pi \delta(\vec{r}) \right] \\ &\quad + \int_{r \leq \varepsilon} d^3r \, f(\vec{r}) \cdot \left[\pm 2ik^3 - k^4 r \pm \dots \right] \\ \int_{r \leq \varepsilon} d^3r \, f(\vec{r}) \cdot (\Delta + k^2) \frac{e^{\pm ikr}}{r} &= -4\pi f(0) + \mathcal{O}(\varepsilon^2) + \mathcal{O}(\varepsilon^3) + \dots . \end{aligned} \quad (10.35)$$

An der Stelle $r = 0$, also für $\varepsilon \rightarrow 0$, verschwinden aber in (10.35) die Terme $\mathcal{O}(\varepsilon^2)$ und höherer Ordnung, sodass an der Stelle $r = 0$ nur der Term $-4\pi f(0)$ verbleibt. Es gilt also tatsächlich

$$(\Delta + k^2) \frac{e^{\pm ikr}}{r} = -4\pi \delta(\vec{r}) .$$

Für r schreiben wir $|\vec{r} - \vec{r}'|$, sodass schließlich

$$\boxed{(\Delta + k^2) \frac{e^{\pm ik|\vec{r} - \vec{r}'|}}{|\vec{r} - \vec{r}'|} = -4\pi \delta(\vec{r} - \vec{r}')} . \quad (10.36)$$

⁶Testfunktionen sind eine spezielle Klasse von Funktionen, d. h., Testfunktionen haben besondere Eigenschaften. Insbesondere darf die Testfunktion $f(\vec{r})$ in unserem Fall bei $r = 0$ keinen Pol besitzen. Die Funktion $f = 1/r^4$ z. B. wäre keine Testfunktion.

10.9 Die δ -Funktion in der Elektrostatik

Vergleiche: Torsten Fließbach, *Elektrodynamik, Lehrbuch zur Theoretischen Physik II*, 4. Auflage, Elsevier – Spektrum, München, 2005, Seite 43 bis Seite 47.

Zur Vereinfachung sei im Folgenden $\int := \int_{-\infty}^{\infty}$.

- Ladungsdichte für eine Punktladung q_i am Ort $\vec{r}_i$:

$$\varrho(\vec{r}) = q_i \delta(\vec{r} - \vec{r}_i) .$$

- Ladungsdichte für N Punktladungen q_i an den Orten $\vec{r}_i$:

$$\varrho(\vec{r}) = \sum_{i=1}^N q_i \delta(\vec{r} - \vec{r}_i) .$$

- Punktladung q_i am Ort $\vec{r}_i$:

$$q_i(\vec{r}_i) = \int \varrho(\vec{r}) d^3r = \int q_i \delta(\vec{r} - \vec{r}_i) d^3r .$$

- Diskrete Ladungsverteilung $Q(r_i)$ der Gesamtladung Q , bestehend aus N Punktladungen q_i an den Orten $\vec{r}_i$ in einem Raumbereich mit dem Volumen V :

$$Q(\vec{r}_i) = \sum_{i=1}^N q_i(\vec{r}_i) = \sum_{i=1}^N \int \underbrace{q_i \delta(\vec{r} - \vec{r}_i)}_{\varrho(\vec{r})} d^3r = Q .$$

- Näherung der Ladungsdichte einer *stetigen* Ladungsverteilung durch eine diskrete Verteilung von N Punktladungen in einem begrenzten Raumbereich mit dem Volumen V :

Der Raumbereich wird in N Teilvolumina ΔV_i aufgeteilt und die Ladung in den einzelnen Teilvolumina wird jeweils durch eine Punktladungen q_i an einem Ort $\vec{r}_i$ ersetzt. Je kleiner die Teilbereiche ΔV_i gewählt werden, desto kleiner ist am Ende der Fehler.

$$\begin{aligned} \varrho(\vec{r}) &= \int \varrho(\vec{r}') \delta(\vec{r} - \vec{r}') d^3r' = \sum_{i=1}^N \int_{\Delta V_i} d^3r' \varrho(\vec{r}') \delta(\vec{r} - \vec{r}') \\ &\approx \sum_{i=1}^N \underbrace{\left[\int_{\Delta V_i} d^3r' \varrho(\vec{r}') \right]}_{= q_i} \delta(\vec{r} - \vec{r}_i) = \sum_{i=1}^N q_i \delta(\vec{r} - \vec{r}_i) . \end{aligned} \quad (10.37)$$

Wir überprüfen (10.37) durch die Berechnung der Gesamtladung Q im Gesamtvolumen V :

$$\int_V \varrho(\vec{r}) d^3r = \int_V \sum_{i=1}^N q_i \delta(\vec{r} - \vec{r}_i) d^3r = \sum_{i=1}^N \int_V q_i \delta(\vec{r} - \vec{r}_i) d^3r = \sum_{i=1}^N q_i = Q . \quad \square$$

11 Das Wichtigste zu reellen Matrizen

Notation:

Wenn wir im Folgenden von Matrizen sprechen, sind nur reelle Matrizen gemeint. Wir verwenden für die Notation von Matrizen einen Mix aus verschiedenen gebräuchlichen Notationen und schreiben für eine gesprochen „**m-Kreuz-n-Matrix A**“ bzw. für eine Matrix **A** vom Typ (m, n)

$$(m \times n)\text{-Matrix } \mathbf{A} = \mathbf{A}^{(m,n)} = (a_{ij})^{(m,n)}$$

mit

$i \in \mathbb{N}$	Zeilenindex $i = \{1, \dots, m\}$,
$j \in \mathbb{N}$	Spaltenindex $j = \{1, \dots, n\}$,
a_{ij}	Elemente (Matrixelemente) oder Koeffizienten von A ,
$m \times n$	Dimension der Matrix A = Anzahl ihrer Matrixelemente .

Für die Determinante einer Matrix **A** schreiben wir

$$\det \mathbf{A} = |\mathbf{A}| = \det (a_{ij}).$$

11.1 Assoziativgesetz für die Matrizenmultiplikation

Wenn die Formate (Anzahl von Zeilen und Spalten) der Matrizen es ermöglichen, gilt für die Matrizenmultiplikation

$$\boxed{\mathbf{A}(\mathbf{BC}) = (\mathbf{AB})\mathbf{C}}.$$

Beweis mit $\mathbf{A} = (a_{ij})^{(m,n)}$, $\mathbf{B} = (b_{jk})^{(n,r)}$, $\mathbf{C} = (c_{kl})^{(r,s)}$:

$$\begin{aligned} \mathbf{A}(\mathbf{BC}) &= \left(\sum_{j=1}^n a_{ij} \left(\sum_{k=1}^r b_{jk} c_{kl} \right) \right) \\ &= \left(\sum_{k=1}^r \left(\sum_{j=1}^n a_{ij} b_{jk} \right) c_{kl} \right) = (\mathbf{AB})\mathbf{C} \\ &= (d_{il})^{(m,s)} = \mathbf{D} \end{aligned}$$

oder bei Summation über doppelt auftretende Indizes

$$(a_{ij}) \cdot \underbrace{\left[(b_{jk}) \cdot (c_{kl}) \right]}_{(g_{jl})} = (d_{il}) = \underbrace{\left[(a_{ij}) \cdot (b_{jk}) \right]}_{(h_{ik})} \cdot (c_{kl}) = (d_{il}) . \quad \square$$

Die Faktoren, d. h. die Matrizen kommutieren dabei im Allgemeinen nicht. Sie dürfen deshalb in ihrer Reihenfolge allgemein nicht vertauscht werden.

11.2 Die Transponierte einer Matrix

Beim Transponieren einer Matrix werden ihre Zeilen und Spalten miteinander vertauscht, d. h. die erste Zeile wird zur ersten Spalte, die zweite Zeile wird zur zweiten Spalte usw. Aus einer

$$(m \times n)\text{-Matrix } \mathbf{A} = (a_{ij})^{(m,n)}$$

wird so deren Transponierte, die

$$(n \times m)\text{-Matrix } \mathbf{B} = (b_{ij})^{(n,m)} = (a_{ji})^{(n,m)} = (a_{ij})^T = \mathbf{A}^T .$$

Für das einzelne Matrixelement bedeutet die Transposition T

$$a_{ij} \xrightarrow{T} a_{ji} = a_{ij}^T = b_{ij} .$$

Obwohl sich ein Matrixelement allein nicht transponieren lässt, schreiben wir hier *symbolisch* a^T für b .

Beispiel:

$$\begin{aligned} \mathbf{A} &= (a_{ij})^{(2,3)} = \begin{pmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{pmatrix} \xrightarrow{T} \\ \mathbf{A}^T &= (a_{ji})^{(3,2)} = \begin{pmatrix} a_{11} & a_{21} \\ a_{12} & a_{22} \\ a_{13} & a_{23} \end{pmatrix} = \begin{pmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \\ b_{31} & b_{32} \end{pmatrix} = (b_{ij})^{(3,2)} = \mathbf{B} . \end{aligned}$$

11.3 Multiplikation transponierter Matrizen

Wenn die Formate der Matrizen es ermöglichen, gilt

$$\boxed{(\mathbf{A} \cdot \mathbf{B})^T = \mathbf{B}^T \cdot \mathbf{A}^T} .$$

Beweis mit

$$\begin{aligned} \mathbf{A} &= (a_{ij})^{(m,n)} \Rightarrow \mathbf{A}^T = (a_{ji})^{(n,m)} , \\ \mathbf{B} &= (b_{jk})^{(n,r)} \Rightarrow \mathbf{B}^T = (b_{kj})^{(r,n)} , \\ \mathbf{A} \cdot \mathbf{B} &= \left(\sum_{j=1}^n a_{ij} \cdot b_{jk} \right) = (c_{ik})^{(m,r)} \Rightarrow (\mathbf{A} \cdot \mathbf{B})^T = (c_{ki})^{(r,m)} : \\ \mathbf{B}^T \cdot \mathbf{A}^T &= \left(\sum_{j=1}^n b_{kj} \cdot a_{ji} \right) = (c_{ki})^{(r,m)} = (\mathbf{A} \cdot \mathbf{B})^T . \quad \square \end{aligned}$$

11.4 Quadratische Matrizen

Als Beispiel verwenden wir o.B.d.A. eine (4×4) -Matrix $\mathbf{A} = (a_{ij})^{(4,4)}$ mit der zugehörigen 4-reihigen Determinante bzw. Determinante 4. Ordnung $|\mathbf{A}| = \det(a_{ij})$.

- Eine quadratische Matrix ist beispielsweise die $n \times n$ -Matrix $\mathbf{A} = (a_{ij})^{(n,n)}$.
- Eine quadratische Matrix $\mathbf{A}$ ist **singulär**, wenn $\det \mathbf{A} = 0$.
- Eine quadratische Matrix ist **regulär**, wenn ihre Determinante von Null verschieden ist:

$$\det \mathbf{A} \neq 0 \quad \Rightarrow \quad \mathbf{A} = (a_{ij})^{(n,n)} \text{ regulär.}$$

Die Begriffe „singuläre Matrix“ und „reguläre Matrix“ sind nur für quadratische Matrizen definiert.

- **Minor** (synonym: **Unterdeterminante**) D_{ij} :

Durch Streichen der i -ten Zeile und der j -ten Spalte einer Determinante $|\mathbf{A}|$ erhalten wir den Minor bzw. die Unterdeterminante D_{ij} zum Element a_{ij} :
beispielsweise

$$|\mathbf{A}| = \begin{vmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{vmatrix} \quad \Rightarrow \quad D_{ij} = D_{32} = \begin{vmatrix} a_{11} & a_{13} & a_{14} \\ a_{21} & a_{23} & a_{24} \\ a_{41} & a_{43} & a_{44} \end{vmatrix}.$$

- **Adjunkte** (synonym: **algebraisches Komplement**) A_{ij} :

Die Adjunkte bzw. das algebraische Komplement A_{ij} erhalten wir durch Multiplikation der Unterdeterminante D_{ij} mit dem Faktor $(-1)^{i+j}$ entsprechend dem sogenannten

$$\begin{array}{cccc} + & - & + & \cdots \\ - & + & - & \cdots \\ + & - & + & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{array} \quad \text{für Determinanten:}$$

Vorzeichen-Schachbrettmuster

$$A_{ij} = (-1)^{i+j} \cdot D_{ij}.$$

- Mit den Adjunkten A_{ij} einer Matrix $\mathbf{A}$ lassen sich n -reihige Determinanten bzw. Determinanten n -ter Ordnung berechnen nach dem **Laplace'schen Entwicklungssatz**:

$$\text{Entwicklung nach der } i\text{-ten Zeile:} \quad \det \mathbf{A} = \sum_{j=1}^n a_{ij} A_{ij} \quad (i = 1, 2, \dots, n),$$

$$\text{Entwicklung nach der } j\text{-ten Spalte:} \quad \det \mathbf{A} = \sum_{i=1}^n a_{ij} A_{ij} \quad (j = 1, 2, \dots, n).$$

Der Wert der 1-reihigen Determinante ist gleich dem Wert des einzigen Matrixelements.

2-reihige Determinanten werden wie folgt berechnet:

$$\begin{aligned}\det \mathbf{A} &= \det (a_{ij})^{(2,2)} = a_{11} a_{22} - a_{12} a_{21} \\ &= a_{11} a_{22} - a_{21} a_{12} .\end{aligned}$$

Nicht nur für 2-reihige Determinanten, wie man hier sehen kann, sondern auch allgemein gilt folglich

$$\det \mathbf{A} = \det \mathbf{A}^T \quad (\text{Stürzen der Determinante } |\mathbf{A}|) .$$

3-reihige Determinanten werden beispielsweise durch Entwicklung nach der 1. Zeile wie folgt berechnet:

$$\begin{aligned}\det \mathbf{A} &= \det (a_{ij})^{(3,3)} \\ &= a_{11} D_{11} - a_{12} D_{12} + a_{13} D_{13} \quad (D_{1j} \text{ Unterdeterminanten zur 1. Zeile}) \\ &= a_{11} A_{11} + a_{12} A_{12} + a_{13} A_{13} \quad (A_{1j} \text{ Adjunkten zur 1. Zeile}) .\end{aligned}$$

4-reihige Determinanten werden beispielsweise durch Entwicklung nach der 1. Zeile wie folgt berechnet:

$$\det \mathbf{A} = \det (a_{ij})^{(4,4)} = \sum_{j=1}^4 a_{1j} A_{1j} = a_{11} A_{11} + a_{12} A_{12} + a_{13} A_{13} + a_{14} A_{14} ,$$

wobei sich die Adjunkten zur 1. Zeile wie folgt aus den 3-reihigen Unterdeterminanten zur 1. Zeile ergeben:

$$A_{11} = D_{11} , \quad A_{12} = -D_{12} , \quad A_{13} = D_{13} , \quad A_{14} = -D_{14} .$$

• **adjungierte Matrix $\mathbf{A}_{\text{adj}}$:**

Die zu einer Matrix $\mathbf{A}$ adjungierte Matrix $\mathbf{A}_{\text{adj}}$ ist die Transponierte $(A_{ij})^T$ der Matrix (A_{ij}) , gebildet aus den Adjunkten A_{ij} von $\mathbf{A}$, also beispielsweise

$$\mathbf{A}_{\text{adj}} = (A_{ij})^T = \begin{pmatrix} A_{11} & A_{12} & A_{13} & A_{14} \\ A_{21} & A_{22} & A_{23} & A_{24} \\ A_{31} & A_{32} & A_{33} & A_{34} \\ A_{41} & A_{42} & A_{43} & A_{44} \end{pmatrix}^T = \begin{pmatrix} A_{11} & A_{21} & A_{31} & A_{41} \\ A_{12} & A_{22} & A_{32} & A_{42} \\ A_{13} & A_{23} & A_{33} & A_{43} \\ A_{14} & A_{24} & A_{34} & A_{44} \end{pmatrix} .$$

11.4.1 Invertieren von Matrizen

Invertieren kann man nur quadratische Matrizen. Eine (quadratische) Matrix $\mathbf{A}$ besitzt genau dann eine Inverse, wenn sie regulär ist, d. h. wenn $|\mathbf{A}| \neq 0$ ist. Daraus folgt mit der **Einheitsmatrix** $\mathbb{1}$ und mit dem **Multiplikationstheorem** für Determinanten $|\mathbf{C}| = |\mathbf{AB}| = |\mathbf{A}| \cdot |\mathbf{B}|$:

$$\mathbf{A} \cdot \mathbf{A}^{-1} = \mathbf{A}^{-1} \cdot \mathbf{A} = \mathbb{1} \quad \Rightarrow$$

$$|\mathbf{A}^{-1} \cdot \mathbf{A}| = |\mathbb{1}| = |\mathbf{A}^{-1}| \cdot |\mathbf{A}| = 1 \quad \Leftrightarrow \quad |\mathbf{A}^{-1}| = \frac{1}{|\mathbf{A}|} \neq 0 .$$

Zunächst zeigen wir am Beispiel der 4×4 -Matrix $\mathbf{A} = (a_{ij})^{(4,4)}$, dass

$$\boxed{\mathbf{A} \cdot \mathbf{A}_{\text{adj}} = \mathbf{A} \cdot (A_{ij})^T = |\mathbf{A}| \cdot \mathbb{1}} : \quad (11.1)$$

$$\begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{pmatrix} \cdot \begin{pmatrix} A_{11} & A_{21} & A_{31} & A_{41} \\ A_{12} & A_{22} & A_{32} & A_{42} \\ A_{13} & A_{23} & A_{33} & A_{43} \\ A_{14} & A_{24} & A_{34} & A_{44} \end{pmatrix} = \begin{pmatrix} |\mathbf{A}| & 0 & 0 & 0 \\ 0 & |\mathbf{A}| & 0 & 0 \\ 0 & 0 & |\mathbf{A}| & 0 \\ 0 & 0 & 0 & |\mathbf{A}| \end{pmatrix} = |\mathbf{A}| \cdot \mathbb{1} . \quad (11.2)$$

Die Matrixelemente $|\mathbf{A}|$ der Diagonalen gehen gemäß des Laplace'schen Entwicklungssatzes beispielsweise zurück auf die Determinante, gebildet aus dem Produkt der 2. Zeile von $\mathbf{A}$ mit der 2. Spalte von $\mathbf{A}_{\text{adj}}$:

$$\begin{pmatrix} a_{21} & a_{22} & a_{23} & a_{24} \end{pmatrix} \cdot \begin{pmatrix} A_{21} \\ A_{22} \\ A_{23} \\ A_{24} \end{pmatrix} = a_{21} A_{21} + a_{22} A_{22} + a_{23} A_{23} + a_{24} A_{24} = |\mathbf{A}|$$

ist die Entwicklung nach der 2. Zeile: $|\mathbf{A}| = \begin{vmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ \mathbf{a_{21}} & \mathbf{a_{22}} & \mathbf{a_{23}} & \mathbf{a_{24}} \\ a_{31} & a_{32} & a_{33} & a_{34} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{vmatrix} .$

Dass die Matrixelemente außerhalb der Diagonalen alle gleich Null sind, ergibt sich aus der Rechenregel für Determinanten, nach der eine Determinante dann den Wert Null hat, wenn zwei Parallelreihen miteinander übereinstimmen. So geht beispielsweise das Produkt der 1. Zeile von $\mathbf{A}$ mit der 3. Spalte von $\mathbf{A}_{\text{adj}}$, also

$$\begin{pmatrix} a_{11} & a_{12} & a_{13} & a_{14} \end{pmatrix} \cdot \begin{pmatrix} A_{31} \\ A_{32} \\ A_{33} \\ A_{34} \end{pmatrix} = a_{11} A_{31} + a_{12} A_{32} + a_{13} A_{33} + a_{14} A_{34} = 0 , \quad (11.3)$$

zurück auf die Determinante

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} & a_{14} \\ a_{21} & a_{22} & a_{23} & a_{24} \\ a_{11} & a_{12} & a_{13} & a_{14} \\ a_{41} & a_{42} & a_{43} & a_{44} \end{vmatrix} = 0 ,$$

weil bei der Bildung der Adjunkten nicht die 1. Zeile der Matrix $\mathbf{A}$ „gestrichen“ wird, wie es für die Entwicklung nach der 1. Zeile erforderlich wäre, sondern die 3. Zeile, sodass im Produkt (11.3) statt der 3. Zeile die 1. Zeile noch einmal und damit doppelt auftritt.

Wer dies nicht sofort überschaut, sollte sich die Mühe machen, eine 3-reihige Determinante $\det(a_{ij})^{(3,3)}$ mit dem Laplace'schen Entwicklungssatz beispielsweise nach der ersten Zeile zu entwickeln, aber nicht mit den Adjunkten zur 1. Zeile sondern mit den Adjunkten zur 2. Zeile. Man stellt dann fest, dass die auf diese eigentlich falsche Weise entwickelte Determinante den Wert Null hat.

Anschließend ist in der gleichen Determinante $\det(a_{ij})^{(3,3)}$ die 2. Zeile durch die 1. Zeile zu ersetzen, sodass die 1. Zeile doppelt auftritt, und dann die Determinante mit der Regel von Sarrus nach der 1. Zeile zu berechnen. Wieder hat die Determinante den Wert Null.

Aber **Achtung!** Dem Vorzeichen-Schachbrettmuster entsprechend ist die Reihenfolge der Vorzeichen der Adjunkten zur 2. Zeile gegenüber der Vorzeichenreihenfolge der Adjunkten zur 1. Zeile umgekehrt. Diese Vorzeichenumkehr tritt immer dann auf, wenn sich die Entwicklungsreihe für die Adjunkten um eine *ungerade* Zahl von der „richtigen“ Entwicklungsreihe für die Determinante unterscheidet, was aber keinen Einfluss auf das Ergebnis in (11.2) hat.

- Um jetzt die Formel zur Berechnung der Inversen von $\mathbf{A}$ mit Hilfe der adjungierten Matrix $\mathbf{A}_{\text{adj}}$ zu erhalten, multiplizieren wir (11.1) von links mit $\mathbf{A}^{-1}$:

$$\underbrace{\mathbf{A}^{-1} \cdot \mathbf{A}}_{=1} \cdot \mathbf{A}_{\text{adj}} = \mathbf{A}^{-1} \cdot |\mathbf{A}| \cdot \mathbb{1} = |\mathbf{A}| \cdot \mathbf{A}^{-1} \quad \Leftrightarrow$$

$$\boxed{\mathbf{A}^{-1} = \frac{1}{|\mathbf{A}|} \cdot \mathbf{A}_{\text{adj}}}$$

- Invertieren von transponierten Matrizen:

$$\boxed{(\mathbf{A}^T)^{-1} = (\mathbf{A}^{-1})^T},$$

denn

$$\underbrace{(\mathbf{A}^T)^{-1} \cdot \mathbf{A}^T}_{=1} = (\mathbf{A}^{-1})^T \cdot \mathbf{A}^T = (\mathbf{A} \cdot \mathbf{A}^{-1})^T = \mathbb{1}^T = \mathbb{1} \quad \square$$

- Invertieren von Matrizenprodukten:

$$\boxed{(\mathbf{A} \cdot \mathbf{B})^{-1} = \mathbf{B}^{-1} \cdot \mathbf{A}^{-1}},$$

denn

$$(\mathbf{A} \cdot \mathbf{B})(\mathbf{B}^{-1} \cdot \mathbf{A}^{-1}) = \mathbf{A} \cdot \underbrace{(\mathbf{B}\mathbf{B}^{-1})}_{=1} \cdot \mathbf{A}^{-1} = \mathbf{A}\mathbf{A}^{-1} = \mathbb{1} \quad \square$$

11.4.2 Orthogonale Matrizen – Drehmatrizen

Orthogonale Matrizen $\mathbf{O}$ sind reguläre quadratische Matrizen mit der Eigenschaft

$$\det \mathbf{O} = \begin{cases} -1 & \text{für Spiegelungsmatrizen,} \\ +1 & \text{für Drehungs- oder Drehmatrizen } \mathbf{D} \text{ im Rechtssystem.} \end{cases}$$

Uns interessieren hier nur die **Drehmatrizen** $\mathbf{D}$. Sie beschreiben Drehungen von einem Rechtssystem in ein Rechtssystem und haben die folgenden Eigenschaften:

- **Orthogonalität:** $\mathbf{D}^T = \mathbf{D}^{-1} \Leftrightarrow \mathbf{D}\mathbf{D}^T = \mathbf{D}^T\mathbf{D} = \mathbb{1}$.

Die Orthogonalität der Drehmatrizen impliziert folglich ihre Reihen-Orthonormalität, d. h. ihre Zeilen- und Spalten-Orthonormalität, was man am einfachsten durch Drehung einer Standardbasis $\{\vec{e}_i\}$ in die Standardbasis $\{\vec{e}_i'\}$ mittels $\mathbf{D}^{-1}$ und die anschließende Rückdrehung mittels $\mathbf{D}$ zeigen kann.

- Produkte orthogonaler Matrizen sind wieder orthogonal.

Nacheinander ausgeführte Drehungen mit den Drehmatrizen $\mathbf{A}$, $\mathbf{B}$, $\mathbf{C}$:

$$\text{erste Drehung : } \mathbf{A} \cdot \vec{x} = \vec{x}',$$

$$\text{zweite Drehung : } \mathbf{B} \cdot \vec{x}' = \vec{x}'',$$

$$\text{Gesamtdrehung : } \mathbf{C} \cdot \vec{x} = \mathbf{B} \cdot (\mathbf{A} \cdot \vec{x}) = (\mathbf{B} \cdot \mathbf{A}) \cdot \vec{x}.$$

12 Rechnen mit komplexen Vektoren und Matrizen bzw. Operatoren

12.1 Rechenregeln

- In bra-ket-Notation (Dirac-Notation) sind basisfreie (abstrakte) Vektoren z. B. der bra-Vektor (Zeilenvektor) $\langle u|$ und der ket-Vektor (Spaltenvektor) $|v\rangle$. Sie bilden das komplexe Standardskalarprodukt $\langle u|v\rangle$. Sie lassen sich in einer vollständigen Orthonormalbasis (VON-Basis), z. B. $\{|a\rangle_i\}$, darstellen bzw. entwickeln, sodass wir dann mit den Entwicklungskoeffizienten u_i und v_i schreiben können:

$$|v\rangle = \sum_i \langle a_i|v\rangle |a_i\rangle = \sum_i v_i |a_i\rangle ,$$

$$|v\rangle := \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_i \\ \vdots \end{pmatrix} = v \quad \Rightarrow \quad |v\rangle \neq v ,$$

$$\begin{aligned} \langle u| &= \sum_i \langle a_i|\langle u|a_i\rangle = \sum_i \langle a_i|u_i^* \stackrel{u_i^* \text{ ist Skalar}}{=} \sum_i u_i^* \langle a_i| \\ &= \sum_i \underbrace{\langle u|a_i\rangle}_{=u_i^*} \langle a_i| \stackrel{(12.1)}{=} \langle u| \sum_i |a_i\rangle \langle a_i| \stackrel{(12.2)}{=} \langle u|\mathbb{1} = \langle u| , \\ \langle u| &:= (u_1^* \quad u_2^* \quad \cdots \quad u_i^* \quad \cdots) = (u^*)^T \quad \Rightarrow \quad \langle u| \neq (u^*)^T . \end{aligned}$$

Basisfreie (abstrakte) Operatoren bezeichnen wir z. B. mit $\hat{A}$, Matrizen aber mit $A = (A_{ij})$ und die Matrixelemente mit A_{ij} . In einer zum Operator $\hat{A}$ passenden VON-Basis $\{|a\rangle_i\}$ erhält $\hat{A}$ die Darstellung

$$\hat{A} = \sum_{i,j} |a_i\rangle \langle a_i|\hat{A}|a_j\rangle \langle a_j| = \sum_{i,j} A_{ij} |a_i\rangle \langle a_j|$$

mit den Matrixelementen

$$\langle a_i|\hat{A}|a_j\rangle = A_{ij}$$

und der aus diesen A_{ij} gebildeten und zum Operator $\hat{A}$ gehörenden Matrix

$$(A_{ij}) = A \quad \Rightarrow \quad \hat{A} \neq A .$$

- Adjungiert (hochgestellter Index $\dagger$) heißt komplex konjugiert (hochgestellter Index $*$) und zusätzlich transponiert (hochgestellter Index T), also sinngemäß

$$\dagger = (*) \wedge T = (T) \wedge * .$$

Beispiel: Matrix A :

$$A^\dagger = (A^*)^T = (A^T)^* ,$$

Beispiel: Zustands-ket-Vektor $|v\rangle$ (entspricht einem Spaltenvektor):

$$(|v\rangle^*)^T = |v\rangle^\dagger = \langle v| .$$

Der resultierende Zustands-bra-Vektor $\langle v|$ entspricht einem Zeilenvektor, dessen Elemente (Komponenten) die komplex konjugierten Elemente des zugehörigen ket-Vektors $|v\rangle$ sind.

Beispiel: Spaltenvektor $\vec{v} = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}$ analog zu $|v\rangle$:

$$(\vec{v}^*)^T = \vec{v}^\dagger = (v_1^*, v_2^*) .$$

- A ist hermitesch, wenn $A = A^\dagger$, beispielsweise

$$A = \begin{pmatrix} a & ic \\ -ic & b \end{pmatrix} = A^\dagger, \quad \begin{pmatrix} a \\ b \\ c \end{pmatrix} \in \mathbb{R} .$$

- A ist antihermitesch, wenn $A = -A^\dagger$ bzw. $iA = (iA)^\dagger$, beispielsweise

$$A = \begin{pmatrix} -ia & c \\ -c & ib \end{pmatrix} = -A^\dagger, \quad \begin{pmatrix} a \\ b \\ c \end{pmatrix} \in \mathbb{R} .$$

- Ein linearer Operator A (bzw. eine lineare Matrix) lässt sich zerlegen in einen hermiteschen Operator A_h und einen antihermiteschen Operator A_a gemäß

$$A = A_h + A_a, \quad A_h = \frac{A + A^\dagger}{2}, \quad A_a = \frac{A - A^\dagger}{2} .$$

- Die Verallgemeinerung des transponierten Produkts $(A \cdot B)^T = B^T \cdot A^T$ aus reellen Matrizen A, B ist das adjungierte Produkt aus den (komplexen) Matrizen A, B

$$(A \cdot B)^\dagger = B^\dagger \cdot A^\dagger .$$

- Adjungieren eines Zustands-ket-Vektors $|v\rangle$:

$$|v\rangle^\dagger = \langle v| .$$

Veranschaulichung: $\begin{pmatrix} ia \\ ib \end{pmatrix}^\dagger = (-ia, -ib) .$

- Das komplexe Standardskalarprodukt (kurz Skalarprodukt)

$$(\alpha^* u_1^* \quad \alpha^* u_2^* \quad \dots) \cdot \begin{pmatrix} \beta v_1 \\ \beta v_2 \\ \vdots \end{pmatrix} = \alpha^* \beta (u_1^* v_1 + u_2^* v_2 + \dots)$$

liefert also mit $\alpha, \beta \in \mathbb{C}$ einen Skalar gemäß

$$\begin{aligned} \langle \alpha u | \beta v \rangle &= \alpha^* \langle u | \beta v \rangle = \alpha^* \beta \langle u | v \rangle \\ &= \left(\langle \beta v | \alpha u \rangle \right)^* = \left(\beta^* \alpha \langle v | u \rangle \right)^* = \alpha^* \beta (\langle v | u \rangle)^* = \alpha^* \beta \langle u | v \rangle . \end{aligned}$$

Das Skalarprodukt ist also antilinear (konjugiert linear) im ersten und linear im zweiten Argument.

- Das Skalarprodukt c_i aus dem Zustandsvektor $|v\rangle$ und dem Basisvektor $|a_i\rangle$, also $c_i = \langle a_i|v\rangle$, ist die Projektion von $|v\rangle$ auf $|a_i\rangle$, gesprochen: „Skalarprodukt v in a_i “. c_i ist somit die komplexe skalare Vektorkomponente von $|v\rangle$ „in Richtung“ des Basisvektors $|a_i\rangle$.
- Adjungieren des Skalarprodukts $\langle u|v\rangle$:

$$\begin{aligned} (\langle u|v\rangle)^\dagger &= |v\rangle^\dagger \cdot \langle u|^\dagger = \langle v|u\rangle = \sum_i v_i^* u_i \\ &= (\langle u|v\rangle)^* = \left(\sum_i u_i^* v_i \right)^* \Rightarrow \end{aligned}$$

$$(\langle u|v\rangle)^\dagger = (\langle u|v\rangle)^* = \langle v|u\rangle.$$

- Adjungieren eines Matrix-Vektor-Produkts:

$$A|u\rangle = |v\rangle \Rightarrow$$

$$(A|u\rangle)^\dagger = (|u\rangle)^\dagger A^\dagger = \langle u|A^\dagger = |v\rangle^\dagger = \langle v|,$$

wobei $A^\dagger$ rechts von $\langle u|$ stehen muss, u. a. weil $\langle u|$ ein Zeilenvektor ist.

- Adjungieren eines Matrix-Matrix-Vektor-Produkts:

$$\left[A (B|u\rangle) \right]^\dagger = |v\rangle^\dagger = \langle v| = (\langle u|B^\dagger) A^\dagger.$$

Hierbei ist die Reihenfolge von Matrizen und Vektor zu beachten.

- Das dyadische Produkt liefert eine Matrix gemäß

$$|u\rangle\langle v| = \begin{pmatrix} u_1 \\ u_2 \\ \vdots \end{pmatrix} \cdot (v_1^* \quad v_2^* \quad \dots) = \begin{pmatrix} u_1 v_1^* & u_1 v_2^* & \dots \\ u_2 v_1^* & u_2 v_2^* & \dots \\ \vdots & \vdots & \ddots \end{pmatrix}.$$

Die Vektoren (Faktoren) des dyadischen Produkts dürfen in ihrer Reihenfolge nicht vertauscht werden, wie man in der Veranschaulichung an einem einfachen Beispiel sofort erkennt.

Adjungieren eines dyadischen Produkts:

$$(|u\rangle\langle v|)^\dagger = \langle v|^\dagger |u\rangle^\dagger = |v\rangle\langle u|.$$

Mit den Basisvektoren $|a_i\rangle$ der VON-Basis¹ $\{|a_i\rangle\}$ ist der Einsoperator $\mathbb{1}$, auch Identitätsoperator genannt, die Summe aus den

$$\text{Projektionsoperatoren } P_i = |a_i\rangle\langle a_i|.$$

¹VON-Basis heißt vollständige Orthonormalbasis und ist gleichbedeutend mit der Bezeichnung vollständiges Orthonormalsystem (VONS). Meistens ist die VON-Basis gemeint, wenn kurz von einer Basis die Rede ist.

Für die orthonormierten Basisvektoren gilt $\langle a_i | a_j \rangle = \langle a_i | a_j \rangle^* = \delta_{ij}$. Wir veranschaulichen die Darstellung des Einsoperators für $i = 1, 2$:

$$\sum_{i=1}^2 P_i = \sum_{i=1}^2 |a_i\rangle \langle a_i| \quad (12.1)$$

$$\begin{aligned} &= \begin{pmatrix} a_1 \\ 0 \end{pmatrix} \cdot (a_1^* \ 0) + \begin{pmatrix} 0 \\ a_2 \end{pmatrix} \cdot (0 \ a_2^*) \\ &= \begin{pmatrix} a_1 a_1^* & 0 \\ 0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & 0 \\ 0 & a_2 a_2^* \end{pmatrix} \\ &= \begin{pmatrix} \langle a_1 | a_1 \rangle & 0 \\ 0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & 0 \\ 0 & \langle a_2 | a_2 \rangle \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = \mathbb{1} . \quad \square \end{aligned} \quad (12.2)$$

- In Analogie zur Darstellung eines Vektors $\vec{v}$ durch seine skalaren Vektorkomponenten v_1, v_2, v_3 in der VON-Basis $\{\vec{e}_1, \vec{e}_2, \vec{e}_3\}$ im $\mathbb{R}^3$ gemäß

$$\begin{aligned} \vec{v} &= (\vec{e}_1 \cdot \vec{v}) \cdot \vec{e}_1 + (\vec{e}_2 \cdot \vec{v}) \cdot \vec{e}_2 + (\vec{e}_3 \cdot \vec{v}) \cdot \vec{e}_3 \\ &= \sum_{i=1}^3 \underbrace{(\vec{e}_i \cdot \vec{v})}_{v_i} \cdot \vec{e}_i = \sum_{i=1}^3 v_i \cdot \vec{e}_i = \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} \end{aligned}$$

wird die Darstellung (Entwicklung) eines (Zustands)vektors $|v\rangle$ nach der VON-Basis $\{|a_i\rangle\}$ im $\mathbb{C}^n$ beschrieben durch den **Entwicklungssatz**

$$\begin{aligned} |v\rangle &= \sum_{i=1}^n \underbrace{\langle a_i | v \rangle}_{v_i} |a_i\rangle = \sum_{i=1}^n v_i |a_i\rangle := \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_i \\ \vdots \end{pmatrix} = \begin{pmatrix} \langle a_1 | v \rangle \\ \langle a_2 | v \rangle \\ \vdots \\ \langle a_i | v \rangle \\ \vdots \end{pmatrix} \quad (12.3) \\ &= \sum_{i=1}^n \underbrace{|a_i\rangle \langle a_i|}_{P_i} |v\rangle = \sum_{i=1}^n P_i |v\rangle \\ &= \mathbb{1} |v\rangle = |v\rangle . \end{aligned}$$

Die Projektionen von $|v\rangle$ auf $|a_i\rangle$ bzw. die Skalarprodukte $|v\rangle$ in $|a_i\rangle$ sind die **Entwicklungskoeffizienten** $v_i \in \mathbb{C}$. In (12.3) steht ganz bewusst nicht das Gleichheitszeichen, sondern das $:=$ -Zeichen („definiert durch“). Wir werden auf diesen Umstand im Zusammenhang mit der Matrixdarstellung von Operatoren im Abschnitt 12.3 zurückkommen.

- Darstellung des **komplexen Standardskalarprodukts** in der VON-Basis $\{a_i\}$:
Wir gehen dabei aus vom basisfreien Standardskalarprodukt $\langle u | v \rangle$ und müssen berücksichtigen, dass das komplexe Standardskalarprodukt **hermitesch** ist gemäß

$$\langle u | v \rangle = (\langle v | u \rangle)^* \xrightarrow{\text{i.A.}} \langle u | v \rangle \neq \langle v | u \rangle : \quad (12.4)$$

$$\begin{aligned}
\langle u|v\rangle &= \sum_{ij} \langle a_i|\langle u|a_i\rangle \langle a_j|v\rangle|a_j\rangle = \sum_{ij} \langle a_i|u_i^* v_j|a_j\rangle \\
&= \sum_{ij} u_i^* v_j \underbrace{\langle a_i|a_j\rangle}_{\delta_{ij}} = \sum_i u_i^* v_i .
\end{aligned}$$

Die Beziehung $\langle u|v\rangle = (\langle v|u\rangle)^*$ zeigt die **Konjugationssymmetrie** des hermiteschen Skalarprodukts.

- Aus der **Hermitezitat** in (12.4) folgt fur das Standardskalarprodukt eines komplexen Vektors $|v\rangle$ mit sich selbst:

$$\langle v|v\rangle = (\langle v|v\rangle)^* \Rightarrow \begin{cases} \langle v|v\rangle \in \mathbb{R} , \\ \langle v|v\rangle \geq 0 \text{ und somit } \mathbf{positiv\ defin\it} . \end{cases} \quad (12.5)$$

(12.5) lasst sich zuruckfuhren auf das Normquadrat einer komplexen Zahl z :

$$z^* \cdot z = (a - ib)(a + ib) = a^2 + b^2 = |z|^2 \Rightarrow \begin{cases} z^* z \in \mathbb{R} , \\ z^* z \geq 0 . \end{cases}$$

12.2 Veranschaulichung des komplexen Standardskalarprodukts und der Multiplikation komplexer Vektoren mit komplexen Matrizen

Summen und Produkte aus komplexen Zahlen werde komplex konjugiert, indem alle Summanden und Faktoren jeweils für sich komplex konjugiert werden. Wir zeigen dies für das Produkt aus zwei komplexen Zahlen:

$$\begin{aligned} [(a + ib)(c + id)]^* &= [ac + i(ad + bc) - bd]^* \\ &= (a - ib)(c - id) = ac - i(ad + bc) - bd , \end{aligned}$$

$$\begin{aligned} [(a - ib)(c - id)]^* &= [ac - i(ad + bc) - bd]^* \\ &= (a + ib)(c + id) = ac + i(ad + bc) - bd , \end{aligned}$$

$$\begin{aligned} [(a + ib)(c - id)]^* &= [ac - i(ad - bc) + bd]^* \\ &= (a - ib)(c + id) = ac + i(ad - bc) + bd . \end{aligned}$$

Das komplexe Standardskalarprodukt (kurz Skalarprodukt) ist definiert durch

$$\langle u|v \rangle := \sum_{i=1}^n u_i^* \cdot v_i , \quad u_i, v_i \in \mathbb{C} .$$

Hierbei ist $|v\rangle$ ein komplexer Spaltenvektor und $\langle u|$ ein komplexer Zeilenvektor. Der Zusammenhang zwischen einem bra-Vektor $\langle w|$ und dem zugehörigen ket-Vektor $|w\rangle$ ist

$$(|w\rangle^*)^T = |w\rangle^\dagger = \langle w| .$$

Der hochgestellte Index $*$ bedeutet *komplex konjugiert*, T bedeutet *transponiert* und † bedeutet *adjungiert*, also sowohl komplex konjugiert als auch transponiert.

Zunächst zeigen wir

$$\langle u|v \rangle = \sum_i u_i^* \cdot v_i = \sum_i (v_i^* \cdot u_i)^* = \langle v|u \rangle^* \quad \Rightarrow \quad \vec{u}^* \cdot \vec{v} = (\vec{v}^* \cdot \vec{u})^*$$

am Beispiel der Vektoren²

$$|u\rangle = \begin{pmatrix} -2i \\ 2 - 3i \end{pmatrix} = \vec{u} , \quad |v\rangle = \begin{pmatrix} 1 - i \\ i \end{pmatrix} = \vec{v} :$$

$$\langle u|v \rangle = (2i, 2 + 3i) \begin{pmatrix} 1 - i \\ i \end{pmatrix} = -1 + 4i ,$$

$$\langle v|u \rangle = (1 + i, -i) \begin{pmatrix} -2i \\ 2 - 3i \end{pmatrix} = -1 - 4i ,$$

$$\langle v|u \rangle^* = -1 + 4i = \langle u|v \rangle . \quad \square$$

²Mit der Notation $\vec{v}$ für Spaltenvektoren und $\vec{v}^T$ für Zeilenvektoren ist $(\vec{u}^*)^T \cdot \vec{v} = \vec{u}^\dagger \cdot \vec{v} = \vec{u}^* \cdot \vec{v}$ das Punktprodukt $\vec{u}^* \cdot \vec{v} = \sum_i u_i^* \cdot v_i$. Man achte auf den Unterschied zwischen der Verwendung des Multiplikationspunktes $\cdot$ und der Verwendung des fettgedruckten Symbols $\cdot$.

Jetzt zeigen wir die Wirkung einer komplexen (komplexwertigen) Matrix innerhalb des Skalarprodukts am Beispiel der Matrix

$$A = \begin{pmatrix} 2+i & 3i \\ -2i & 1-i \end{pmatrix}, \quad A^\dagger = \begin{pmatrix} 2-i & 2i \\ -3i & 1+i \end{pmatrix} :$$

$$\begin{aligned} \langle u | \underbrace{A|v\rangle}_{=|w\rangle} &= \langle u|w\rangle = \langle u| \cdot (A|v\rangle) = \\ &= (2i, 2+3i) \cdot \begin{pmatrix} 2+i & 3i \\ -2i & 1-i \end{pmatrix} \cdot \begin{pmatrix} 1-i \\ i \end{pmatrix} \\ &= (2i, 2+3i) \cdot \begin{pmatrix} -i \\ -1-i \end{pmatrix} = 3-5i \end{aligned}$$

mit

$$A|v\rangle = |w\rangle = \begin{pmatrix} -i \\ -1-i \end{pmatrix}.$$

Zum gleichen Ergebnis kommen wir mit der Rechnung

$$\begin{aligned} \underbrace{\langle u|A}_{=\langle m|} |v\rangle &= \langle m|v\rangle = (\langle u|A) \cdot |v\rangle = \\ (2i, 2+3i) \cdot \begin{pmatrix} 2+i & 3i \\ -2i & 1-i \end{pmatrix} \cdot \begin{pmatrix} 1-i \\ i \end{pmatrix} &= \\ (4, -1+i) \cdot \begin{pmatrix} 1-i \\ i \end{pmatrix} &= 3-5i. \end{aligned}$$

Mit

$$\begin{aligned} (A|v\rangle)^\dagger &= |w\rangle^\dagger = \langle w| = |v\rangle^\dagger A^\dagger = \langle v|A^\dagger \\ &= (1+i, -i) \cdot \begin{pmatrix} 2-i & 2i \\ -3i & 1+i \end{pmatrix} = (i, -1+i) \end{aligned} \quad (12.6)$$

und

$$A^\dagger|u\rangle = \begin{pmatrix} 2-i & 2i \\ -3i & 1+i \end{pmatrix} \cdot \begin{pmatrix} -2i \\ 2-3i \end{pmatrix} = \begin{pmatrix} 4 \\ -1-i \end{pmatrix}$$

zeigen wir schließlich

$$\begin{aligned} \langle u|w\rangle^* &= \langle w|u\rangle = \langle v|A^\dagger|u\rangle = 3+5i : \\ (1+i, -i) \cdot \begin{pmatrix} 2-i & 2i \\ -3i & 1+i \end{pmatrix} \cdot \begin{pmatrix} -2i \\ 2-3i \end{pmatrix} &= \\ = (\langle v|A^\dagger) \cdot |u\rangle &= (i, -1+i) \cdot \begin{pmatrix} -2i \\ 2-3i \end{pmatrix} = 3+5i, \\ = \langle v| \cdot (A^\dagger|u\rangle) &= (1+i, -i) \cdot \begin{pmatrix} 4 \\ -1-i \end{pmatrix} = 3+5i. \quad \square \end{aligned}$$

Wichtig für das Verständnis und für die Praxis sind die aus (12.6) abgeleiteten und im Einklang mit $(A \cdot B)^\dagger = B^\dagger \cdot A^\dagger$ stehenden Beziehungen

$$\langle m| = \langle u|A = |u\rangle^\dagger A = (A^\dagger|u\rangle)^\dagger = |m\rangle^\dagger \quad \Rightarrow$$

$$\begin{aligned} \langle m| &= \langle u|A \quad \Leftrightarrow \quad |m\rangle = A^\dagger|u\rangle, \\ |w\rangle &= A|v\rangle \quad \Leftrightarrow \quad \langle w| = \langle v|A^\dagger. \end{aligned}$$

Im Zusammenhang mit hermiteschen Matrizen werden uns diese Beziehungen noch beschäftigen.

12.3 Matrixdarstellung von Operatoren

Nach:

Christian B. Lang und Norbert Pucker, *Mathematische Methoden in der Physik*, Hochschultaschenbuch, Spektrum Akademischer Verlag, Heidelberg, Berlin, 1998, Seite 425.

Siehe auch:

Torsten Fließbach, *Quantenmechanik, Lehrbuch zur Theoretischen Physik III*, 4. Auflage, Elsevier-Spektrum Akademischer Verlag, München, 2005, Seite 235 bis Seite 237

und

http://schwalbe.org.chemie.uni-frankfurt.de/sites/default/files/attachements/mathematische_methoden_in_der_nmr-spektroskopie/skript_zur_ubung_1.pdf.

In diesem Abschnitt zeigen wir, dass man einen Operator A als Matrix auffassen und dementsprechend mit ihm umgehen darf. Dazu gehen wir mit den Entwicklungskoeffizienten u_j von $|u\rangle$ und v_i von $|v\rangle$ aus von der (basisfreien) **Operatorgleichung**

$$|v\rangle = \hat{A} |u\rangle . \quad (12.7)$$

Zunächst zeigen wir die Entwicklung dieser Operatorgleichung nach der VON-Basis $\{|a_i\rangle\}$:

$$\begin{aligned} \sum_j \underbrace{\langle a_j | v \rangle}_{v_j} |a_j\rangle &= \hat{A} \sum_j \underbrace{\langle a_j | u \rangle}_{u_j} |a_j\rangle \\ \sum_j v_j |a_j\rangle &= \hat{A} \sum_j u_j |a_j\rangle = \sum_j u_j \hat{A} |a_j\rangle . \end{aligned}$$

Durch das Skalarprodukt mit $|a_i\rangle$ werden jetzt die Komponenten v_i „herausprojiziert“:

$$\sum_j v_j \underbrace{\langle a_i | a_j \rangle}_{\delta_{ij}} = \sum_j u_j \langle a_i | \hat{A} | a_j \rangle .$$

In Komponentenschreibweise ergibt dies

$$\boxed{v_i = \sum_j \langle a_i | \hat{A} | a_j \rangle \cdot u_j} \quad (12.8)$$

und mit der Notation für die **Matrixelemente** $\langle a_i | \hat{A} | a_j \rangle = A_{ij}$ für die Matrix

$$\boxed{(\langle a_i | \hat{A} | a_j \rangle) = (A_{ij})} . \quad (12.9)$$

Damit können wir (12.8) als **Matrixgleichung** in der Basis $\{|a_i\rangle\}$ formulieren:

$$v_i = \sum_j A_{ij} \cdot u_j \quad \Rightarrow \quad \underbrace{\begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_i \\ \vdots \end{pmatrix}}_{\text{Matrixgleichung}} = (A_{ij}) \cdot \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_j \\ \vdots \end{pmatrix} =: \underbrace{|v\rangle = \hat{A} |u\rangle}_{\text{basisfrei}} . \quad (12.10)$$

Die Matrixelemente $\langle a_i | \hat{A} | a_j \rangle$ sind Skalarprodukte. Sie wichten die $|u\rangle$ -Komponenten u_j bei der Summierung zur $|v\rangle$ -Komponente v_i .

Zum besseren Verständnis zeigen wir noch, wie sich der Operator $\hat{A}$ im

Skalarprodukt $\langle v | \hat{A} | u \rangle$

darstellt. Durch Einschleiben von zwei Identitätsoperatoren erhalten wir daraus in Komponentenschreibweise

$$\langle v | \hat{A} | u \rangle = \langle v | \mathbb{1} \hat{A} \mathbb{1} | u \rangle = \sum_{i,j} \underbrace{\langle v | a_i \rangle}_{v_i^*} \underbrace{\langle a_i | \hat{A} | a_j \rangle}_{A_{ij}} \underbrace{\langle a_j | u \rangle}_{u_j} = \sum_{i,j} v_i^* A_{ij} u_j \quad (12.11)$$

mit den skalaren Vektorkomponenten $\langle v | a_i \rangle = v_i^*$ und $\langle a_j | u \rangle = u_j$. Wenn wir die Vektoren und die Matrix (in Matrixschreibweise) vollständig ausschreiben, erkennen wir wieder die Matrixdarstellung (A_{ij}) des Operators $\hat{A}$:

$$\begin{aligned} \sum_{i,j} v_i^* A_{ij} u_j &= \begin{pmatrix} v_1^* & v_2^* & \cdots & v_i^* & \cdots \end{pmatrix} \begin{pmatrix} A_{11} & A_{12} & \cdots & A_{1j} & \cdots \\ A_{21} & A_{22} & \cdots & A_{2j} & \cdots \\ \vdots & \vdots & \ddots & \vdots & \\ A_{i1} & A_{i2} & \cdots & A_{ij} & \cdots \\ \vdots & \vdots & & \vdots & \ddots \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_j \\ \vdots \end{pmatrix} = \\ &= \begin{pmatrix} \langle v | a_1 \rangle & \langle v | a_2 \rangle & \cdots & \langle v | a_i \rangle & \cdots \end{pmatrix} \begin{pmatrix} \langle a_1 | \hat{A} | a_1 \rangle & \langle a_1 | \hat{A} | a_2 \rangle & \cdots & \langle a_1 | \hat{A} | a_j \rangle & \cdots \\ \langle a_2 | \hat{A} | a_1 \rangle & \langle a_2 | \hat{A} | a_2 \rangle & \cdots & \langle a_2 | \hat{A} | a_j \rangle & \cdots \\ \vdots & \vdots & \ddots & \vdots & \\ \langle a_i | \hat{A} | a_1 \rangle & \langle a_i | \hat{A} | a_2 \rangle & \cdots & \langle a_i | \hat{A} | a_j \rangle & \cdots \\ \vdots & \vdots & & \vdots & \ddots \end{pmatrix} \begin{pmatrix} \langle a_1 | u \rangle \\ \langle a_2 | u \rangle \\ \vdots \\ \langle a_j | u \rangle \\ \vdots \end{pmatrix} \\ &= \langle v | \mathbb{1} \hat{A} \mathbb{1} | u \rangle = \langle v | \hat{A} | u \rangle . \end{aligned}$$

Jetzt zeigen wir die im Grunde genommen äquivalente Variante der „eigentlichen“ Matrixdarstellung des Operators $\hat{A}$. Dafür benutzen wir wieder die Operatorgleichung (12.7):

$$\begin{aligned} |v\rangle &= \hat{A} |u\rangle \\ \mathbb{1} |v\rangle &= \mathbb{1} \hat{A} \mathbb{1} |u\rangle \\ \sum_i |a_i\rangle \langle a_i | v \rangle &= \sum_{i,j} |a_i\rangle \underbrace{\langle a_i | \hat{A} | a_j \rangle}_{A_{ij}} \langle a_j | u \rangle \\ \sum_i \langle a_i | v \rangle |a_i\rangle &= \underbrace{\sum_{i,j} A_{ij} |a_i\rangle \langle a_j |}_{\hat{A}} u \end{aligned} \quad (12.12)$$

$$|v\rangle = \hat{A} |u\rangle . \quad \square \quad (12.13)$$

Wie wir sehen, ist die „eigentliche“ **Matrixdarstellung des Operators $\hat{A}$**

$$\boxed{\hat{A} = \mathbb{1} \hat{A} \mathbb{1} = \sum_{i,j} |a_i\rangle \langle a_i | \hat{A} | a_j \rangle \langle a_j | = \sum_{i,j} A_{ij} |a_i\rangle \langle a_j |} . \quad (12.14)$$

Bei einem Vergleich der Matrixgleichung in (12.10) mit (12.12), der Operatorgleichung in der Basis $\{|a_i\rangle\}$, und der basisfreien Operatorgleichung (12.13) stellen wir fest, dass

$$\hat{A} = \sum_{i,j} A_{ij} |a_i\rangle \langle a_j| \neq (A_{ij}) .$$

Dies liegt daran, dass wir für (12.12) den linken Teil des Entwicklungssatzes (12.3)

$$|v\rangle = \sum_i v_i |a_i\rangle$$

verwendet haben und für (12.10) den rechten Teil

$$|v\rangle := \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_i \\ \vdots \end{pmatrix} .$$

Wenn wir nämlich die Matrixgleichung (12.10) gegenüberstellen der Operatorgleichung (12.12), berechnet nur für die Komponente v_i , sehen wir

$$(12.10) \quad v_i = \sum_j A_{ij} u_j \quad \neq \quad (12.12) \quad v_i |a_i\rangle = \left(\sum_j A_{ij} u_j \right) |a_i\rangle .$$

(12.15)

Hierbei sind $v_i = \langle a_i | v \rangle$ und $u_j = \langle a_j | u \rangle$. Als Analogie im $\mathbb{R}^3$ ergibt (12.15) z. B. für die Komponente mit $i = 2$

$$v_2 = A_{21}u_1 + A_{22}u_2 + A_{23}u_3 \quad \neq \quad v_2 \cdot \vec{e}_2 = (A_{21}u_1 + A_{22}u_2 + A_{23}u_3) \cdot \vec{e}_2 .$$

Im Gegensatz zur Matrixgleichung (12.10) beinhaltet die Operatorgleichung (12.12) nämlich den zum Entwicklungskoeffizienten bzw. zur skalaren Vektorkomponente v_i gehörenden Basisvektor $|a_i\rangle$.

13 Eigenwertgleichung einer 2-reihigen reellen Matrix – Verallgemeinerungen für n -reihige Matrizen

Siehe auch:

Lothar Papula, *Mathematik für Ingenieure und Naturwissenschaftler Band 2*, 10. Auflage, Vieweg, Braunschweig, Wiesbaden, 2001, Seite 121 bis Seite 140

und

Christian B. Lang und Norbert Pucker, *Mathematische Methoden in der Physik*, Spektrum, Heidelberg, Berlin, 1998, Seite 204 bis 211, Seite 414 bis 422.

Wir gehen aus von der linearen Abbildungsgleichung (Transformationsgleichung)

$$A\vec{x} = \vec{y}$$

im $\mathbb{R}^2$ bzw. in der Ebene. Dabei seien die Funktion A eine 2-reihige (quadratische) reelle Matrix und $\vec{x} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$ und $\vec{y} = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}$ die Ortsvektoren zur Beschreibung der Punkte $P_{\vec{x}} := \{x_1, x_2\}$ und $P_{\vec{y}} := \{y_1, y_2\}$. Die Wirkung von A auf den Ortsvektor $\vec{x}$ kann allgemein sowohl in einer Drehung als auch in einer Längenänderung bestehen, woraus dann der Ortsvektor $\vec{y}$ resultiert (s. Abb. 13.1 a). Beispielsweise liefern

$$A = \begin{pmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{pmatrix} = \begin{pmatrix} 2 & -1 \\ -3 & 4 \end{pmatrix}, \quad \vec{x} = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$$

$$A\vec{x} = \begin{pmatrix} 2 & -1 \\ -3 & 4 \end{pmatrix} \begin{pmatrix} 2 \\ 1 \end{pmatrix} = \begin{pmatrix} 3 \\ -2 \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = \vec{y}.$$

Wir gehen jetzt der Frage nach, ob zur Matrix A Vektoren $\vec{x}$ existieren, sodass $\vec{y}$ kollinear zu $\vec{x}$ verläuft (s. Abb. 13.1 b).

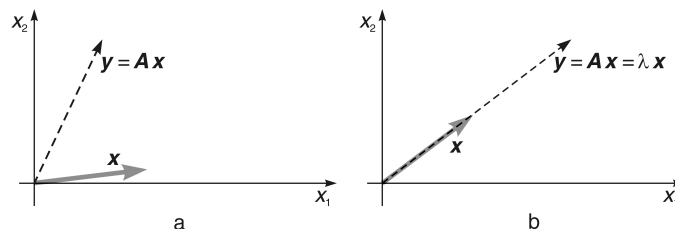


Abb. 13.1

a Die Matrix A bewirkt sowohl eine Längenänderung als auch eine Drehung am Vektor $\vec{x}$ und überführt diesen so in den Vektor $\vec{y}$. Der Vektor $\vec{x}$ ist folglich kein Eigenvektor zu A .

b Die Matrix A überführt den Vektor $\vec{x}$ (ohne Drehung) in den zu $\vec{x}$ kollinearen Vektor $\vec{y}$. Der Vektor $\vec{x}$ ist somit ein Eigenvektor zur Matrix A .

In diesem Fall resultiert mit der Einheitsmatrix $E = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$

$$A\vec{x} = \lambda\vec{x} = \lambda \cdot E\vec{x} = \vec{y}, \quad \lambda \in \mathbb{R}. \quad (13.1)$$

Um $\vec{x}$ zu bestimmen bilden wir aus (13.1) durch die Äquivalenzumformung

$$A \cdot \vec{x} = \lambda E \cdot \vec{x} \Leftrightarrow (A - \lambda E) \cdot \vec{x} = 0$$

das homogene lineare Gleichungssystem (LGS)

$$\left[\begin{pmatrix} 2 & -1 \\ -3 & 4 \end{pmatrix} - \begin{pmatrix} \lambda & 0 \\ 0 & \lambda \end{pmatrix} \right] \cdot \vec{x} = \begin{pmatrix} 2-\lambda & -1 \\ -3 & 4-\lambda \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} \stackrel{!}{=} \begin{pmatrix} 0 \\ 0 \end{pmatrix}. \quad (13.2)$$

Ein homogenes LGS besitzt nur dann nicht-triviale Lösungen $\vec{x}$, wenn die Koeffizientendeterminante $\det(A - \lambda E)$ verschwindet¹:

$$\begin{aligned}\det(A - \lambda E) &= \det \begin{pmatrix} 2 - \lambda & -1 \\ -3 & 4 - \lambda \end{pmatrix} = (2 - \lambda)(4 - \lambda) - [(-1)(-3)] \\ &= \lambda^2 - 6\lambda + 5 \stackrel{!}{=} 0\end{aligned}$$

$A - \lambda E$ ist die *charakteristische Matrix* von A und $\lambda^2 - 6\lambda + 5 \stackrel{!}{=} 0$ heißt die *charakteristische Gleichung* bzw. das *charakteristische Polynom* von A . Es ist in diesem Fall eine quadratische Gleichung in λ und somit nur erfüllt für

$$\lambda \Rightarrow \begin{cases} 3 + \sqrt{9 - 5} = \lambda_1 = 5, \\ 3 - \sqrt{9 - 5} = \lambda_2 = 1. \end{cases}$$

Die Lösungen λ_1 und λ_2 dieser quadratischen Gleichung, d.h. die Nullstellen des charakteristischen Polynoms von A sind die *Eigenwerte der Matrix* A .

Allgemein können wir für die 2-reihige Determinante der Matrix A

$$\det A = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11}a_{22} - a_{12}a_{21}$$

und für das charakteristische Polynom der Matrix A

$$\begin{aligned}\det(A - \lambda E) &= \begin{vmatrix} a_{11} - \lambda & a_{12} \\ a_{21} & a_{22} - \lambda \end{vmatrix} = (a_{11} - \lambda)(a_{22} - \lambda) - a_{12}a_{21} \\ &= \lambda^2 - \underbrace{(a_{11} + a_{22})}_{\text{Sp}(A)} \lambda + \underbrace{(a_{11}a_{22} - a_{12}a_{21})}_{\det A}\end{aligned}$$

$$\det(A - \lambda E) = \lambda^2 - \text{Sp}(A) \cdot \lambda + \det A = 0 \quad (13.3)$$

schreiben. Durch Faktorisierung (Zerlegung in Linearfaktoren) erhält das charakteristische Polynom die Produktform

$$\det(A - \lambda E) = (\lambda - \lambda_1)(\lambda - \lambda_2) = \lambda^2 - (\lambda_1 + \lambda_2) \cdot \lambda + \lambda_1 \lambda_2. \quad (13.4)$$

Der Koeffizientenvergleich zwischen (13.3) und (13.4) zeigt, wie Spur und Determinante von A mit den Eigenwerten von A zusammenhängen:

$$\begin{aligned}\text{Sp}(A) &= a_{11} + a_{22} = \lambda_1 + \lambda_2 \\ \det A &= a_{11}a_{22} - a_{12}a_{21} = \lambda_1 \lambda_2.\end{aligned}$$

¹Die Determinante einer Matrix M verschwindet, wenn M singulär (linear abhängig) ist. Folglich ist es erforderlich, dass die charakteristische Matrix $A - \lambda E$ singulär ist.

Mit den Eigenwerten λ_1 und λ_2 können wir jetzt das homogene LGS (13.2) lösen.

Für $\lambda_1 = 5$ resultiert

$$\begin{pmatrix} 2 - \lambda_1 & -1 \\ -3 & 4 - \lambda_1 \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \Rightarrow$$
$$\left. \begin{array}{l} \text{(I)} \quad -3x_1 - x_2 = 0 \\ \text{(II)} \quad -3x_1 - x_2 = 0 \end{array} \right\} \Leftrightarrow x_2 = -3x_1.$$

x_1 ist frei wählbar. Wir setzen $x_1 = \alpha$ und erhalten für den normierten Lösungsvektor $\vec{x}_1$ zum Eigenwert λ_1

$$\vec{x}_1(\lambda_1) = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \alpha \begin{pmatrix} 1 \\ -3 \end{pmatrix} = \frac{1}{\sqrt{10}} \begin{pmatrix} 1 \\ -3 \end{pmatrix}.$$

Für $\lambda_2 = 1$ resultiert

$$\begin{pmatrix} 2 - \lambda_2 & -1 \\ -3 & 4 - \lambda_2 \end{pmatrix} \cdot \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \Rightarrow$$
$$\left. \begin{array}{l} \text{(I)} \quad x_1 - x_2 = 0 \\ \text{(II)} \quad -3x_1 + 3x_2 = 0 \end{array} \right\} \Leftrightarrow x_2 = x_1.$$

x_1 ist frei wählbar. Wir setzen $x_1 = \beta$ und erhalten für den normierten Lösungsvektor $\vec{x}_2$ zum Eigenwert λ_2

$$\vec{x}_2(\lambda_2) = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \beta \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix}.$$

Die Lösungsvektoren $\vec{x}_1$ und $\vec{x}_2$ sind die *Eigenvektoren bezüglich der Matrix A*. Die Eigenvektoren liegen im Kern der charakteristischen Matrix $A - \lambda E$.

Weil α und β frei gewählt werden durften, hätten wir sie auch gleich 1 setzen können. Es ist zwar Konvention, die Eigenvektoren zu normieren, aber nicht unbedingt notwendig. Deshalb führen wir die Probe mit

$$\alpha = 1 \Rightarrow \tilde{\vec{x}}_1 = \begin{pmatrix} 1 \\ -3 \end{pmatrix} \quad \text{und} \quad \beta = 1 \Rightarrow \tilde{\vec{x}}_2 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$$

durch:

$$A \cdot \tilde{\vec{x}}_1 = \begin{pmatrix} 2 & -1 \\ -3 & 4 \end{pmatrix} \begin{pmatrix} 1 \\ -3 \end{pmatrix} = \begin{pmatrix} 5 \\ -15 \end{pmatrix} = 5 \cdot \begin{pmatrix} 1 \\ -3 \end{pmatrix} = \lambda \cdot \tilde{\vec{x}}_1, \quad \lambda = \lambda_1 = 5 \in \mathbb{R},$$
$$A \cdot \tilde{\vec{x}}_2 = \begin{pmatrix} 2 & -1 \\ -3 & 4 \end{pmatrix} \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \begin{pmatrix} 1 \\ 1 \end{pmatrix} = 1 \cdot \begin{pmatrix} 1 \\ 1 \end{pmatrix} = \lambda \cdot \tilde{\vec{x}}_2, \quad \lambda = \lambda_2 = 1 \in \mathbb{R}.$$

Wirkt eine Matrix A auf einen ihrer Eigenvektoren $\vec{x}_i$, so erzeugt sie einen zu $\vec{x}_i$ kollinearen Vektor $\vec{y}_i$, dessen Länge das λ_i -fache der Länge von $\vec{x}_i$ beträgt. Hierbei ist λ_i der zum Eigenvektor $\vec{x}_i$ gehörende Eigenwert der Matrix A . Die Länge der Eigenvektoren $\vec{x}_i$ kann frei gewählt werden. Konventionsgemäß werden die Eigenvektoren aber auf die Länge 1 normiert.

Merke!

Eine Eigenwertgleichung beschreibt die Wirkung einer Matrix A auf einen ihrer Eigenvektoren $\vec{v}$. Diese Wirkung ist die Multiplikation des Eigenvektors mit dem zugehörigen Eigenwert λ . A produziert folglich den Vektor $\lambda\vec{v}$ der parallel (kollinear) zum Eigenvektor $\vec{v}$ ist.

Verallgemeinerungen für n -reihige Matrizen

Die für 2-reihige Matrizen gewonnenen Erkenntnisse lassen sich auf n -reihige Matrizen verallgemeinern.

- Insbesondere gilt für die Eigenwerte n -reihiger Matrizen A

$$\lambda_1 + \lambda_2 + \lambda_3 + \cdots + \lambda_n = \text{Sp}(A) , \quad \lambda_1 \lambda_2 \lambda_3 \cdots \lambda_n = \det A .$$

- Zu n *verschiedenen* Eigenwerten gehören n linear unabhängige Eigenvektoren.
- Ein k -facher Eigenwert (Vielfachheit k) hat mindestens einen und höchstens k linear unabhängige Eigenvektoren.
- Die Eigenwerte einer Dreiecks- oder Diagonalmatrix A sind identisch mit ihren Hauptdiagonalelementen. Dies resultiert aus den Rechenregeln für Determinanten beispielsweise wie folgt:

$$\begin{aligned} \det(A - \lambda E) &= \det \begin{pmatrix} a_{11} - \lambda_1 & 0 & 0 \\ a_{21} & a_{22} - \lambda_2 & 0 \\ a_{31} & a_{32} & a_{33} - \lambda_3 \end{pmatrix} \\ &= (a_{11} - \lambda_1)(a_{22} - \lambda_2)(a_{33} - \lambda_3) \stackrel{!}{=} 0 \\ \Rightarrow \quad \lambda_1 &= a_{11} , \quad \lambda_2 = a_{22} , \quad \lambda_3 = a_{33} . \end{aligned}$$

- Die Eigenvektoren bezüglich einer Diagonalmatrix sind die den Eigenwerten λ_i entsprechenden kanonischen Einheitsvektoren $|e_i\rangle$ der Standardbasis $\{|e_i\rangle\}$.
Beispiel zur Veranschaulichung mit dem Eigen- bzw. Basisvektor $|e_2\rangle = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$:

$$A|e_i\rangle = \lambda_i|e_i\rangle \quad \Rightarrow \quad \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \begin{pmatrix} 0 \\ \lambda_2 \end{pmatrix} = \lambda_2 \cdot \begin{pmatrix} 0 \\ 1 \end{pmatrix} . \quad \square$$

- Zusätzlich gilt für *reelle* symmetrische Matrizen $A = A^T$:
Alle n Eigenwerte sind reell und die Eigenvektoren von verschiedenen Eigenwerten sind orthogonal.²

²Die Beweise finden sich unter
www.mpi-inf.mpg.de/departments/d1/teaching/ss10/MFI2/kap46.pdf
oder unter dem Suchbegriff „Eigenwerte und Eigenvektoren symmetrischer Matrizen“.

14 Hilbertraum, Hermitesche Matrizen (Operatoren) und ihre Eigenschaften

Die Feststellungen über hermitesche Matrizen gelten auch für hermitesche Operatoren. Abhängig vom Zusammenhang werden wir deshalb im Folgenden entweder von Matrizen oder von Operatoren sprechen. Weil die hermiteschen Operatoren in der Quantenmechanik auf Zustandsvektoren, d. h. auf Elemente des Hilbert-Raums $\mathcal{H}$ wirken, beginnen wir diesen Abschnitt mit der

Definition des Hilbert-Raums:

Im Allgemeinen sind Hilberträume $\mathcal{H}$ *abzählbar-unendlichdimensionale, lineare Vektorräume über $\mathbb{C}$ mit einem Standardskalarprodukt*. Weiterhin sind Hilberträume *vollständig* bezüglich der durch $\|\Psi\| = \sqrt{\langle \Psi | \Psi \rangle}$ definierten *Norm*, wobei $|\Psi\rangle$ abstrakte, basisfreie Hilbert-Raum-Vektoren (Hilbert-Raum-Elemente) sind.

Die Hilbert-Raum-Vektoren $|\Psi\rangle$ lassen sich nach ihren quadratsummablen Koordinaten entwickeln, die wiederum eine den Hilbertraum vollständig aufspannende Orthonormalbasis bilden. Der Hilbert-Raum zu einem Operator mit einem vollständigen Satz von n linear unabhängigen Eigenvektoren ist ebenfalls n -dimensional.

Hilbert-Räume sind z. B. der $\mathbb{R}^n$, der $\mathbb{C}^n$, der quadratsummable Folgenraum ℓ^2 , aber auch der L^2 -Raum der quadratintegrierbaren Funktionen.

Hermitesche Matrizen (Operatoren) und ihre Eigenschaften

- Hermitesche Matrizen A sind quadratisch und komplex. Sie lassen sich zerlegen in einen Realteil, die symmetrische reelle Matrix B , und in einen Imaginärteil, die antisymmetrische (schiefsymmetrische) Matrix C . Für eine n -reihige hermitesche Matrix

$$A = A_{n \times n} = B_{n \times n} + i C_{n \times n} \quad \Rightarrow \quad A = B + i C$$

gilt folglich

$$A^\dagger = (B + i C)^\dagger = B^\dagger + C^\dagger i^\dagger = B^T + C^T i^* = B - C(-i) = B + i C = A ,$$

$$A = A^\dagger .$$

Alle Hauptdiagonalelemente sind reell. Eine hermitesche Matrix ist z. B.

$$\begin{aligned} A &= \begin{pmatrix} -7 & 3i & 1-4i \\ -3i & 2 & 8-i \\ 1+4i & 8+i & 6 \end{pmatrix} \\ &= \begin{pmatrix} -7 & 0 & 1 \\ 0 & 2 & 8 \\ 1 & 8 & 6 \end{pmatrix} + \begin{pmatrix} 0 & 3i & -4i \\ -3i & 0 & -i \\ 4i & i & 0 \end{pmatrix} . \end{aligned}$$

- Die Determinante einer hermiteschen Matrix ist reell:

$$\det A \in \mathbb{R} .$$

- Hermitesche Matrizen (**lineare hermitesche Operatoren**) besitzen nur **reelle Eigenwerte**:¹

In (12.4) verwenden wir für $\langle u|$

$$|u\rangle = A|v\rangle \quad \Rightarrow \quad \langle u| = |u\rangle^\dagger = (A|v\rangle)^\dagger = \langle v|A^\dagger = \langle v|A$$

und erhalten

$$\langle u|v\rangle = \langle u|v\rangle^* = \langle v|A|v\rangle = (\langle v|A|v\rangle)^*.$$

Mit dem Eigenwert λ von A gilt die Eigenwertgleichung

$$A|v\rangle = \lambda|v\rangle,$$

woraus

$$\begin{aligned} \langle v|A|v\rangle &= \langle v|\lambda|v\rangle \\ &= \lambda\langle v|v\rangle = (\lambda\langle v|v\rangle)^* \end{aligned}$$

folgt. Wegen $\langle v|v\rangle \in \mathbb{R}$ muss also auch $\lambda \in \mathbb{R}$ gelten.

- Die **Eigenvektoren** zu *verschiedenen* Eigenwerten einer hermiteschen Matrix (einem linearen hermiteschen Operator) sind **orthogonal**:¹

„Sei

$$A|u\rangle = \lambda_1|u\rangle, \quad A|v\rangle = \lambda_2|v\rangle, \quad \lambda_1 \neq \lambda_2.$$

Aus der zweiten Gleichung folgt $\langle v|A = \lambda_2\langle v|$. Bildet man in dieser Gleichung das Skalarprodukt mit $|u\rangle$, in der ersten Gleichung das Skalarprodukt mit $\langle v|$ und subtrahiert anschließend die beiden Gleichungen voneinander, so erhält man

$$\langle v|A|u\rangle - \langle v|A|u\rangle = \lambda_1\langle v|u\rangle - \lambda_2\langle v|u\rangle,$$

also

$$(\lambda_1 - \lambda_2)\langle v|u\rangle = 0,$$

also $\langle v|u\rangle = 0$.“

- Die **Eigenvektoren** bezüglich einer hermiteschen Matrix (einem linearen hermiteschen Operator) spannen eine **vollständige Orthonormalbasis** (VON-Basis) in $\mathcal{H}$ auf.² Kurz gesagt, die Eigenvektoren aller Eigenwerte einer hermiteschen Matrix (eines hermiteschen Operators) spannen den gesamten zugehörigen Hilbert-Raum auf.
- Der Eigenvektor bzw. „die Eigenvektoren zu einem bestimmten Eigenwert λ bilden einen Vektorraum, einen Unterraum von $\mathcal{H}$, den Eigenraum $\mathcal{H}_\lambda$ zum Eigenwert λ .“² Insofern stellt jeder einzelne Eigenvektor einen eindimensionalen Unterraum von $\mathcal{H}$ dar.

¹Zitiert aus bzw. nach: Jan-Markus Schwindt, *Tutorium Quantenmechanik*, Springer Spektrum, Berlin, Heidelberg, 2013, Seite 27 bis Seite 30.

²Zitiert aus bzw. nach: Jan-Markus Schwindt, *Tutorium Quantenmechanik*, Springer Spektrum, Berlin, Heidelberg, 2013, Seite 27 bis Seite 30.

- Jeder durch eine quantenmechanische Messung zu beschreibende Zustand eines Objekts ist mathematisch gesehen ein **Strahl** bzw. die Vektorenmenge $\{\alpha|\Psi\rangle\}$ mit $\alpha \in \mathbb{C}$ und folglich ein eindimensionaler Unterraum von $\mathcal{H}$. $|\Psi\rangle$ ist dabei der im Allgemeinen gemäß $\langle\Psi|\Psi\rangle = 1$ normierte Zustandsvektor und kann als Repräsentant des zu beschreibenden Zustands gewählt werden. Die Zustandsvektoren (Zustände) $|\Psi\rangle$ wiederum sind im Allgemeinen Linearkombinationen aus Eigenvektoren $|a_i\rangle$ bezüglich einer Observablen, d. h. bezüglich eines Operators $\hat{A}$.

15 Unitäre Matrizen (Operatoren)

Unitäre Operatoren bezeichnen wir mit $\hat{U}$ und die zugehörigen unitären Matrizen (Operatorenmatrizen) mit $(U_{ij}) = U$. Zwei sehr einfache Beispiele für unitäre Matrizen sind die Einheitsmatrix und $U = \begin{pmatrix} 0 & i \\ i & 0 \end{pmatrix}$. Vereinfachend verwenden wir in diesem Abschnitt für die (Zustands-)Vektoren gelegentlich die Darstellung

$$\begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_i \\ \vdots \end{pmatrix} =: |v\rangle .$$

15.1 Eigenschaften unitärer Matrizen (Operatoren)

- Unitär heißt im weitesten Sinne normerhaltend. Das aber bedeutet, dass $\hat{U}$ bzw. U auch das Skalarprodukt invariant lassen:

$$\begin{aligned} \langle \Phi | \Psi \rangle &\stackrel{!}{=} \langle \hat{U} \Phi | \hat{U} \Psi \rangle \\ &= \langle \Phi | \underbrace{\hat{U}^\dagger \hat{U}}_{\mathbb{1}} | \Psi \rangle \\ &= \langle \Phi | U^\dagger U | \Psi \rangle = \langle \Phi | \Psi \rangle . \end{aligned}$$

- Daraus folgt

$$\hat{U}^\dagger \hat{U} = U^\dagger U = \mathbb{1} \quad \Leftrightarrow \quad U^\dagger = U^{-1} , \quad \hat{U}^\dagger = \hat{U}^{-1} .$$

- Damit ist U stets invertierbar und regulär.
- Inverse und Produkt unitärer Matrizen (Operatoren) sind unitär.
- Der Betrag der Determinante von U ist stets

$$|\det U| = 1 .$$

- Unitär im Komplexen entspricht orthogonal im Reellen, denn Zeilen und Spalten unitärer Matrizen (Operatoren) sind orthonormiert. Mit der unitären Matrix $U \in \mathbb{C}^{n \times n}$ und der orthogonalen Matrix $O \in \mathbb{R}^{n \times n}$ heißt das

$$U^\dagger U = U^{-1} U = \mathbb{1} \quad \longleftrightarrow \quad O^T O = O^{-1} O = \mathbb{1} .$$

- $|u_i\rangle$ seien die Eigenzustände bezüglich $\hat{U}$ und $\lambda_i \in \mathbb{C}$ die zugehörigen Eigenwerte. Aus der entsprechenden Eigenwertgleichung

$$\hat{U}|u_i\rangle = \lambda_i|u_i\rangle$$

erhalten wir das Skalarprodukt (gleicher Zeilen oder gleicher Spalten)

$$\begin{aligned} \langle \hat{U} u_i | \hat{U} u_i \rangle &= \langle u_i | U^\dagger U | u_i \rangle = \langle u_i | u_i \rangle = 1 \\ &= \langle \lambda_i u_i | \lambda_i u_i \rangle = \lambda_i^* \lambda_i \langle u_i | u_i \rangle = \langle u_i | u_i \rangle = 1 \\ \Rightarrow \quad \lambda_i^* \lambda_i &= |\lambda_i|^2 = 1 . \end{aligned}$$

Das Betragsquadrat jedes Eigenwerts λ_i des unitären Operators $\hat{U}$ bzw. der unitären Matrix U ist gleich 1.

15.2 Unitäre Transformation

Nach: Christian B. Lang und Norbert Pucker, *Mathematische Methoden in der Physik*, Hochschultaschenbuch, Spektrum Akademischer Verlag, Heidelberg, Berlin, 1998, Seite 430 und Seite 431.

Bei der unitären Transformation wird durch eine unitäre Matrix (einen unitären Operator) ein normerhaltender Darstellungswechsel bewirkt. Die unitäre Transformation ändert also das Basissystem im Vektorraum z. B. im Sinne einer Drehung oder Spiegelung. Die Orthogonalitätseigenschaften des Basissystems bleiben dabei erhalten.

Analog zu einer Drehmatrix vermittelt so ein unitärer Operator $\hat{U}$ die unitäre Transformation zwischen den zwei verschiedenen Orthonormalbasen $\{|a_i\rangle\}$ und $\{|b_i\rangle\}$ gemäß

$$\hat{U}|a_i\rangle = |b_i\rangle . \quad (15.1)$$

Die Matrixelemente des unitären Operators $\hat{U}$ (s. Matrixdarstellung von Operatoren) sind damit

$$\langle a_i|\hat{U}|a_j\rangle = \langle a_i|b_j\rangle = U_{ij} .$$

Die Umkehroperation von (15.1) ist

$$\hat{U}^\dagger \hat{U}|a_i\rangle = U^\dagger U|a_i\rangle = \hat{U}^\dagger |b_i\rangle = |a_i\rangle .$$

Ein Vektor $|v\rangle$ hat in den beiden verschiedenen Orthonormalbasen $\{|a_i\rangle\}$ und $\{|b_i\rangle\}$ die verschiedenen Komponentendarstellungen

$$\sum_i \alpha_i |a_i\rangle = |v\rangle \equiv |\widetilde{v}\rangle = \sum_i \beta_i |b_i\rangle .$$

Der Darstellungswechsel des Vektors $|v\rangle$ von der einen in die andere Orthonormalbasis mittels unitärer Transformation erfolgt in Komponentenschreibweise wie folgt:

$$\begin{aligned} \beta_j &= \langle b_j|v\rangle = \langle b_j|\sum_i \alpha_i |a_i\rangle \\ &= \sum_i \alpha_i \langle b_j|a_i\rangle \\ \beta_j &= \sum_i U_{ji}^* \alpha_i . \end{aligned} \quad (15.2)$$

Für die zu den U_{ij} adjungierten Matrixelementen haben wir dabei U_{ji}^* geschrieben, weil für die zu (U_{ij}) adjungierte Matrix $(U_{ij})^\dagger = (U_{ji})^* = (U_{ji}^*)$ gilt. Die Matrixelemente U_{ji}^* erhalten wir aus den Matrixelementen U_{ij} wie folgt:

$$\begin{aligned} U_{ij} &= \langle a_i|b_j\rangle && \text{mit Zeilenindex } i \text{ und Spaltenindex } j \\ \Rightarrow & \langle a_i|b_j\rangle^\dagger = \langle a_i|b_j\rangle^* = \langle b_j|a_i\rangle \\ \Rightarrow & \langle b_j|a_i\rangle = U_{ji}^* && \text{mit Zeilenindex } j \text{ und Spaltenindex } i . \end{aligned}$$

In Matrixschreibweise ist (15.2)

$$\begin{pmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_j \\ \vdots \end{pmatrix} =: |\widetilde{v}\rangle = \hat{U}^\dagger |v\rangle := U^\dagger \begin{pmatrix} \alpha_1 \\ \alpha_2 \\ \vdots \\ \alpha_i \\ \vdots \end{pmatrix} .$$

Die Darstellung einer Matrix bzw. eines Operators ändert sich ebenfalls bei einem Basiswechsel. Betrachten wir also die unitäre Transformation des Matrixelements A_{ij} eines Operators $\hat{A}$. Die Transformation erfolge von der Darstellung in der Basis $\{|b_i\rangle\}$ in die Darstellung in der Basis $\{|a_i\rangle\}$. Mit (15.1) und den Komponenten von $\hat{U}$ und $\hat{A}$ erhalten wir in Komponentenschreibweise

$$\begin{aligned} A_{ij} &= \langle b_i | \hat{A} | b_j \rangle = \langle \hat{U} a_i | \hat{A} | \hat{U} a_j \rangle = \langle a_i | \hat{U}^\dagger \hat{A} \hat{U} | a_j \rangle = (U^\dagger A U)_{ij} = \tilde{A}_{ij} \\ &= \sum_{k,l} (U^\dagger)_{ik} A_{kl} U_{lj} . \end{aligned} \quad (15.3)$$

Hierbei ist also A_{ij} das Matrixelement der Matrix A in der Basis $\{|b_i\rangle\}$ und $\tilde{A}_{ij}$ das Matrixelement der Matrix $\tilde{A}$ in der Basis $\{|a_i\rangle\}$. $U_{lj} = \langle a_l | b_j \rangle$, d. h. $|b_j\rangle$ in $|a_l\rangle$, ist die Projektion von $|b_j\rangle$ auf $|a_l\rangle$. In Matrixschreibweise ist schließlich (15.3) kurz

$$U^\dagger A U = \tilde{A} .$$

Dies entspricht dem basisfreien Operatorenprodukt

$$\hat{U}^\dagger \hat{A} \hat{U} = \tilde{\hat{A}} .$$

Die unitäre Transformation von Polynomen $(A + B)$ erfolgt durch

$$\widetilde{(A + B)} = U^\dagger (A + B) U = U^\dagger (A U + B U) = U^\dagger A U + U^\dagger B U = \tilde{A} + \tilde{B}$$

und von Potenzreihen $(A \cdot B)$ durch

$$\widetilde{(A \cdot B)} = U^\dagger (A \cdot B) U = U^\dagger (A \cdot U U^\dagger \cdot B) U = U^\dagger A U \cdot U^\dagger B U = \tilde{A} \cdot \tilde{B} .$$

Merke!

Die unitäre Transformation von der Basis $\{|a_i\rangle\}$ in die Basis $\{|b_i\rangle\}$ erfolgt durch

$$\hat{U}|a_i\rangle = |b_i\rangle$$

mit der Umkehroperation $|a_i\rangle = \hat{U}^\dagger|b_i\rangle$.

Die unitäre Transformation eines Vektors $|v\rangle$ erfolgt durch

$$\hat{U}^\dagger|v\rangle = |\widetilde{v}\rangle$$

mit der Umkehroperation $|v\rangle = U|\widetilde{v}\rangle$.

Die unitäre Transformation einer Matrix A bzw. eines Operators $\hat{A}$ erfolgt durch

$$U^\dagger A U = \widetilde{A} \quad \longleftrightarrow \quad \hat{U}^\dagger \hat{A} \hat{U} = \widetilde{\hat{A}}$$

mit der Umkehroperation $A = U \widetilde{A} U^\dagger \quad \longleftrightarrow \quad \hat{A} = \hat{U} \widetilde{\hat{A}} \hat{U}^\dagger$.

Dabei ist $(A_{ij}) = A$ die zum Operator $\hat{A}$ gehörende und $(\widetilde{A}_{ij}) = \widetilde{A}$ die zum Operator $\widetilde{\hat{A}}$ gehörende Matrix. $U(\dots)U^\dagger$ ist die zu $U^\dagger(\dots)U$ inverse Operation.

15.3 Diagonalisierung von Matrizen (Operatoren)

Nach: Christian B. Lang und Norbert Pucker, *Mathematische Methoden in der Physik*, Hochschultaschenbuch, Spektrum Akademischer Verlag, Heidelberg, Berlin, 1998, Seite 418 und Seite 419.

In diesem Abschnitt folgen wir im Wesentlichen der Argumentation von Lang und Pucker (s. Literaturhinweis). Wie wir in unseren bisherigen Betrachtungen über den Umgang mit Matrizen und Operatoren unter Berücksichtigung der Matrixdarstellung von Operatoren feststellen konnten, lassen sich die Ergebnisse für Matrizen mühelos auf Operatoren übertragen. Deshalb gehen wir jetzt vereinfachend nur von einer Matrix A aus.

Diese Matrix A habe die spezielle Eigenschaft, dass ihre Eigenvektoren $|v_i\rangle$ orthogonal aufeinander stehen und gleichzeitig auch eine zu A passende VON-Basis $\{|v_i\rangle\}$ bilden. Dies gilt stets für *hermitesche* und folglich auch für *symmetrische reelle* Matrizen. Selbstverständlich erfüllen diese Basisvektoren (Eigenvektoren) $|v_i\rangle$ die Eigenwertgleichung

$$A|v_i\rangle = \lambda_i|v_i\rangle.$$

Wenn wir aus den orthonormierten Eigen- bzw. Basisvektoren $|v_i\rangle$ eine Matrix bilden derart, dass die Basisvektoren die Spalten der Matrix darstellen, so erhalten wir die *unitäre* Matrix U sinngemäß durch

$$U := (|v_1\rangle \quad |v_2\rangle \quad \cdots \quad |v_i\rangle \quad \cdots),$$

denn wegen der Orthonormalität der $|v_i\rangle$ gemäß $\langle v_i|v_j\rangle = \delta_{ij}$ gilt

$$U^\dagger U = \begin{pmatrix} (v_1)_1^* & (v_1)_2^* & \cdots & (v_1)_j^* & \cdots \\ (v_2)_1^* & (v_2)_2^* & \cdots & (v_2)_j^* & \cdots \\ \vdots & \vdots & \ddots & \vdots & \\ (v_i)_1^* & (v_i)_2^* & \cdots & (v_i)_j^* & \cdots \\ \vdots & \vdots & & \vdots & \ddots \end{pmatrix} \begin{pmatrix} (v_1)_1 & (v_2)_1 & \cdots & (v_i)_1 & \cdots \\ (v_1)_2 & (v_2)_2 & \cdots & (v_i)_2 & \cdots \\ \vdots & \vdots & \ddots & \vdots & \\ (v_1)_j & (v_2)_j & \cdots & (v_i)_j & \cdots \\ \vdots & \vdots & & \vdots & \ddots \end{pmatrix} = \mathbb{1} \quad (15.4)$$

Wie man sieht, sind die Elemente der Matrix U die Entwicklungskoeffizienten $(v_i)_j$ der Eigenvektoren $|v_i\rangle$ in einer anderen VON-Basis, z. B. $\{|a_j\rangle\}$:

$$|v_i\rangle = \sum_j \langle a_j | v_i \rangle |a_j\rangle = \sum_j (v_i)_j |a_j\rangle := \begin{pmatrix} (v_i)_1 \\ (v_i)_2 \\ \vdots \\ (v_i)_j \\ \vdots \end{pmatrix}.$$

Wären die $|a_j\rangle$ selbst die Eigenvektoren zur Matrix A , hätte U die Gestalt einer Diagonalmatrix.

Jetzt multiplizieren wir U in (15.4) mit der Matrix A .¹ Entsprechend der Eigenwertgleichung

$$A|v_i\rangle = \lambda_i |v_i\rangle := \lambda_i \begin{pmatrix} (v_i)_1 \\ (v_i)_2 \\ \vdots \\ (v_i)_j \\ \vdots \end{pmatrix} = \begin{pmatrix} \lambda_i (v_i)_1 \\ \lambda_i (v_i)_2 \\ \vdots \\ \lambda_i (v_i)_j \\ \vdots \end{pmatrix}$$

bezüglich der Eigenvektoren (Spalten) $|v_i\rangle$ von U gilt für die Elemente der resultierenden Matrix $U^\dagger A U$, d. h. für „Zeile k von $U^\dagger$ mal λ_i mal Spalte i von U “

$$\sum_j (v_k)_j^* \lambda_i (v_i)_j = \lambda_i \sum_j (v_k)_j^* (v_i)_j = \lambda_i \delta_{ki},$$

sodass

$$U^\dagger A U = \begin{pmatrix} \lambda_1 & & & 0 \\ & \lambda_2 & & \\ & & \ddots & \\ 0 & & & \lambda_i & \\ & & & & \ddots \end{pmatrix} = \Lambda.$$

Die Hauptdiagonalelemente der Diagonalmatrix Λ sind die Eigenwerte der Matrix A .

¹Die Multiplikation der Matrix U mit der Matrix A entspricht der Wirkung des Operators $\hat{A}$ auf den unitären Operator $\hat{U}$.

Damit haben wir ein Verfahren zur Diagonalisierung von hermiteschen bzw. symmetrischen reellen Matrizen A (Operatoren $\hat{A}$) gefunden:

1. Wir bilden die orthonormierten Eigenvektoren $|v_i\rangle$ zu A (zu $\hat{A}$).
2. Diese Eigenvektoren bilden die Spalten der unitären Matrix U . Auf die Reihenfolge der Eigenvektoren bzw. Spalten ist zu achten. Sie ist nicht beliebig, denn eine Vertauschung der Reihenfolge der Eigenvektoren $|v_i\rangle$ in U bewirkt eine entsprechende Vertauschung der Reihenfolge der Diagonalelemente λ_i in Λ .
3. Das Matrixprodukt $U^\dagger A U$ bewirkt dann die Diagonalisierung von A und liefert die Diagonalmatrix Λ , deren Hauptdiagonalelemente die Eigenwerte λ_i von A (von $\hat{A}$) sind.

Beispiel

Betrachten wir die hermitesche Matrix

$$A = \begin{pmatrix} 8 & -3i \\ 3i & 0 \end{pmatrix}.$$

Ihre Eigenwerte sind

$$\begin{aligned} \lambda_1 &= 4 + \sqrt{16 + 9} = 9, \\ \lambda_2 &= 4 - \sqrt{16 + 9} = -1, \end{aligned}$$

also erwartungsgemäß reell. Die (normierten) Eigenvektoren zu A sind

$$|v_1\rangle_{(\lambda_1)} := \frac{1}{\sqrt{10}} \begin{pmatrix} 3 \\ i \end{pmatrix},$$

$$|v_2\rangle_{(\lambda_2)} := \frac{1}{\sqrt{10}} \begin{pmatrix} i \\ 3 \end{pmatrix}.$$

$|v_1\rangle$ und $|v_2\rangle$ sind tatsächlich orthogonal, denn

$$\langle v_1 | v_2 \rangle = \frac{1}{\sqrt{10}} \begin{pmatrix} 3 \\ i \end{pmatrix}^\dagger \cdot \frac{1}{\sqrt{10}} \begin{pmatrix} i \\ 3 \end{pmatrix} = \frac{1}{10} (-3i + 3i) = 0.$$

Die Eigenvektoren sind die Spalten der unitären Matrix U , die A wie folgt diagonalisiert:

$$U^\dagger A U = \begin{pmatrix} \frac{3}{\sqrt{10}} & \frac{-i}{\sqrt{10}} \\ \frac{-i}{\sqrt{10}} & \frac{3}{\sqrt{10}} \end{pmatrix} \begin{pmatrix} 8 & -3i \\ 3i & 0 \end{pmatrix} \begin{pmatrix} \frac{3}{\sqrt{10}} & \frac{i}{\sqrt{10}} \\ \frac{i}{\sqrt{10}} & \frac{3}{\sqrt{10}} \end{pmatrix} = \begin{pmatrix} 9 & 0 \\ 0 & -1 \end{pmatrix} = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} = \Lambda.$$

Die Umkehroperation dazu ist

$$U \Lambda U^\dagger = U (U^\dagger A U) U^\dagger = U U^\dagger A U U^\dagger = \mathbb{1} A \mathbb{1} = A.$$

Eine Änderung der Reihenfolge der Eigenvektoren $|v_i\rangle$ in U bewirkt eine entsprechende Änderung der Reihenfolge der Diagonalelemente λ_i in Λ :

$$\begin{pmatrix} |v_2\rangle & |v_1\rangle \end{pmatrix} =: \check{U} = \begin{pmatrix} \frac{i}{\sqrt{10}} & \frac{3}{\sqrt{10}} \\ \frac{3}{\sqrt{10}} & \frac{i}{\sqrt{10}} \end{pmatrix} \Rightarrow \check{U}^\dagger A \check{U} = \begin{pmatrix} -1 & 0 \\ 0 & 9 \end{pmatrix} = \begin{pmatrix} \lambda_2 & 0 \\ 0 & \lambda_1 \end{pmatrix}.$$

16 Spur einer Matrix

- Die Spur einer Matrix A ist die Summe ihrer Hauptdiagonalelemente. Für die Darstellung der Spur in der VON-Basis $\{|a_i\rangle\}$ gilt somit

$$\text{Sp}\{A\} = \sum_i \langle a_i | A | a_i \rangle = \sum_i A_{ii} .$$

- Die Spur einer Matrix ist unabhängig von der verwendeten VON-Basis (hier $\{|a_i\rangle\}$ und $\{|b_k\rangle\}$) von $\mathcal{H}$:

$$\begin{aligned} \text{Sp}\{A\} &= \sum_i \langle a_i | A | a_i \rangle \\ &= \sum_i \sum_{k,l} \langle a_i | b_k \rangle \langle b_k | A | b_l \rangle \underbrace{\langle b_l | a_i \rangle}_{=1} \\ &= \sum_i \sum_{k,l} \underbrace{\langle b_l | a_i \rangle}_{=1} \langle a_i | b_k \rangle \langle b_k | A | b_l \rangle \\ &= \sum_{k,l} \langle b_l | \left(\underbrace{\sum_i | a_i \rangle \langle a_i |}_{=1} \right) | b_k \rangle \langle b_k | A | b_l \rangle \\ &= \sum_{k,l} \underbrace{\langle b_l | b_k \rangle}_{\delta_{kl}} \langle b_k | A | b_l \rangle \\ &= \sum_{k,l} \delta_{kl} \langle b_k | A | b_l \rangle \\ \text{Sp}\{A\} &= \sum_k \langle b_k | A | b_k \rangle . \quad \square \end{aligned}$$

- Die zyklische Invarianz der Spur:

Zunächst zeigen wir

$$\text{Sp}\{AB\} = \text{Sp}\{BA\} :$$

$$\begin{aligned} \text{Sp}\{AB\} &= \sum_i \langle a_i | AB | a_i \rangle = \sum_{i,k} \langle a_i | A | a_k \rangle \langle a_k | B | a_i \rangle \\ &= \sum_k \sum_i \langle a_k | B | a_i \rangle \langle a_i | A | a_k \rangle \\ &= \sum_k \langle a_k | BA | a_k \rangle = \text{Sp}\{BA\} . \quad \square \end{aligned}$$

Daraus folgt die zyklische Invarianz der Spur

$$\begin{aligned} \text{Sp}\{A \cdot BC\} &= \text{Sp}\{BC \cdot A\} \\ &= \text{Sp}\{B \cdot CA\} = \text{Sp}\{CA \cdot B\} \end{aligned}$$

und für die unitäre Transformation der Spur einer Matrix A

$$\text{Sp}\{\tilde{A}\} = \text{Sp}\{U^\dagger A U\} = \text{Sp}\{\underbrace{U U^\dagger}_{=1} A\} = \text{Sp}\{A\} .$$

•

$$\text{Sp}\{AB\} = \sum_{i,j} A_{ij} B_{ji} \neq \text{Sp}\{A\} \cdot \text{Sp}\{B\} = \sum_i A_{ii} \cdot \sum_j B_{jj} .$$

- Spur des dyadischen Produkts

$$D = |u\rangle\langle v| , \quad |u\rangle \text{ orthogonal zu } |v\rangle \quad \Rightarrow \quad \langle u|v\rangle = \langle v|u\rangle = 0 :$$

$$\begin{aligned} \text{Sp}\{D\} &= \sum_i \langle a_i|u\rangle\langle v|a_i\rangle \\ &= \sum_i \langle v|a_i\rangle\langle a_i|u\rangle = \langle v| \underbrace{\sum_i |a_i\rangle\langle a_i|}_{=1} |u\rangle \\ \text{Sp}\{D\} &= \langle v|u\rangle = 0 . \end{aligned}$$

Veranschaulichendes Beispiel:

$$D = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} (0 \quad 0 \quad 1) = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix} \quad \Rightarrow \quad \text{Sp}\{D\} = 0 .$$

17 Kombinatorik, Verteilungsmöglichkeiten und ihre Wahrscheinlichkeit

In diesem Kapitel werden einige einfache Beispiele gezeigt zur Bestimmung der Anzahl und der Wahrscheinlichkeit von Verteilungsmöglichkeiten, die sich bei der Anordnung oder Aufteilung von Teilchenmengen nach bestimmten Kriterien ergeben.¹ Den Elementen in der mathematischen Kombinatorik entsprechen in der Physik häufig Teilchen. Deshalb werden wir dort, wo es sinnvoll erscheint, den Begriff „Teilchen“ anstatt des Begriffs „Element“ verwenden.

- Wie groß ist die Anzahl der Verteilungsmöglichkeiten, wenn man N unterscheidbare Teilchen auf 2 Töpfe verteilt?

Das erste Teilchen fällt entweder in den einen oder in den anderen Topf und liefert somit zwei Möglichkeiten. Für jede dieser beiden Möglichkeiten fällt das zweite Teilchen wieder in den einen oder den anderen Topf, sodass aus den zwei Möglichkeiten vier werden. Und so verdoppelt sich die Anzahl der Möglichkeiten beim Hinzufügen jedes weiteren Teilchens. Z. B. resultieren dann bei 3 Teilchen $2^3 = 8$ Verteilungsmöglichkeiten.

Wir stellen fest: Es gibt

$$2^N \text{ Möglichkeiten,} \quad (17.1)$$

eine Menge von N unterscheidbaren Teilchen in zwei Teilmengen N_1 und N_2 mit $N_1 + N_2 = N$ aufzuteilen. Hierbei spielt die Reihenfolge der Teilchen in den Teilmengen *keine* Rolle, denn die Teilchen wurden in *einer* beliebigen Reihenfolge auf die zwei Töpfe verteilt.

Wie man sich analog überlegen kann, gibt es allgemein

$$m^N \text{ Möglichkeiten,} \quad (17.2)$$

N unterscheidbare Teilchen auf m Töpfe zu verteilen, wenn die Reihenfolge bzw. Anordnung der Teilchen in den Töpfen keine Rolle spielt.

- **Permutationen:**

Permutation heißt Vertauschung im Sinne von Umordnung. Die Permutation von N unterscheidbaren Elementen liefert

$$N! = 1 \cdot 2 \cdot 3 \cdot 4 \cdots N$$

Anordnungen (Reihenfolgen, Permutationen). Z. B. ergeben die 3 Elemente a, b und c die folgenden $3! = 1 \cdot 2 \cdot 3 = 6$ verschiedenen Permutationen:

$$abc, acb, bac, bca, cab, cba.$$

¹Ergänzende Erläuterungen zu diesem Thema finden sich z. B. in M. Alonso und E. J. Finn, Quantenphysik und Statistische Physik, 4. Aufl., Oldenbourg-Verlag, München, Wien, 2005, Seite 472 bis Seite 474 und in Lehr- und Übungsbuch Mathematik, Band IV, Fachbuchverlag Leipzig - Köln, 13. Aufl., 1992, Seite 256 bis Seite 270.

- **Variationen ohne Wiederholung:**

Die Anzahl $V_N^{(k)}$ aller Variationen von N Elementen in geordneten Gruppen (mit Berücksichtigung der Reihenfolge) zu je k Elementen ohne Wiederholung ist

$$V_N^{(k)} = N \cdot (N-1) \cdot (N-2) \cdot (N-3) \cdot \dots \cdot (N-k+1) = \frac{N!}{(N-k)!}, \quad (17.3)$$

$$V_N^{(k)} = \binom{N}{k} k!. \quad (17.4)$$

Das kann man sich so vorstellen: Wir wollen aus N Elementen alle möglichen k -elementigen Gruppen bilden. Im ersten Schritt, d. h. bei der Anordnung des ersten Elements in den Gruppen, können noch alle N Elemente verwendet werden, sodass wir mit dem ersten Schritt N Anordnungen aus einem Element erhalten. Jedes dieser N Elemente wird im zweiten Schritt nacheinander mit allen Elementen außer mit sich selbst, also jeweils mit $N-1$ Elementen verbunden. So erhalten wir mit dem zweiten Schritt $N \cdot (N-1)$ Anordnungen aus zwei Elementen. Im dritten Schritt wird jede dieser $N \cdot (N-1)$ Anordnungen wieder nacheinander mit allen Elementen außer mit den eigenen beiden verbunden. Mit dem dritten Schritt erhalten wir folglich $N \cdot (N-1) \cdot (N-2)$ Anordnungen aus je drei Elementen. Wir setzen diese Prozedur fort bis zum k -ten Schritt, mit dem wir dann schließlich $N \cdot (N-1) \cdot (N-2) \cdot \dots \cdot (N-k+1)$ Anordnungen aus jeweils k Elementen erhalten.

Beispiel:

$N = 3$ (Elemente a, b und c), $k = 2$, $V_N^{(k)} = 6$:
 ab, ac, ba, bc, ca, cb .

- **Variationen mit Wiederholung:**

Die Anzahl ${}^wV_N^{(k)}$ aller Variationen von N Elementen in geordneten Gruppen (mit Berücksichtigung der Reihenfolge) zu je k Elementen mit Wiederholung ist

$${}^wV_N^{(k)} = N^k. \quad (17.5)$$

Die Prozedur zur Bildung der Variationen mit Wiederholung ist ähnlich der Variation ohne Wiederholung. Der Unterschied besteht lediglich darin, dass bei jedem der k Schritte alle N Elemente mit den zuvor gebildeten Anordnungen verbunden werden, also auch die jeweils gleichen Elemente miteinander, wodurch die Wiederholungen entstehen.

Beispiel:

$N = 3$ (Elemente a, b und c), $k = 2$, ${}^wV_N^{(k)} = 9$:
 $aa, ab, ac, ba, bb, bc, ca, cb, cc$.

• **Kombinationen ohne Wiederholung:**

Wir teilen jetzt N unterscheidbare Teilchen auf zwei Kammern so auf, dass stets in der Kammer I N_1 Teilchen und in der Kammer II N_2 Teilchen vorhanden sind, wobei $N_2 = N - N_1$ sei. Da die Bildung der Kombinationen ohne Wiederholung, d. h. ohne Zurücklegen erfolgen soll, können die unterscheidbaren Teilchen in den einzelnen Kammern nicht mehrfach vorkommen.

Beispiel: $N = 3$ (Teilchen a , Teilchen b , Teilchen c), $N_1 = 2$, $N_2 = 1$

I	II	
a	b	c
a	c	b
b	a	c
b	c	a
c	a	b
c	b	a

$N! = 3! = 6$ Verteilungen

Wie wir sehen, ist hier die Reihenfolge der Teilchen innerhalb der Kammern von Bedeutung.

Wenn die Reihenfolge der Teilchen in den Kammern keine Rolle spielen soll, muss durch N_i dividiert werden. In unserem Beispiel mit $N = 3$, $N_1 = 2$ und $N_2 = 1$ resultieren dann

$$C_N(N_1) = \frac{N!}{N_1! \cdot N_2!} = \frac{N!}{N_1! \cdot (N - N_1)!} = \binom{N}{N_1} \quad (17.6)$$

$$= \frac{3!}{2! \cdot 1!} = \frac{6}{2} = 3 \text{ Verteilungen :} \quad (17.7)$$

$$ab|c, \quad ac|b, \quad bc|a.$$

Üblicherweise schreibt man für die Anzahl $C_N^{(k)}$ der Kombinationen von N Elementen zur k -ten Klasse ohne Wiederholung

$$C_N^{(k)} = \frac{N!}{k!(N - k)!} = \binom{N}{k}. \quad (17.8)$$

Kombiniert man N unterscheidbare Teilchen ohne Wiederholung zu Gruppen von $N_i = 1, 2, 3, \dots, N$ Teilchen, so resultieren nach sukzessiver Anwendung von (17.6)

$$C_N = \binom{N}{1} + \binom{N}{2} + \binom{N}{3} + \dots + \binom{N}{N} = 2^N - 1 \quad (17.9)$$

Kombinationen. Dass dies $2^N - 1$ ergibt, kann mit dem weiter unten aufgeführten binomischen Satz (17.19) wie folgt gezeigt werden:

$$2^N = (1 + 1)^N = \underbrace{\binom{N}{0}}_{=1} + \binom{N}{1} + \binom{N}{2} + \binom{N}{3} + \dots + \binom{N}{N}. \quad (17.10)$$

- Jetzt zeigen wir (17.6) mit Hilfe einer anderen Argumentation:

Wenn wir N_1 unterscheidbare Teilchen aus N unterscheidbaren Teilchen auswählen, haben wir bei der Wahl des ersten Teilchens N Wahlmöglichkeiten, bei der Wahl des zweiten Teilchens aus den verbliebenen $N - 1$ Teilchen nur noch $N - 1$ Möglichkeiten usw. Die Anzahl von Möglichkeiten bei der Auswahl von N_1 unterscheidbaren Teilchen aus N unterscheidbaren Teilchen ist demzufolge

$$N \cdot (N - 1) \cdot (N - 2) \cdot \dots \cdot (N - N_1 + 1) = \frac{N!}{(N - N_1)!} = x. \quad (17.11)$$

Hierbei spielt aber die Reihenfolge der N_1 Teilchen eine Rolle, denn in x sind alle möglichen Reihenfolgen der N_1 Teilchen enthalten. Es resultieren also y verschiedene Gruppen mit jeweils $N_1!$ Möglichkeiten. Die y Gruppen unterscheiden sich in der Teilchenzusammensetzung der N_1 aus N Teilchen, wobei jede Gruppe aus $N_1!$ verschiedenen Anordnungen derselben N_1 Teilchen besteht. Das bedeutet: y Gruppen $\cdot N_1!$ Möglichkeiten = maximale Anzahl x von Möglichkeiten bei Unterscheidbarkeit der Teilchen bzw.

$$y \cdot N_1! = x = \frac{N!}{(N - N_1)!} \Leftrightarrow y = \frac{N!}{N_1! \cdot (N - N_1)!} = \binom{N}{N_1}. \quad (17.12)$$

Wir stellen fest:

$y = \binom{N}{N_1} = C_N(N_1)$ ist die Anzahl von Teilchen-Kombinationsmöglichkeiten bei der Auswahl von N_1 Teilchen aus N Teilchen, wobei die Reihenfolge der N_1 Teilchen in der jeweiligen Kombination keine Rolle spielt.

- **Kombinationen mit Wiederholung:**

Für die Anzahl ${}^w C_N^{(k)}$ der Kombinationen von N Elementen zur k -ten Klasse mit Wiederholung schreibt man üblicherweise

$${}^w C_N^{(k)} = \frac{(k + N - 1)!}{k!(N - 1)!} = \binom{N + k - 1}{k}. \quad (17.13)$$

Plausibilisierungen von (17.13) sind in den Abschnitten ?? und ?? zu finden. Dort wird die Bildung von Kombinationen mit Wiederholung so interpretiert, dass dabei die Unterscheidbarkeit der Teilchen verloren geht. ${}^w C_N^{(k)}$ ist dann die Anzahl der Realisierungsmöglichkeiten für die Aufteilung von k nicht unterscheidbaren Teilchen auf N Zellen.

- (17.6) ist die Anzahl der Kombinationsmöglichkeiten, wenn N unterscheidbare Teilchen auf zwei Kammern aufgeteilt werden und dabei die Reihenfolge der Teilchen innerhalb der Kammern nicht berücksichtigt wird.

Wir werden jetzt zeigen, wie man die Anzahl $C_N(N_i)$ der Kombinationsmöglichkeiten bei der Aufteilung von N unterscheidbaren Teilchen auf m Kammern berechnet, wenn die Anzahl der Teilchen in den Kammern $N_1, N_2, N_3, \dots, N_n$ beträgt und wenn $\sum_{i=1}^m N_i = N$ gilt. Die Reihenfolge der Teilchen bei der Aufteilung auf die Kammern bzw. innerhalb der Kammern soll dabei wieder keine Rolle spielen.

In der Herleitung von (17.12) standen uns bei der Aufteilung bezüglich der ersten Kammer N Teilchen zur Verfügung. Daraus resultierte

$$\frac{N!}{N_1! \cdot (N - N_1)!} \quad (17.14)$$

Bei der Aufteilung bezüglich der zweiten Kammer stehen uns dann nur noch $N - N_1$ Teilchen zur Verfügung, sodass wir aus (17.14) für die neue Situation

$$\frac{(N - N_1)!}{N_2! \cdot (N - N_1 - N_2)!} \quad (17.15)$$

erhalten und bei der Aufteilung auf die dritte Kammer analog

$$\frac{(N - N_1 - N_2)!}{N_3! \cdot (N - N_1 - N_2 - N_3)!} \quad (17.16)$$

So verfahren wir weiter, bis zur Aufteilung auf die m -te Kammer mit

$$\frac{(N - N_1 - N_2 - \dots - N_{m-1})!}{N_m! \cdot (N - N_1 - N_2 - N_3 - \dots - N_m)!} \quad (17.17)$$

Wir stellen fest:

Das Produkt aus (17.14), (17.15) und aus allen folgenden Anzahlen der Kombinationsmöglichkeiten bis (17.17) ergibt die Gesamtzahl $C_N(N_i)$ der Kombinationsmöglichkeiten bei der Aufteilung von N unterscheidbaren Teilchen auf m Kammern, wenn die Reihenfolge der Teilchen innerhalb der Kammern keine Rolle spielt:

$$C_N(N_i) = \frac{N!}{N_1! \cdot N_2! \cdot N_3! \cdot \dots \cdot N_m!} = \frac{N!}{\prod_{i=1}^m N_i!} \quad (17.18)$$

Die Anzahl $C_N(N_i)$ gilt hierbei für eine bestimmte Verteilung bzw. Folge der Teilchenzahlen N_i unter der Bedingung $\sum_{i=1}^m N_i = N$ und kann als Gewicht² dieser Verteilung angesehen werden.

Anders gesagt:

$C_N(N_i)$ in (17.18) ist gleich der Anzahl von Permutationen von N Teilchen, die sich aus N_1 Teilchen einer 1. Art, N_2 Teilchen einer 2. Art, N_3 Teilchen einer 3. Art, $\dots$, N_m Teilchen einer m -ten Art zusammensetzen, wenn $\sum_{i=1}^m N_i = N$ gilt und wenn innerhalb der m Gruppen die jeweiligen N_i Teilchen selbst nicht permutiert werden.

²Das Gewicht einer Verteilung ist proportional zu ihrer Wahrscheinlichkeit, denn die Wahrscheinlichkeit ist allgemein das auf 1 normierte Gewicht.

- Im folgenden Beispiel werden wir den binomischen Satz gebrauchen. Deshalb wird er hier zur Erinnerung mit den anschließend verwendeten Größen angeschrieben:

$$\begin{aligned}
(p_1 + p_2)^N &= \binom{N}{0} p_1^N + \binom{N}{1} p_1^{N-1} \cdot p_2 + \binom{N}{2} p_1^{N-2} \cdot p_2^2 + \dots \\
&\quad + \binom{N}{N-1} p_1 \cdot p_2^{N-1} + \binom{N}{N} p_2^N \\
&= \sum_{N_1=0}^N \binom{N}{N_1} p_1^{N-N_1} \cdot p_2^{N_1} \\
&\quad \text{und weil } p_1 \text{ und } p_2 \text{ vertauschbar sind:}
\end{aligned}$$

$$(p_1 + p_2)^N = \sum_{N_1=0}^N \binom{N}{N_1} p_1^{N_1} \cdot p_2^{N-N_1} . \quad (17.19)$$

Wir betrachten jetzt zur Veranschaulichung zunächst nur zwei Teilchen, die wieder auf zwei Kammern gemäß $N_1 + N_2 = N = 2$ aufgeteilt werden sollen.

Das Volumen der Kammer I sei $V_1 = 3/4$, das Volumen der Kammer II sei $V_2 = 1/4$ und das Gesamtvolumen ist folglich $V = 1$. Weil das Volumen V_1 dreimal so groß wie das Volumen V_2 ist, soll auch die Wahrscheinlichkeit p_1 , dort ein Teilchen anzutreffen, dreimal so groß sein wie die Wahrscheinlichkeit p_2 , ein Teilchen im Volumen V_2 anzutreffen. Wir erhalten also für die Wahrscheinlichkeit, ein Teilchen in Kammer I anzutreffen, $p_1 = V_1/V = 3/4$, für die Wahrscheinlichkeit, ein Teilchen in Kammer II anzutreffen, $p_2 = V_2/V = 1/4$ und für die Wahrscheinlichkeit, im Gesamtvolumen $V_1 + V_2 = V$ ein Teilchen anzutreffen, folglich die Summe

$$p_1 + p_2 = 1 . \quad (17.20)$$

Bei der Aufteilung von zwei unterscheidbaren Teilchen auf zwei Kammern ergeben sich die vier Kombinationsmöglichkeiten mit ihren zugehörigen Wahrscheinlichkeiten wie folgt:

$N_1=2, N_2=0$	$N_1=1, N_2=1$	$N_1=1, N_2=1$	$N_1=0, N_2=2$
$p_1^{N_1} p_2^{N_2} = p_1^2$	$p_1^{N_1} p_2^{N_2} = p_1^1 p_2^1$	$p_1^{N_1} p_2^{N_2} = p_1^1 p_2^1$	$p_1^{N_1} p_2^{N_2} = p_2^2$
$\frac{3}{4} \cdot \frac{3}{4} = \frac{9}{16}$	$\frac{3}{4} \cdot \frac{1}{4} = \frac{3}{16}$	$\frac{3}{4} \cdot \frac{1}{4} = \frac{3}{16}$	$\frac{1}{4} \cdot \frac{1}{4} = \frac{1}{16}$
$P_N(N_1) = \frac{9}{16}$	$P_N(N_1) = P_N(1) = \frac{6}{16}$		$P_N(N_1) = \frac{1}{16}$
$\sum_{N_1=0}^2 P_N(N_1) = \frac{16}{16} = 1$			

Hierbei ist $P_N(N_1)$ die Wahrscheinlichkeit, eine Kombination mit N_1 Teilchen in Kammer I anzutreffen.

Wir stellen fest:

Jede einzelne Kombination mit $N_1 + N_2 = N$ besitzt die Wahrscheinlichkeit

$$p_1^{N_1} \cdot p_2^{N_2} , \quad (17.21)$$

sodass die Wahrscheinlichkeit für die $C_N(N_1)$ Kombinationen mit der bestimmten Teilchenzahl N_1 (in Kammer I)

$$P_N(N_1) = C_N(N_1) \cdot p_1^{N_1} \cdot p_2^{N_2} \quad (17.22)$$

$$P_N(N_1) = \frac{N!}{N_1! \cdot (N - N_1)!} \cdot p_1^{N_1} \cdot p_2^{N - N_1} \quad (17.23)$$

$$= \binom{N}{N_1} \cdot p_1^{N_1} \cdot p_2^{N_2} \quad (17.24)$$

ist. Die Summe der Wahrscheinlichkeiten aller Kombinationen unter Einbeziehung aller möglichen Teilchenzahlen N_i also von $N_i = 0$ bis $N_i = N$, ist mit $p_1 + p_2 = 1$ definitionsgemäß

$$(p_1 + p_2)^N = 1^N = \sum_{N_1=0}^N \binom{N}{N_1} p_1^{N_1} \cdot p_2^{N - N_1} \quad (17.25)$$

$$= \sum_{N_1=0}^N P_N(N_1) = 1 . \quad \square \quad (17.26)$$

- Ein Beispiel zur Vertiefung:

Teilchen a , Teilchen b und Teilchen c seien wieder die $N = 3$ unterscheidbaren Teilchen, die auf die beiden Kammern I und II aufgeteilt werden sollen. Die Anzahl der sich daraus ergebenden Teilchenkombinationen auf die zwei Kammern, wenn die Reihenfolge der Teilchen in den Kammern *keine* Rolle spielt, ist gemäß (17.1), unserem ersten Beispiel, $2^N = 2^3 = 8$. Dabei läuft die Teilchenzahl N_1 in Kammer I von 0 bis N .

$C_N(N_1) = \binom{N}{N_1}$ in (17.6) ist aber die Anzahl der Kombinationen nur für *eine bestimmte* Teilchenzahl N_1 , wobei die Reihenfolge der Teilchen ebenfalls keine Rolle spielt. Wir können deshalb für die Anzahl der Kombinationen 2^N mit (17.6) auch schreiben:

$$\sum_{N_1=0}^N C_N(N_1) = \sum_{N_1=0}^N \frac{N!}{N_1! \cdot (N - N_1)!} = \sum_{N_1=0}^N \binom{N}{N_1} \quad (17.27)$$

$$= 1 + 3 + 3 + 1 = 8 = 2^N . \quad (17.28)$$

In der folgenden Tabelle sind die acht Kombinationen mit ihren zugehörigen Wahrscheinlichkeiten aufgeführt.

	I	II	N_1	$p_1^{N_1} \cdot p_2^{N_2}$
1.	$a \quad b \quad c$	—	3	$\frac{3}{4} \cdot \frac{3}{4} \cdot \frac{3}{4} = \frac{27}{64}$
2.	$a \quad b$	c	2	$\frac{3}{4} \cdot \frac{3}{4} \cdot \frac{1}{4} = \frac{9}{64}$
3.	$a \quad c$	b	2	$\frac{3}{4} \cdot \frac{3}{4} \cdot \frac{1}{4} = \frac{9}{64}$
4.	$b \quad c$	a	2	$\frac{3}{4} \cdot \frac{3}{4} \cdot \frac{1}{4} = \frac{9}{64}$
5.	a	$b \quad c$	1	$\frac{3}{4} \cdot \frac{1}{4} \cdot \frac{1}{4} = \frac{3}{64}$
6.	b	$a \quad c$	1	$\frac{3}{4} \cdot \frac{1}{4} \cdot \frac{1}{4} = \frac{3}{64}$
7.	c	$a \quad b$	1	$\frac{3}{4} \cdot \frac{1}{4} \cdot \frac{1}{4} = \frac{3}{64}$
8.	—	$a \quad b \quad c$	0	$\frac{1}{4} \cdot \frac{1}{4} \cdot \frac{1}{4} = \frac{1}{64}$
				$\sum_{N_1=0}^3 P_N(N_1) = \frac{64}{64} = 1$

Die Wahrscheinlichkeiten für die Kombinationen mit bestimmter Teilchenzahl N_1 in Kammer I sind:

$$N_1 = 3 \quad : \quad P_N(N_1) = P_N(3) = 1 \cdot \frac{27}{64} = \frac{27}{64} ,$$

$$N_1 = 2 \quad : \quad P_N(N_1) = P_N(2) = 3 \cdot \frac{9}{64} = \frac{27}{64} ,$$

$$N_1 = 1 \quad : \quad P_N(N_1) = P_N(1) = 3 \cdot \frac{3}{64} = \frac{9}{64} ,$$

$$N_1 = 0 \quad : \quad P_N(N_1) = P_N(0) = 1 \cdot \frac{1}{64} = \frac{1}{64} .$$

- Mit (17.23) lässt sich das Gewicht (17.18) einer Verteilung von N Teilchen auf m Kammern ebenfalls allgemein als Wahrscheinlichkeit ausdrücken, wenn die Besetzungswahrscheinlichkeiten p_i der einzelnen Kammern bekannt sind:

$$P = \frac{N!}{N_1! \cdot N_2! \cdot N_3! \cdot \dots \cdot N_m!} \cdot p_1^{N_1} \cdot p_2^{N_2} \cdot p_3^{N_3} \cdot \dots \cdot p_m^{N_m} \quad (17.29)$$

$$= M \cdot p_1^{N_1} \cdot p_2^{N_2} \cdot p_3^{N_3} \cdot \dots \cdot p_m^{N_m} \quad (17.30)$$

$$P = N! \prod_i \frac{p_i^{N_i}}{N_i!} . \quad (17.31)$$

- Anstelle der Wahrscheinlichkeiten p_i verwenden wir jetzt die Gewichte g_i der m Kammern. Dies können z. B. die Entartungsgrade der Kammern sein. Zu ihrer Veranschaulichung nehmen wir an, dass die Kammern in Zellen unterteilt sind. So wäre dann z. B. die i -te Kammer in g_i Zellen unterteilt, auf die sich dann die N_i Teilchen gemäß (17.2) in $g_i^{N_i}$ verschiedenen Kombinationen aufteilen können. So liefern die Gewichte g_i bzw. die g_i Zellen *für jede* Kammer mit der Teilchenzahl N_i die Anzahl von $g_i^{N_i}$ zusätzlichen Realisierungsmöglichkeiten.

Das Gewicht $W(N_i, g_i)$ bzw. die Gesamtanzahl der Realisierungsmöglichkeiten aus den N unterscheidbaren Teilchen unter Berücksichtigung aller m Kammern und ihrer Gewichte g_i bei einer bestimmten Verteilung bezüglich der Teilchenzahlen N_i erhält man folglich durch Multiplikation von (17.18) mit allen $g_i^{N_i}$:

$$W(N_i, g_i) = \frac{N!}{N_1! \cdot N_2! \cdot \dots \cdot N_m!} \cdot g_1^{N_1} \cdot g_2^{N_2} \cdot \dots \cdot g_m^{N_m} \quad (17.32)$$

$$= C_N(N_i) \cdot g_1^{N_1} \cdot g_2^{N_2} \cdot g_3^{N_3} \cdot \dots \cdot g_m^{N_m} \quad (17.33)$$

$$W(N_i, g_i) = N! \prod_i \frac{g_i^{N_i}}{N_i!}, \quad \text{Bedingung: } \sum_i N_i = N. \quad (17.34)$$

18 Gesetz der großen Zahlen – Binomialverteilung

(Siehe auch Torsten Fließbach, Statistische Physik, Lehrbuch zur Theoretischen Physik IV, 4. Aufl., Spektrum Akademischer Verlag, Heidelberg, 2007, Seite 10 bis Seite 21, Franz Embacher, Random Walk in einer Dimension, Universität Wien, homepage.univie.ac.at/ franz.embacher/Lehre/aussermathAnw2012/RandomWalk.pdf und

Wolfgang Nolting, Grundkurs Theoretische Physik 6, Statistische Physik, 4. Aufl., Springer-Verlag, Berlin, Heidelberg, New York, 2005, Seite 6 bis Seite 12.)

Um von vornherein Klarheit zu schaffen, leiten wir dieses Kapitel mit einigen Hinweisen zum Sprachgebrauch ein:

Wir werden im Folgenden vom Mittelwert und vom Erwartungswert einer Messgröße x sprechen. Mit dem Begriff Mittelwert $\bar{x}$ bezeichnen wir den real bzw. empirisch bestimmbaren *arithmetischen Mittelwert* aus einer *endlichen* Anzahl N von Messwerten. Mit dem Begriff Erwartungswert $\langle x \rangle$ bezeichnen wir den nur theoretisch bzw. fiktiv bestimmbaren *wahren Mittelwert* aus der Anzahl $N \rightarrow \infty$ von Messwerten. Sowohl die Standardabweichung bzw. Schwankung

$$\Delta x \equiv \sigma = \sqrt{\frac{\sum (\bar{x} - x_i)^2}{N - 1}} \quad (18.1)$$

der Einzelmessung x_i als auch die Standardabweichung bzw. Schwankung

$$\sigma_m = \sqrt{\frac{\sum (\bar{x} - x_i)^2}{N(N - 1)}} \quad (18.2)$$

des arithmetischen Mittels $\bar{x}$ basieren auf dem wahren Mittelwert bzw. Erwartungswert $\langle x \rangle$.¹ Weil wir es in der Thermodynamik bzw. in der statistischen Physik mit $N \gg 1$ bzw. $N \rightarrow \infty$ zu tun haben, werden wir nicht die exakten Formeln (18.1) und (18.2), sondern ausschließlich deren Näherung gemäß $\lim_{N \rightarrow \infty} (N - 1) = N$ verwenden.

Zum schnelleren Verständnis wiederholen wir zunächst die erforderlichen grundlegenden mathematischen „Werkzeuge“:

- In der Thermodynamik bzw. in der statistischen Physik betrachten wir meistens Systeme mit Teilchenzahlen N in der Größenordnung von 10^{23} , also mit sehr großen Teilchenzahlen.²
- Wir müssen unterscheiden zwischen dem Mittelwert $\bar{x}$ der Messwerte x_i und dem Erwartungswert $\langle x \rangle$ der Messgröße x und werden den Unterschied zwischen diesen beiden Begriffen am einfachen Beispiel der beim Würfeln erzielten Punktzahlen illustrieren. Dabei setzen wir voraus, dass die verwendeten Würfel so konstruiert

¹Eine leicht verständliche Herleitung dieser Zusammenhänge findet man z. B. im Abschnitt „Messgenauigkeit und Messfehler“ des Springer-Lehrbuchs von Wolfgang Demtröder, Experimentalphysik 1, Mechanik und Wärme, 3. Auflage, Springer, Berlin, Heidelberg, New York, 2003, Seite 29 bis 36.

²Unter Normbedingungen (273, 15 K, 1013, 25 hPa) besteht 1 ml Gas aus ca. $2,7 \cdot 10^{19}$ Teilchen. 1 ml Wasser besteht aus ca. $3,35 \cdot 10^{22}$ Molekülen und 1 cm³ Kupfer bei 25°C aus ca. $0,85 \cdot 10^{23}$ Atomen.

sind, dass alle Seitenflächen beim Würfeln mit der gleichen Wahrscheinlichkeit und unabhängig voneinander oben liegen können. Die Punktzahlen auf den sechs Seitenflächen des Würfels sind $x_i \in \{1, 2, 3, 4, 5, 6\}$, $i = 1, 2, \dots, 6$.³ Beim Würfeln mit nur *einem* Würfel sei n die Gesamtanzahl der Würfe und N_i die Anzahl der Würfe mit der Punktzahl x_i unter der Bedingung $n = \sum_i N_i$.

Wenn wir mit *einem* Würfel z. B. sieben ($n=7$) mal werfen (würfeln) und dabei nacheinander die Punktzahlen 5, 1, 4, 2, 2, 6, 1 erzielen, ist der **Mittelwert** $\bar{x}$ aus den Punktzahlen x_i

$$\bar{x} = \frac{5 + 1 + 4 + 2 + 2 + 6 + 1}{7} = \frac{21}{7} = 3 \quad (18.3)$$

$$= \frac{(2 \cdot 1) + (2 \cdot 2) + (1 \cdot 4) + (1 \cdot 5) + (1 \cdot 6)}{7} \quad (18.4)$$

$$= \frac{(N_1 \cdot x_1) + (N_2 \cdot x_2) + (N_4 \cdot x_4) + (N_5 \cdot x_5) + (N_6 \cdot x_6)}{n}, \quad (18.5)$$

$$\bar{x} = \frac{\sum_i N_i x_i}{n}. \quad (18.6)$$

Wenn wir aber immer öfter würfeln, d. h. mit wachsendem n bzw. im Grenzfall $n \rightarrow \infty$, stellen wir fest, dass der Mittelwert einem bestimmten Wert, dem **Erwartungswert** $\langle x \rangle$, zustrebt:

$$\langle x \rangle = \lim_{n \rightarrow \infty} \frac{\sum_i N_i x_i}{n}. \quad (18.7)$$

Während der Mittelwert in der Praxis (real) ermittelt werden kann, ist der Erwartungswert demzufolge von theoretischer Natur.

Den Erwartungswert für unser Beispiel finden wir, wenn wir gemäß den Voraussetzungen hinsichtlich der Würfeigenschaften von der (theoretischen) Wahrscheinlichkeit $p_i = \frac{1}{6}$ für das Auftreten jeder Punktzahl x_i ausgehen, sodass konventionsgemäß für die Wahrscheinlichkeit des Auftretens aller sechs möglichen Punktzahlen insgesamt $\sum_i p_i = \sum_{i=1}^6 p_i = 6 \cdot \frac{1}{6} = 1$ resultiert. Es soll also jede mögliche Punktzahl von 1 bis 6 mit der gleichen Wahrscheinlichkeit von $\frac{1}{6}$ auftreten, sodass

$$\langle x \rangle = \sum_{i=1}^6 p_i \cdot x_i = \sum_{i=1}^6 \frac{1}{6} \cdot x_i = \frac{1 + 2 + 3 + 4 + 5 + 6}{6} = \frac{21}{6} = 3,5. \quad (18.8)$$

In (18.8) wird die Grenzwertbildung $n \rightarrow \infty$ deshalb nicht sichtbar, weil sie bereits in der (theoretischen) Wahrscheinlichkeit $p_i = \frac{1}{6}$ enthalten ist. $p_i = \frac{1}{6}$ erhält man nämlich nur dann exakt, wenn man mit einem „idealen“ Würfel unendlich mal ($n \rightarrow \infty$) würfelt.

³Im Fall des Würfels ist der Laufindex i gleich der Punktzahl der zugehörigen Seitenfläche. Das ist allgemein nicht der Fall. In der Thermodynamik bzw. in der statistischen Physik sind die x_i z. B. Energiewerte von Quantenzuständen.

- Der Mittelwert, den wir z. B. durch n -maliges Würfeln (in einer zeitlichen Abfolge) mit nur einem Würfel erzielen, heißt **Zeitmittel**. Wenn wir aber mit einer Schar bzw. einem Ensemble aus $N = n$ gleichartigen Würfeln nur einmal würfeln, erhalten wir prinzipiell den gleichen Mittelwert, bezeichnen ihn aber als Ensemble- oder **Scharmittel**. So lässt sich z. B. in der Thermodynamik der Mittelwert einer Messgröße eines idealen N -Teilchensystems⁴ berechnen, indem man entweder $n = N$ Messungen nacheinander an nur *einem* Teilchen durchführt und anschließend mittelt (Zeitmittel) oder indem man gleichzeitig alle N Teilchen misst und diese N Messwerte mittelt (Scharmittel). Praktischerweise verwenden wir im Folgenden das Scharmittel.
- Von der Gesamtteilchenzahl N haben also bei einer Messung jeweils N_i Teilchen den Wert x_i . Dann gilt unter der Bedingung $\sum_i N_i = N$ für die (theoretische) **Wahrscheinlichkeit**, ein Teilchen mit dem Wert x_i zu finden,

$$p(x_i) = p_i = \lim_{N \rightarrow \infty} \frac{N_i}{N} , \quad (18.9)$$

sodass konventionsgemäß $\sum_i p_i = 1$.

Weil wir es in der statistischen Physik, wie bereits betont, mit sehr großen Teilchenzahlen zu tun haben, werden wir im Folgenden, wenn auch nicht ganz korrekt, auf die Angabe des Limes für $N \rightarrow \infty$ verzichten und vereinfachend schreiben

$$p_i = \frac{N_i}{N} , \quad (18.10)$$

- Der **Erwartungswert** der Messgröße x über alle N Teilchen ist damit

$$\langle x \rangle = \frac{\sum_i N_i \cdot x_i}{N} = \sum_i p_i \cdot x_i . \quad (18.11)$$

Wie man sieht, ist die Summe X der Messwerte aller N Teilchen

$$X = \sum_i N_i \cdot x_i = N \cdot \langle x \rangle . \quad (18.12)$$

- Die Messwerte x_i weichen vom Erwartungswert jeweils um $x_i - \langle x \rangle$ ab. Weil der Erwartungswert dieser Abweichungen gemäß

$$\begin{aligned} \frac{\sum_i N_i (x_i - \langle x \rangle)}{N} &= \sum_i p_i (x_i - \langle x \rangle) = \sum_i p_i x_i - \langle x \rangle \sum_i p_i \\ &= \langle x \rangle - \langle x \rangle = 0 \end{aligned} \quad (18.13)$$

verschwindet, verwenden wir die **mittlere quadratische Abweichung** bzw.

⁴In einem idealen N -Teilchensystem können die Teilchen unabhängig voneinander verschiedene Werte der Messgröße mit einer dem jeweiligen Wert entsprechenden Wahrscheinlichkeit annehmen.

Schwankung⁵ Δx :

$$(\Delta x)^2 = \sum_i p_i (x_i - \langle x \rangle)^2 = \sum_i \frac{N_i}{N} (x_i - \langle x \rangle)^2 \Rightarrow \quad (18.14)$$

$$\Delta x = \sqrt{\sum_i p_i (x_i - \langle x \rangle)^2} \quad (18.15)$$

$$= \sqrt{\sum_i p_i x_i^2 - 2\langle x \rangle \sum_i p_i x_i + \langle x \rangle^2 \sum_i p_i} , \quad (18.16)$$

$$\boxed{\Delta x = \sqrt{\langle x^2 \rangle - \langle x \rangle^2}} . \quad (18.17)$$

- Das Verhältnis zwischen Schwankung und Erwartungswert der Messgröße x ist die

$$\text{relative Schwankung} \quad \frac{\Delta x}{\langle x \rangle} . \quad (18.18)$$

- Die Summe aller N Abweichungsquadrate vom Erwartungswert $\langle x \rangle$ ist in Analogie zu (18.14) das Schwankungsquadrat $(\Delta X)^2$ des Gesamtsystems. Unter Berücksichtigung von (18.14) hat es die Form

$$N \cdot (\Delta x)^2 = \sum_i N_i (x_i - \langle x \rangle)^2 = (\Delta X)^2 . \quad (18.19)$$

Die Schwankung des Gesamtsystems ist folglich

$$\boxed{\Delta X = \sqrt{N} \Delta x} . \quad (18.20)$$

- Weil wir die Teilchen als voneinander unabhängig angenommen haben, gelten Erwartungswert $\langle x \rangle$, Schwankung Δx und relative Schwankung $\Delta x / \langle x \rangle$ gleichermaßen für jedes einzelne der N Teilchen. Davon ausgehend erhalten wir für das (makroskopische) Gesamtsystem aus N Teilchen den Erwartungswert

$$\boxed{\langle X \rangle = N \cdot \langle x \rangle} . \quad (18.21)$$

- Analog zu (18.18) ist damit die relative Schwankung für das Gesamtsystem

$$\boxed{\frac{\Delta X}{\langle X \rangle} = \frac{\sqrt{N} \cdot \Delta x}{N \cdot \langle x \rangle} = \frac{1}{\sqrt{N}} \cdot \frac{\Delta x}{\langle x \rangle} = \frac{1}{\sqrt{N}} \cdot C} . \quad (18.22)$$

⁵Die Schwankung Δx heißt in der Statistik auch Standardabweichung und das Schwankungsquadrat $(\Delta x)^2$ auch Varianz.

In diesem Zusammenhang ist $C = \Delta x / \langle x \rangle$ eine Konstante. (18.22) bezeichnet man als das **Gesetz der großen Zahlen**. Es bringt zum Ausdruck, dass die relative Schwankung eines idealen N -Teilchensystems mit wachsender Teilchenzahl N immer kleiner wird und im Grenzfall $N \rightarrow \infty$ gegen Null geht. Das liegt daran, dass der Erwartungswert $\langle X \rangle$ mit N schneller wächst als die Schwankung ΔX , welche nur mit $\sqrt{N}$ wächst.

Was das Gesetz der großen Zahlen in der statistischen Physik bzw. der Thermodynamik bedeutet, wollen wir an einem einfachen Modell illustrieren. Dazu betrachten wir ein ideales N -Teilchensystem, dessen Teilchen sich (unabhängig voneinander) entweder mit der Wahrscheinlichkeit p_1 im Quantenzustand 1 oder mit der Wahrscheinlichkeit p_2 im Quantenzustand 2 befinden.⁶ Die Anzahl der Teilchen des Systems, die sich im Quantenzustand 1 befinden, bezeichnen wir mit N_1 , sodass sich folglich $N - N_1 = N_2$ Teilchen im Quantenzustand 2 befinden. Ein derartiges N -Teilchensystem wird durch die **Binomialverteilung** (17.23)

$$P_N(N_1) = \underbrace{\frac{N!}{N_1! \cdot (N - N_1)!}}_{C_N(N_1)} \cdot p_1^{N_1} \cdot p_2^{N-N_1} \Rightarrow \quad (18.23)$$

$$\sum_{N_1=0}^N P_N(N_1) = \underbrace{(p_1 + p_2)}_{=1}^N = 1 \quad (18.24)$$

beschrieben.⁷ $C_N(N_1)$ ist hierbei die Anzahl von Teilchenkombinationen für ein bestimmtes N_1 und $P_N(N_1)$ deren Wahrscheinlichkeit (s. Kapitel 17).

Für die weitere Diskussion werden wir den Erwartungswert $\langle N_1 \rangle$ und die Schwankung ΔN_1 der Binomialverteilung benötigen:

$$\langle N_1 \rangle = \sum_{N_1=0}^N N_1 \cdot P_N(N_1) = N p_1 \quad , \quad (18.25)$$

denn

$$p_1 \frac{\partial}{\partial p_1} \left[\sum_{N_1=0}^N \frac{N!}{N_1! \cdot (N - N_1)!} p_1^{N_1} p_2^{N-N_1} \right] = p_1 \frac{\partial}{\partial p_1} \left[(p_1 + p_2)^N \right] , \quad (18.26)$$

$$\sum_{N_1=0}^N N_1 \cdot \frac{N!}{N_1! \cdot (N - N_1)!} p_1^{N_1} p_2^{N-N_1} = p_1 N (p_1 + p_2)^{N-1} , \quad (18.27)$$

$$\langle N_1 \rangle = N p_1 \quad \square \quad (18.28)$$

⁶Diese Quantenzustände entsprechen z.B. den Energieniveaus, die von den N Teilchen besetzt werden können.

⁷Es handelt sich bei diesem System um eine *Schar* aus N voneinander unabhängigen Teilchen. Jedes dieser Teilchen „schlüpft“ entweder mit der Wahrscheinlichkeit p_1 in den Quantenzustand 1 oder mit der Wahrscheinlichkeit p_2 in den Quantenzustand 2. Ein dazu analoges System im Sinne des *Zeitmittels* ist z.B. der eindimensionale Random Walk *eines* Teilchens. Dieser Random Walk besteht aus $N = n$ Schritten, wobei sich das Teilchen bei jedem Schritt mit p_1 um eine dem Wert des (Quanten)zustands 1 entsprechende Strecke vor- oder mit p_2 um eine dem Wert des (Quanten)zustands 2 entsprechende Strecke zurückbewegt.

und

$$\boxed{\Delta N_1 = \sqrt{N p_1 p_2}} \quad , \quad (18.29)$$

denn

$$p_1 \frac{\partial}{\partial p_1} \left\{ p_1 \frac{\partial}{\partial p_1} \sum_{N_1=0}^N P_N(N_1) \right\} = p_1 \frac{\partial}{\partial p_1} \left\{ p_1 \frac{\partial}{\partial p_1} [(p_1 + p_2)^N] \right\} \quad , \quad (18.30)$$

$$\underbrace{\sum_{N_1=0}^N N_1^2 \cdot P_N(N_1)}_{= \langle N_1^2 \rangle} = p_1 N + p_1^2 N(N-1) \quad , \quad (18.31)$$

woraus mit $1 - p_1 = p_2$

$$\boxed{\langle N_1^2 \rangle = N p_1 p_2 + \langle N_1 \rangle^2 = N p_1 p_2 + N^2 p_1^2} \quad (18.32)$$

folgt, sodass

$$\Delta N_1 = \sqrt{\langle N_1^2 \rangle - \langle N_1 \rangle^2} = \sqrt{(p_1 N + p_1^2 N^2 - p_1^2 N) - p_1^2 N^2} = \sqrt{N p_1 p_2} \quad . \quad \square \quad (18.33)$$

Auf analoge Weise erhält man aus

$$P_N(N_2) = \frac{N!}{N_2! \cdot (N - N_2)!} \cdot p_1^{N-N_2} \cdot p_2^{N_2} \quad (18.34)$$

schließlich

$$\langle N_2 \rangle = N p_2 \quad , \quad (18.35)$$

$$\langle N_2^2 \rangle = N p_1 p_2 + \langle N_2 \rangle^2 = N p_1 p_2 + N^2 p_2^2 \quad , \quad (18.36)$$

$$\Delta N_2 = \sqrt{N p_1 p_2} \quad . \quad (18.37)$$

Die relative Schwankung von N_1 in der Binomialverteilung ist folglich

$$\frac{\Delta N_1}{\langle N_1 \rangle} = \frac{\sqrt{N p_1 p_2}}{N p_1} = \frac{1}{\sqrt{N}} \cdot \sqrt{\frac{p_2}{p_1}} = \frac{1}{\sqrt{N}} \cdot A \quad . \quad (18.38)$$

In diesem Zusammenhang ist $A = \sqrt{p_2/p_1}$ eine Konstante.

Bisher haben wir lediglich die Binomialverteilung der N Teilchen auf zwei Zustände gemäß deren Besetzungswahrscheinlichkeiten p_1 und p_2 dargestellt. Jetzt ordnen wir dem Zustand $i = 1$ den Wert $x_i = x_1$ und dem Zustand $i = 2$ den Wert $x_i = x_2$ zu und untersuchen, welche Werte $X = \sum_{i=1}^2 N_i \cdot x_i$ mit welcher Wahrscheinlichkeit das N -Teilchensystem infolge der Binomialverteilung besitzen kann.

N_i ist die Teilchenzahl im Zustand i und kann die Zahlen $0, 1, 2, 3, \dots, N$ annehmen. So können sich in einem System aus $N = 20$ Teilchen z.B. $N_1 = 16$ Teilchen im Zustand 1 und $N_2 = 4$ Teilchen im Zustand 2 befinden. Es sind alle Kombinationen

aus N_1 und N_2 erlaubt, die die Bedingung $N_1 + N_2 = N$ erfüllen. Zur Vereinfachung, aber auch wegen der Analogie zum eindimensionalen Random Walk, wählen wir für den (Quanten)zustand 1 den Wert

$$x_1 = +1$$

mit der Besetzungswahrscheinlichkeit p_1 und für den (Quanten)zustand 2 den Wert

$$x_2 = -1$$

mit der Besetzungswahrscheinlichkeit p_2 . Damit können wir jetzt konkrete X -Werte unseres Systems und deren Wahrscheinlichkeit $P_N(N_1)$ berechnen, wenn wir p_1 und p_2 kennen. Die N_1 Teilchen im Zustand 1 liefern $N_1 \cdot (+1)$ und die N_2 Teilchen im Zustand 2 liefern $N_2 \cdot (-1)$, sodass

$$X = N_1 \cdot (+1) + N_2 \cdot (-1) = N_1 - N_2 = 2N_1 - N . \quad (18.39)$$

Mit dem Erwartungswert $\langle N_1 \rangle = Np_1$ erhalten wir sofort den Erwartungswert für X gemäß⁸

$$\langle X \rangle = \langle 2N_1 - N \rangle = 2\langle N_1 \rangle - N , \quad (18.40)$$

$$\langle X \rangle = 2Np_1 - N = N(p_1 - p_2) \quad (18.41)$$

und mit dem Erwartungswert $\langle N_1^2 \rangle = Np_1 p_2 + N^2 p_1^2$ auch den Erwartungswert von X^2 gemäß

$$\langle X^2 \rangle = \langle (2N_1 - N)^2 \rangle = \langle 4N_1^2 - 4N N_1 + N^2 \rangle \quad (18.42)$$

$$= 4\langle N_1^2 \rangle - 4N\langle N_1 \rangle + N^2 \quad (18.43)$$

$$= 4Np_1 p_2 + 4N^2 p_1^2 - 4N^2 p_1 + N^2 , \quad (18.44)$$

$$\langle X^2 \rangle = 4Np_1 p_2 + (2Np_1 - N)^2 . \quad (18.45)$$

Damit ist die Schwankung von X

$$\Delta X = \sqrt{\langle X^2 \rangle - \langle X \rangle^2} = 2\sqrt{Np_1 p_2} . \quad (18.46)$$

Die relative Schwankung der Systemgesamtwerte X ist folglich

$$\frac{\Delta X}{\langle X \rangle} = \frac{2\sqrt{Np_1 p_2}}{N(p_1 - p_2)} = \frac{1}{\sqrt{N}} \cdot \frac{2\sqrt{p_1 p_2}}{p_1 - p_2} = \frac{1}{\sqrt{N}} \cdot B . \quad (18.47)$$

In diesem Zusammenhang ist $B = \frac{2\sqrt{p_1 p_2}}{p_1 - p_2}$ eine Konstante.

⁸Man erhält $\langle X \rangle$, auch wenn man von $\langle X \rangle = \langle N_1 - N_2 \rangle$ ausgeht und die Erwartungswerte $\langle N_1 \rangle$ und $\langle N_2 \rangle$ einsetzt. Bei der Bestimmung von $\langle X^2 \rangle$ muss man jedoch entscheiden, ob man ausgeht von X in Abhängigkeit von N_2 oder wie in unserem Fall ausgeht von $X = 2N_1 - N$, also von X in Abhängigkeit von N_1 . Verwendet man nämlich $\langle X^2 \rangle = \langle (N_1 - N_2)^2 \rangle$ und setzt sowohl $\langle N_1 \rangle$ als auch $\langle N_2 \rangle$ ein, resultiert das falsche Ergebnis $\sqrt{2Np_1 p_2}$.

Wir stellen fest:

Die Gleichungen (18.38) und (18.47) besitzen die gleiche Gestalt wie (18.22) und repräsentieren deshalb ebenfalls das Gesetz der großen Zahl.

Weil der Gleichung (18.47), aber auch bereits der Gleichung (18.39) die Binomialverteilung der N_i zugrunde liegt, entspricht die Verteilungsfunktion der Systemgesamtwerte X ebenfalls der Binomialverteilung, sodass beide Verteilungsfunktionen prinzipiell die gleichen Eigenschaften besitzen. Vereinfachend dürfen wir uns deshalb im Folgenden allein auf die Untersuchung der Eigenschaften der Binomialverteilung von N_1 für große Teilchenzahlen N beschränken. Die dabei gewonnenen Erkenntnisse lassen sich dann auf die Verteilungsfunktion der X übertragen.

Die Binomialverteilung ist bereits normiert gemäß

$$\sum_{N_1=0}^N P_N(N_1) = \sum_{N_1=0}^N \frac{N!}{N_1! \cdot (N - N_1)!} p_1^{N_1} p_2^{N-N_1} = (p_1 + p_2)^N = 1 . \quad (18.48)$$

Im Folgenden werden wir ausnutzen, dass das Maximum der Binomialverteilung an derselben Stelle liegt wie das Maximum des Logarithmus der Binomialverteilung. Aus Gründen der Bequemlichkeit und weil $P_N(N_1)$ sehr empfindlich von N_1 abhängt, gehen wir zwischenzeitlich über zur logarithmischen Darstellung der Binomialverteilung:

$$\ln P_N(N_1) = \ln N! - \ln N_1! - \ln(N - N_1)! + N_1 \cdot \ln p_1 + (N - N_1) \cdot \ln p_2 . \quad (18.49)$$

Für große N dürfen wir die Fakultäten mit Hilfe der Stirling-Formel $N! \approx N^N / e^N$ nähern, sodass

$$\begin{aligned} \ln P_N(N_1) \approx \\ N \ln N - N_1 \ln N_1 - (N - N_1) \ln(N - N_1) + N_1 \ln p_1 + (N - N_1) \ln p_2 . \end{aligned} \quad (18.50)$$

Zur Bestimmung des Maximums von $\ln P_N(N_1)$ an der Stelle $\hat{N}_1$ nehmen wir näherungsweise an, dass N_1 eine kontinuierliche Variable ist, sodass

$$\left. \frac{d \ln P_N(N_1)}{d N_1} \right|_{\hat{N}_1} \stackrel{!}{=} 0 = -\ln \hat{N}_1 + \ln(N - \hat{N}_1) + \ln p_1 - \ln p_2 , \quad (18.51)$$

$$\left. \frac{d^2 \ln P_N(N_1)}{d N_1^2} \right|_{\hat{N}_1} = -\frac{1}{\hat{N}_1} - \frac{1}{N - \hat{N}_1} < 0 \Rightarrow \text{Maximum} . \quad (18.52)$$

Aus (18.51) resultiert

$$\ln \frac{\hat{N}_1}{N - \hat{N}_1} = \ln \frac{p_1}{p_2} = \ln \frac{p_1}{1 - p_1} , \quad (18.53)$$

$$\Leftrightarrow \hat{N}_1(1 - p_1) = (N - \hat{N}_1)p_1 \quad (18.54)$$

$$\hat{N}_1 - \hat{N}_1 p_1 = N p_1 - \hat{N}_1 p_1 , \quad (18.55)$$

$$\hat{N}_1 = N p_1 \quad (18.56)$$

und mit $N p_1 = \langle N_1 \rangle$

$$\boxed{\hat{N}_1 = \langle N_1 \rangle} . \quad (18.57)$$

Wir stellen fest:

Bei großer Teilchenzahl N liegt das Maximum der Binomialverteilung an der Stelle des Erwartungswerts von N_1 . Anders gesagt, die wahrscheinlichste Teilchenzahl $N_1 = \hat{N}_1$ ist auch der Erwartungswert der Teilchenzahl N_1 , nämlich $\langle N_1 \rangle$.

Schließlich kommen wir zur Diskussion der physikalisch gesehen wohl wichtigsten Eigenschaft der Binomialverteilung, zu ihrem Verhalten im Bereich ihres Maximums, unter den Voraussetzungen $N \gg 1$ und $Np_1 p_2 \gg 1$. Zu diesem Zweck führen wir einige (mühsame!) Umformungen und Näherungen durch:

Die Taylor-Entwicklung von $\ln P_N(N_1)$ an der Stelle $\hat{N}_1$ (des Maximums) bis zur zweiten Ordnung ist

$$\begin{aligned} \ln P_N(N_1) \approx & \ln P_N(\hat{N}_1) + \left[\frac{d \ln P_N(N_1)}{dN_1} \right]_{\hat{N}_1} (N_1 - \hat{N}_1) \\ & + \frac{1}{2} \left[\frac{d^2 \ln P_N(N_1)}{dN_1^2} \right]_{\hat{N}_1} (N_1 - \hat{N}_1)^2. \end{aligned} \quad (18.58)$$

Gemäß (18.51) verschwindet der Term erster Ordnung an der Stelle des Maximums und wir erhalten mit (18.52)

$$\ln P_N(N_1) \approx \ln P_N(\hat{N}_1) + \frac{1}{2} \left(-\frac{1}{\hat{N}_1} - \frac{1}{N - \hat{N}_1} \right) (N_1 - \hat{N}_1)^2 \quad (18.59)$$

$$\approx \ln P_N(\hat{N}_1) - \frac{1}{2} \left(\frac{1}{Np_1} + \frac{1}{Np_2} \right) (N_1 - \hat{N}_1)^2 \quad (18.60)$$

$$\ln P_N(N_1) \approx \ln P_N(\hat{N}_1) - \frac{1}{2} \frac{(N_1 - \hat{N}_1)^2}{Np_1 p_2}. \quad (18.61)$$

Nach (18.29) ist $Np_1 p_2 = (\Delta N_1)^2$, sodass⁹

$$\ln P_N(N_1) = \ln P_N(\hat{N}_1) - \frac{(N_1 - \hat{N}_1)^2}{2(\Delta N_1)^2}. \quad (18.62)$$

Wir verlassen jetzt die logarithmische Darstellung von $P_N(N_1)$ und erhalten so aus (18.62) schließlich

$$P_N(N_1) = P_N(\hat{N}_1) \cdot \exp \left(-\frac{(N_1 - \hat{N}_1)^2}{2(\Delta N_1)^2} \right). \quad (18.63)$$

Wie man sieht, hat (18.63) große Ähnlichkeit mit der Wahrscheinlichkeitsdichtefunktion w der Gauß'schen Normalverteilung

$$w(x, x_0, \sigma) = \frac{1}{\sigma \cdot \sqrt{2\pi}} \cdot e^{-\frac{(x-x_0)^2}{2\sigma^2}} \Rightarrow P = \int_{-\infty}^{+\infty} w(x, x_0, \sigma) dx = 1. \quad (18.64)$$

Um die Wahrscheinlichkeit für das betrachtete N -Teilchensystem mit N_1 Teilchen im Zustand 1 zu zeigen, greifen wir auf die diskrete Betrachtungsweise der Wahrscheinlichkeitsdichte $w(x_i)$ für das Auftreten der Werte x_i zurück:

$$w(x_i) = \frac{1}{N_{ges}} \cdot \frac{\Delta N_i(x_i)}{\Delta x_i} \Rightarrow$$

⁹Wenn auch nicht ganz korrekt, verwenden wir im Folgenden wieder das Gleichheitszeichen.

$$P = \sum_i \underbrace{\frac{1}{N_{ges}} \cdot \frac{\Delta N_i(x_i)}{\Delta x_i}}_{=w(x_i)} \cdot \Delta x_i = \frac{1}{N_{ges}} \underbrace{\sum_i \Delta N_i(x_i)}_{=N_{ges}} = 1 .$$

Hierbei ist $\Delta N_i(x_i)$ die Anzahl dafür, wie oft der Wert x_i der Zufallsgröße X auftritt, und $w(x_i) \cdot \Delta x_i$ ist die Wahrscheinlichkeit für das Auftreten des Wertes x_i . Mit den Umbenennungen

$$N_{ges} \rightarrow N, \quad \Delta N_i \rightarrow \Delta N, \quad x_i \rightarrow N_1, \quad \Delta x_i \rightarrow \Delta(N_1)$$

folgt daraus:

$$P(x_i) = w(x_i) \cdot \Delta x_i = \frac{1}{N_{ges}} \cdot \frac{\Delta N_i(x_i)}{\Delta x_i} \cdot \Delta x_i \quad (18.65)$$

$$\Rightarrow P_N(N_1) = w(N_1) \cdot \underbrace{\Delta(N_1)}_{\equiv 1} = \frac{1}{N} \cdot \frac{\Delta N(N_1)}{\Delta(N_1)} \cdot \underbrace{\Delta(N_1)}_{\equiv 1} . \quad (18.66)$$

Hier ist $\Delta(N_1)$ nicht die Schwankung ΔN_1 von N_1 sondern die Differenz zwischen zwei benachbarten Teilchenzahlen N_1 .

Wegen $\Delta(N_1) \equiv 1$ können wir mit (18.63) und (18.66) mühelos die (diskrete) Wahrscheinlichkeit $P_N(N_1)$ in die kontinuierliche Wahrscheinlichkeitsdichte $w_N(N_1)$ überführen, indem wir für große N einfach N_1 als kontinuierliche Variable betrachten und unter Verwendung der Schwankung ΔN_1 schreiben:

$$w_N(N_1) = \frac{1}{N} \cdot \frac{dN(N_1)}{dN_1} = \widetilde{P}_N(\widehat{N}_1) \cdot \exp \left(- \frac{(N_1 - \widehat{N}_1)^2}{2(\Delta N_1)^2} \right) . \quad (18.67)$$

Dabei ist aus $P_N(\widehat{N}_1)$ der Normierungsfaktor $\widetilde{P}_N(\widehat{N}_1)$ geworden. Mit der Substitution

$$N_1 - \widehat{N}_1 = x \quad \Rightarrow \quad dN_1 = dx ,$$

den Integrationsgrenzen¹⁰

$$N \rightarrow \infty \quad \Rightarrow \quad \begin{cases} N_1 \rightarrow \infty & \Rightarrow N_1 \gg \widehat{N}_1 & \Rightarrow N_1 - \widehat{N}_1 = x \rightarrow +\infty \\ N_1 = 0 & \Rightarrow N_1 \ll \widehat{N}_1 & \Rightarrow N_1 - \widehat{N}_1 = x \rightarrow -\infty \end{cases}$$

und dem Standardintegral

$$\int_{-\infty}^{+\infty} dx e^{-\alpha x^2} = \sqrt{\frac{\pi}{\alpha}} \quad (18.68)$$

bestimmen wir den Normierungsfaktor $\widetilde{P}_N(\widehat{N}_1)$ in (18.67) wie folgt:

$$\begin{aligned} \int_{-\infty}^{+\infty} dx w_N(x) &= \int_{-\infty}^{+\infty} dx \widetilde{P}_N(\widehat{N}_1) \cdot \exp \left(- \frac{x^2}{2(\Delta N_1)^2} \right) \\ &= \widetilde{P}_N(\widehat{N}_1) \cdot \int_{-\infty}^{+\infty} dx \exp \left(- \frac{x^2}{2(\Delta N_1)^2} \right) \stackrel{!}{=} 1 , \end{aligned} \quad (18.69)$$

¹⁰Für $(N_1 - \widehat{N}_1) = x \rightarrow \infty$ verschwindet der Integrand und für $x \ll -1$ wird er vernachlässigbar klein, sodass man in sehr guter Näherung für den Fall $N_1 = 0$ wegen $\widehat{N}_1 \gg 1$ als untere Integrationsgrenze $(N_1 - \widehat{N}_1) = -\widehat{N}_1 = x \rightarrow -\infty$ nutzen kann.

$$\widetilde{P}_N(\widehat{N}_1) \cdot \sqrt{\pi \cdot 2(\Delta N_1)^2} \stackrel{!}{=} 1 \quad \Leftrightarrow \quad \widetilde{P}_N(\widehat{N}_1) = \frac{1}{\Delta N_1 \cdot \sqrt{2\pi}} . \quad (18.70)$$

Die Wahrscheinlichkeitsdichtefunktion in ihrer endgültigen Form ist somit

$$\boxed{w_N(N_1) = \frac{1}{\Delta N_1 \cdot \sqrt{2\pi}} \cdot \exp\left(-\frac{(N_1 - \widehat{N}_1)^2}{2(\Delta N_1)^2}\right), \quad N \gg 1} . \quad (18.71)$$

Der Vergleich mit (18.64) zeigt:

$N_1 = x$ ist die unabhängige Variable,

$\widehat{N}_1 = x_0$ ist die Stelle des Maximums und

$\Delta N_1 = \sigma$ ist die Schwankung bzw. Standardabweichung der (kontinuierlichen) Wahrscheinlichkeitsdichtefunktion¹¹ w .

Wir stellen fest:

Für große N lässt sich die Binomialverteilung in der Umgebung ihres Maximums in guter Näherung durch die Gauß'sche Normalverteilung darstellen und im Grenzfall $N \rightarrow \infty$ geht die Binomialverteilung über in die Gauß'sche Normalverteilung.

Die weitere Diskussion und die Beispielrechnung gestalten sich einfacher, wenn wir in (18.71) ΔN_1 durch $\sqrt{N p_1 p_2}$ ersetzen, sodass

$$w_N(N_1) = \frac{1}{\sqrt{2\pi \cdot N p_1 p_2}} \cdot \exp\left(-\frac{(N_1 - \widehat{N}_1)^2}{2N p_1 p_2}\right) . \quad (18.72)$$

Daraus folgt

$$w_N(\widehat{N}_1) = \frac{1}{\sqrt{2\pi \cdot N p_1 p_2}} . \quad (18.73)$$

¹¹In der diskreten Betrachtungsweise spricht man nicht von Dichtefunktion sondern von Verteilung.

Zur Veranschaulichung betrachten wir als einfache Beispiele die Binomialverteilungen mit den Wahrscheinlichkeiten $p_1 = 0,2$ und $p_2 = 0,8$ für $N = 3$, $N = 6$ und $N = 12$:

$N = 3$	Kombinationen	$C_N(N_1)$	$P_N(N_1) \approx$	X
$N_1 = 0$	-1 -1 -1	1	0,51	-3
$N_1 = 1$	-1 -1 1 -1 1 -1 1 -1 -1	3	0,38	-1
$N_1 = 2$	-1 1 1 1 -1 1 1 1 -1	3	0,10	1
$N_1 = 3$	1 1 1	1	0,01	3
$\langle N_1 \rangle = 0,6$ $\Delta N_1 \approx 0,693$		$\sum_{N_1=0}^3 C_N(N_1) = 8$	$\sum_{N_1=0}^3 P_N(N_1) = 1$	$\langle X \rangle = -1,8$ $\Delta X \approx 1,386$

Tabelle 18.1 Binomialverteilung in einem idealen 3-Teilchensystem mit den Wahrscheinlichkeiten $p_1 = 0,2$ für $x_1 = +1$ und $p_2 = 0,8$ für $x_2 = -1$. Die relativen Schwankungen betragen

$$\frac{\Delta N_1}{\langle N_1 \rangle} \approx 1,155 \quad , \quad \left| \frac{\Delta X}{\langle X \rangle} \right| \approx 0,770 \quad .$$

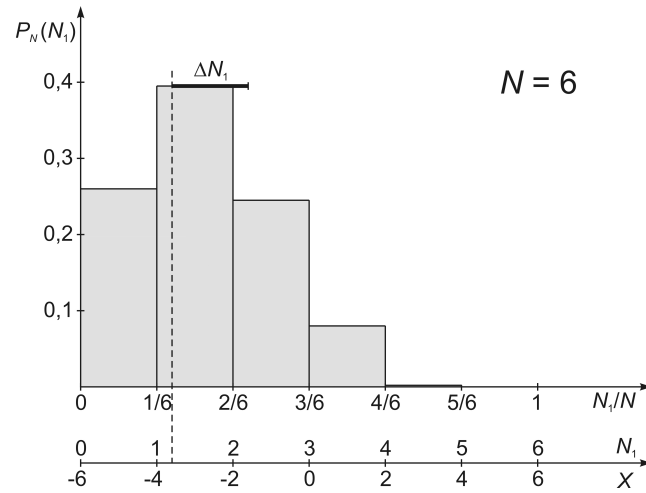


Abb. 18.1 Histogramm der Binomialverteilung für $N = 6$, $p_1 = 0,2$ und $p_2 = 0,8$. Die Wahrscheinlichkeiten $P_N(N_1)$ sind aufgetragen über der gemäß $(\sum_{N_1=0}^N N_1)/N = N/N = 1$ auf 1 normierten Achse N_1/N , aber auch über N_1 und X . Die gestrichelte Linie markiert die Stelle der Erwartungswerte $\langle N_1 \rangle = 1,2$ und $\langle X \rangle = -3,6$. Die Summe der Zellenflächeninhalte repräsentiert die Gesamtwahrscheinlichkeit $\sum_{N_1=0}^N P_N(N_1) = 1$ und ist ebenfalls auf 1 normiert.

$N = 6$	$C_N(N_1)$	$P_N(N_1) \approx$	X
$N_1 = 0$	1	0,262	-6
$N_1 = 1$	6	0,393	-4
$N_1 = 2$	15	0,246	-2
$N_1 = 3$	20	0,082	0
$N_1 = 4$	15	0,015	2
$N_1 = 5$	6	0,002	4
$N_1 = 6$	1	0,000	6
$\langle N_1 \rangle = 1,2$	$\sum_{N_1=0}^6 C_N(N_1) = 64$	$\sum_{N_1=0}^6 P_N(N_1) = 1$	$\langle X \rangle = -3,6$
$\Delta N_1 \approx 0,980$			$\Delta X \approx 1,960$

Tabelle 18.2 Binomialverteilung in einem idealen 6-Teilchensystem mit den Wahrscheinlichkeiten $p_1 = 0,2$ für $x_1 = +1$ und $p_2 = 0,8$ für $x_2 = -1$. Die relativen Schwankungen betragen

$$\frac{\Delta N_1}{\langle N_1 \rangle} \approx 0,817 \quad , \quad \left| \frac{\Delta X}{\langle X \rangle} \right| \approx 0,544 \quad .$$

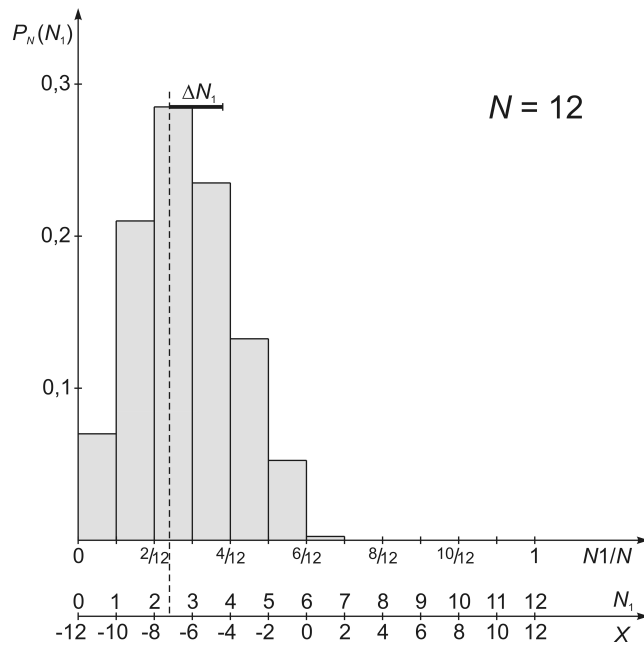


Abb. 18.2 Histogramm der Binomialverteilung für $N = 12$, $p_1 = 0,2$ und $p_2 = 0,8$.

$N = 12$	$C_N(N_1)$	$P_N(N_1) \approx$	X
$N_1 = 0$	1	0,069	-12
$N_1 = 1$	12	0,206	-10
$N_1 = 2$	66	0,284	-8
$N_1 = 3$	220	0,236	-6
$N_1 = 4$	495	0,133	-4
$N_1 = 5$	792	0,053	-2
$N_1 = 6$	924	0,016	0
$N_1 = 7$	792	0,003	2
$N_1 = 8$	495	0,000	4
$N_1 = 9$	220	0,000	6
$N_1 = 10$	66	0,000	8
$N_1 = 11$	12	0,000	10
$N_1 = 12$	1	0,000	12
$\langle N_1 \rangle = 2,4$	$\sum_{N_1=0}^{12} C_N(N_1) = 4096$	$\sum_{N_1=0}^{12} P_N(N_1) = 1$	$\langle X \rangle = -7,2$
$\Delta N_1 \approx 1,386$			$\Delta X \approx 2,771$

Tabelle 18.3 Binomialverteilung in einem idealen 12-Teilchensystem mit den Wahrscheinlichkeiten $p_1 = 0,2$ für $x_1 = +1$ und $p_2 = 0,8$ für $x_2 = -1$. Die relativen Schwankungen betragen

$$\frac{\Delta N_1}{\langle N_1 \rangle} \approx 0,577 \quad , \quad \left| \frac{\Delta X}{\langle X \rangle} \right| \approx 0,385 \quad .$$

Wie wir sehen, wachsen die Erwartungswerte $\langle N_1 \rangle$ bzw. $\langle X \rangle$ proportional zur System-Teilchenzahl N an, die Schwankungen ΔN_1 bzw. ΔX aber nur proportional zu $\sqrt{N}$. Folglich werden die relativen Schwankungen $\Delta N_1 / \langle N_1 \rangle$ bzw. $|\Delta X / \langle X \rangle|$ mit wachsendem N immer kleiner. Um diesen Sachverhalt zu verdeutlichen, wurden die Histogramme in entsprechender Weise skaliert:

Der Wertebereich N_1 der Binomialverteilung läuft von 0 bis N . Wir skalierten ihn aber entsprechend N_1/N , sodass er, unabhängig von der Systemteilchenzahl N , für alle Systeme von 0 bis 1 läuft. Dadurch werden die Histogrammsäulen um den Faktor $1/N$ schmaler. Folglich liegen die Erwartungswerte $\langle N_1 \rangle$ auf der $\frac{N_1}{N}$ -Achse immer an der Stelle $\langle N_1 \rangle / N = N p_1 / N = p_1$. Gemäß der Forderung, dass der Flächeninhalt unterhalb der „Kurve“ des Histogramms stets gleich 1 sei, mussten die Histogrammsäulen im Gegenzug um den Faktor N höher werden. Die zu einem bestimmten Wahrscheinlichkeitswert $P_N(N_1)$ gehörende Strecke auf der Ordinate verlängert sich also durch die neue Skalierung der Histogramme um den Faktor N . Die relative Schwankung $\Delta N_1 / \langle N_1 \rangle$ wird in den Histogrammen (s. Abb. 18.1 und Abb. 18.2) repräsentiert durch die *graphische Länge* der dort eingezeichneten Strecken ΔN_1 .

Beim Vergleich unserer drei Beispiele wird trotz der kleinen Teilchenzahlunterschiede zwischen $N = 3, 6$ und 12 deutlich, dass die Schwankung ΔN_1 zwar mit wachsender Teilchenzahl anwächst, die relative Schwankung der Binomialverteilung dabei aber rasch abnimmt.

Wir schlussfolgern:

Die relative Schwankung der Binomialverteilung geht im Grenzfall $N \rightarrow \infty$ gegen Null. Das bedeutet, dass bei Systemen großer Teilchenzahl N die Werte N_1 bzw. die Systemgesamtwerte (Observablen) X mit zunehmendem Abstand vom Erwartungswert $\langle N_1 \rangle$ bzw. $\langle X \rangle$ sehr schnell an Wahrscheinlichkeit für ihre Realisierung verlieren. Im Grenzfall $N \rightarrow \infty$ liegt das System allein in jenen Teilchenkombinationen vor, welche den Erwartungswert $\langle N_1 \rangle$ liefern, sodass nur noch die Observable $\langle X \rangle$ in Erscheinung tritt bzw. gemessen werden kann.

Die Schwankung ΔN_1 wird allgemein als das Maß für die *Breite* der Gauß'schen Wahrscheinlichkeitsdichtefunktion verwendet, die *relative Schwankung* $\Delta N_1 / \langle N_1 \rangle$ sinngemäß auch als Maß für die relative Breite.

Analog dazu und übereinstimmend mit der normierten Darstellung der Binomialverteilung in unseren Histogrammen ist die bezüglich N „relativierte“ Größe $\Delta N_1 / N$ das Maß für die *relative Breite* der Gauß'schen Normalverteilung. Verdeutlichen wir uns diesen Sachverhalt abschließend durch ein kleines Rechenbeispiel¹²:

Gegeben:

$$N = 6,25 \cdot 10^{22}, \quad p_1 = 0,2, \quad p_2 = 0,8, \quad \hat{N}_1 = \langle N_1 \rangle = N p_1 = 1,25 \cdot 10^{22}.$$

Für die Schwankung (Standardabweichung) von N_1 ermitteln wir

$$\Delta N_1 = \sqrt{N p_1 p_2} = 10^{11}.$$

Das bedeutet, dass ca. $\frac{2}{3}$ aller $\sum_{N_1=0}^N C_N(N_1)$ möglichen Kombinationen der Binomialverteilung im Bereich $(\hat{N}_1 - \Delta N_1) \leq \hat{N}_1 \leq (\hat{N}_1 + \Delta N_1)$ liegen. Diese Schwankung

¹²Vergleiche mit Wolfgang Nolting, Grundkurs Theoretische Physik 6, Statistische Physik, 5. Aufl., Springer-Verlag, Berlin, Heidelberg, 2005, Seite 11 und Seite 12.

der Teilchenzahl N_1 von $\pm 10^{11}$ um ihren Erwartungswert ist im Verhältnis zur Größe des Erwartungswerts von $1,25 \cdot 10^{22}$ gemäß der *relativen Schwankung*

$$\frac{\Delta N_1}{\langle N_1 \rangle} = \frac{\Delta N_1}{\hat{N}_1} = 8 \cdot 10^{-12}$$

sehr klein. In ähnlicher Weise ergibt das Verhältnis von Schwankung zur Systemteilchenzahl N die *relative Breite* der Verteilung:

$$\frac{\Delta N_1}{N} = 1,6 \cdot 10^{-12}.$$

Auch für ein makroskopisches System ist die Schwankung $\Delta N_1 = 10^{11}$ eine große Zahl. Und dennoch ist sie im Verhältnis zum Definitionsbereich von N_1 , nämlich $0 \leq N_1 \leq N = 6,25 \cdot 10^{22}$, verschwindend klein.

Abschließend liefert uns das Rechenbeispiel einige in Tabelle 18.4 aufgeführte Wahrscheinlichkeitsdichtewerte $w_N(N_1)$ mit ihrem Verhältnis zur maximalen Wahrscheinlichkeitsdichte $w_N(\hat{N}_1)$ für zunehmende Abstände vom Maximum $\hat{N}_1$. Die Werte zeigen eindrucksvoll, wie die Wahrscheinlichkeit des Auftretens von Kombinationen und Werten zu beiden Seiten des Bereichs $(\hat{N}_1 - \Delta N_1) \leq \hat{N}_1 \leq (\hat{N}_1 + \Delta N_1)$ extrem stark abfällt.¹³

$N_1 - \hat{N}_1$	$w_N(N_1)$	$\frac{w_N(N_1)}{w_N(\hat{N}_1)}$
0	$w_N(\hat{N}_1) \approx 3,9894 \cdot 10^{-12}$	= 1
$\sqrt{2} \cdot 10^6$	$\approx 3,9894 \cdot 10^{-12}$	$\approx 1,00$
$\sqrt{2} \cdot 10^9$	$\approx 3,9890 \cdot 10^{-12}$	$\approx 1,00$
$\sqrt{2} \cdot 10^{10}$	$\approx 3,9497 \cdot 10^{-12}$	$\approx 0,99$
$\Delta N_1 = 10^{11}$	$\approx 2,4197 \cdot 10^{-12}$	$\approx 0,61$
$\sqrt{2} \cdot 10^{11}$	$\approx 1,4676 \cdot 10^{-12}$	$\approx 0,37$
$\sqrt{2} \cdot 10^{12}$	$\approx 1,4841 \cdot 10^{-55}$	$\approx 3,72 \cdot 10^{-44}$
$\sqrt{2} \cdot 10^{13}$	≈ 0	≈ 0

Tabelle 18.4 Wahrscheinlichkeitsdichten $w_N(N_1)$ und ihr Verhältnis zur maximalen Wahrscheinlichkeitsdichte $w_N(\hat{N}_1)$ der Binomialverteilung eines N -Teilchensystems mit $N = 6,25 \cdot 10^{22}$, $p_1 = 0,2$ und $p_2 = 0,8$, dargestellt bezüglich des Abstands $N_1 - \hat{N}_1$ vom Maximum $\hat{N}_1$.

¹³Für diejenigen, die immer noch zweifeln, hier ein „hässlich hinkendes“ Gleichnis: In einem kleinen Heuhaufen sei eine Nadel leicht zu finden. Wenn aber der Heuhaufen ($\hat{= N$ oder $\hat{= \langle N_1 \rangle}$) schneller wächst als die Nadel ($\hat{= \Delta N_1}$), wird es immer schwieriger die Nadel zu finden, weil die wachsende Nadel im Verhältnis zum Heuhaufen immer kleiner wird ($\hat{= \frac{\Delta N_1}{N}}$ oder $\hat{= \frac{\Delta N_1}{\langle N_1 \rangle}}$).

Anmerkung

Man könnte bei der Herleitung des Gesetzes der großen Zahlen auch auf den Gedanken kommen, nicht von der *mittleren quadratischen Abweichung* (Standardabweichung, Schwankung)¹⁴

$$\Delta x = \sqrt{\sum_j p_j (x_j - \langle x \rangle)^2} = \sqrt{\sum_j p_j |x_j - \langle x \rangle|^2}, \quad (18.74)$$

sondern von der *mittleren absoluten Abweichung*

$$\widetilde{\Delta x} = \langle |x_j - \langle x \rangle| \rangle = \left\langle \sqrt{(x_j - \langle x \rangle)^2} \right\rangle = \sum_j p_j |x_j - \langle x \rangle| \quad (18.75)$$

auszugehen. Dabei wäre zu klären, ob $\widetilde{\Delta x}$ möglicherweise schneller mit N wächst bzw. größer ist als Δx , wodurch das Gesetz der großen Zahlen evl. in Frage gestellt würde. Aber tatsächlich gilt

$$\text{mittlere absolute Abweichung } \widetilde{\Delta x} \leq \Delta x \text{ Standardabweichung}, \quad (18.76)$$

$$\frac{\sum_{j=1}^N \sqrt{|x_j - \langle x \rangle|^2}}{N} \leq \sqrt{\frac{\sum_{j=1}^N (x_j - \langle x \rangle)^2}{N}}, \quad (18.77)$$

$$\frac{\sum_{j=1}^N |x_j - \langle x \rangle|}{N} \leq \sqrt{\frac{\sum_{j=1}^N |x_j - \langle x \rangle|^2}{N}}. \quad (18.78)$$

Wir zeigen¹⁵ dies für $N = 2$, also für die zwei Werte

$$x_1 \Rightarrow |x_1 - \langle x \rangle| = a \quad \text{und} \quad x_2 \Rightarrow |x_2 - \langle x \rangle| = b$$

in Analogie zu (18.78):

$$\text{arithmetisches Mittel} \quad \frac{a+b}{2} \leq \sqrt{\frac{a^2+b^2}{2}} \quad \text{quadratisches Mittel}, \quad (18.79)$$

Quadrieren ergibt

$$\frac{a^2 + b^2 + 2ab}{4} \leq \frac{a^2 + b^2}{2}, \quad (18.80)$$

$$2ab \leq a^2 + b^2, \quad (18.81)$$

$$2ab = a^2 + b^2 \quad \text{für } a = b, \quad (18.82)$$

$$2ab = 2a(a+x) < a^2 + b^2 = a^2 + (a+x)^2 \quad \text{für } a \neq b = (a+x). \quad (18.83)$$

Man kann dies nach Belieben auch für $N > 2$ zeigen.

¹⁴Mit der Standardabweichung finden große Abweichungen vom Erwartungswert überproportional Berücksichtigung.

¹⁵Hinweise zur arithmetisch-quadratischen Ungleichung und zur entsprechenden Beweisführung sind u. a. zu finden bei Jan Pöschko, Ungleichungen, 2009, www.math.tugraz.at/OeMO/A-Kurs/Unterlagen/Ungleichungen

19 Kapitalwachstum bei stetiger Reinvestition der Zinsen

$K_0 = a_0 + b_0 + c_0 + \dots$ sei das Startkapital.

ζ sei der beispielsweise jährliche Zinssatz, z. B. $\zeta = 5\% = 0,05$.

Dann beträgt das Kapital nach 1 Jahr gemäß $x = 1$

$$K_1 = K_0 + \zeta \cdot K_0 ,$$

$$K_1 = K_0 (1 + \zeta) .$$

Nach 2 Jahren gemäß $x = 2$ beträgt das Kapital

$$K_2 = K_1 + \zeta \cdot K_1 = K_0 (1 + \zeta) + \zeta \cdot [K_0 (1 + \zeta)]$$

Ausklammern von $K_0 (1 + \zeta)$

$$= K_0 (1 + \zeta) (1 + \zeta) ,$$

$$K_2 = K_0 (1 + \zeta)^2 .$$

Nach 3 Jahren gemäß $x = 3$ beträgt das Kapital

$$K_3 = K_2 + \zeta \cdot K_2 = K_0 (1 + \zeta)^2 + \zeta \cdot [K_0 (1 + \zeta)^2]$$

Ausklammern von $K_0 (1 + \zeta)^2$

$$= K_0 (1 + \zeta)^2 (1 + \zeta) ,$$

$$K_3 = K_0 (1 + \zeta)^3 .$$

Die Fortsetzung dieses (diskreten) Verfahrens liefert nach Verallgemeinerung auf die kontinuierliche Betrachtungsweise

$$\boxed{K_x = K_0 (1 + \zeta)^x , \quad \zeta \geq 0 , \quad x \in \mathbb{R} \wedge x \geq 0} .$$

Hierbei ist die Exponentialfunktion $f(x) = (1 + \zeta)^x \geq 1$ der Wachstumsfaktor (siehe Abb. 19.1). Wie man sieht, gibt es nur bei einem Zinssatz $\zeta > 0$ ein („positives“) Wachstum.

Mit $K_0 = a_0 + b_0 + c_0 + \dots$ gilt

$$K_x = (K_0)(1 + \zeta)^x$$

$$= (a_0 + b_0 + c_0 + \dots)(1 + \zeta)^x$$

$$= a_0 (1 + \zeta)^x + b_0 (1 + \zeta)^x + c_0 (1 + \zeta)^x + \dots ,$$

sodass es hinsichtlich des Gewinns gleichgültig ist, ob man K_0 als ganzes oder stückweise (zu den gleichen Konditionen) anlegt.

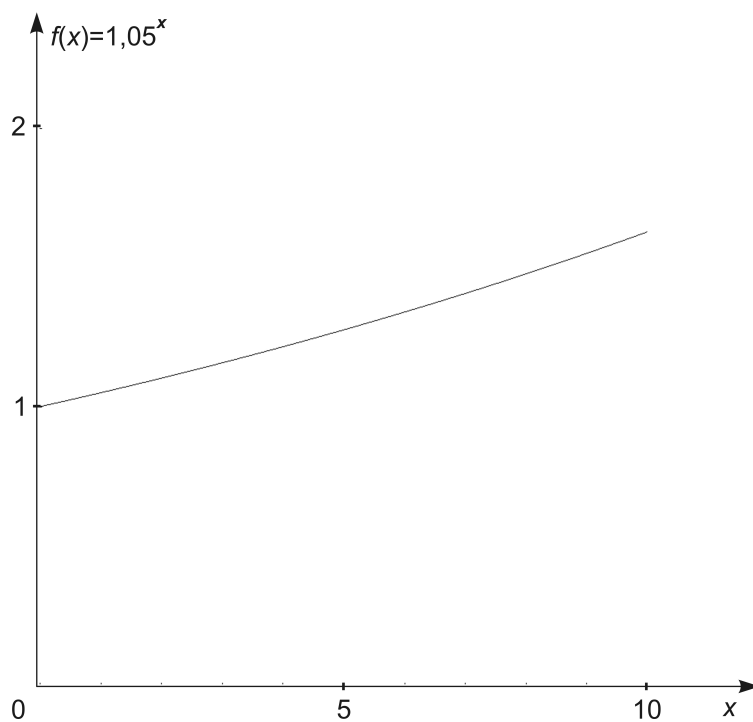
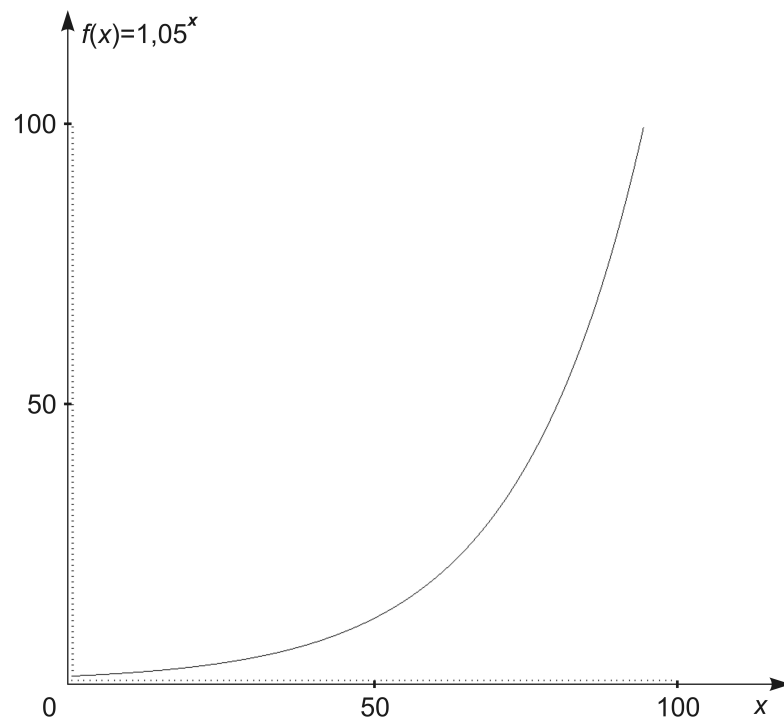


Abb. 19.1 Der Wachstumsfaktor $f(x) = (1 + \zeta)^x = (1 + 0,05)^x = 1,05^x$ für den jährlichen Zinssatz von 5 %, dargestellt über der Laufzeit x in Jahren. Bei der Laufzeit von 100 Jahren (oben) wird das exponentielle Wachstum deutlich sichtbar. Nach einer Laufzeit von 10 Jahren (unten) ist das Startkapital K_0 durch die Reinvestition der Zinsen um den Faktor 1,63 angewachsen.