

Skriptum Elemente der Geometrie 2015

T. von der Twer

Inhaltsverzeichnis

Kapitel 1. Elementare Vorkenntnisse	1
1. Grundlegende Bezeichnungen, Definitionen und Resultate	1
2. Sätze über Dreiecke im Umfeld der Pythagorasbeziehung	4
3. Konvexe Mengen	7
Kapitel 2. Zentralperspektive (Zentralprojektion)	11
1. Das Abbildungsprinzip der Zentralprojektion	11
2. Inzidenzerhaltung, Längenverzerrung, Fluchtpunkt von Geraden	13
3. Fluchtgeraden von Ebenen, Winkel und Winkelmesspunkte	14
4. Die analytische Methode	21
Kapitel 3. Grundzüge Analytischer Geometrie	25
1. Weiteres über Dreiecke: Umkreis, Inkreis, Schwerpunkt, Ankreise	25
2. Punkte, Ortsvektoren, freie Vektoren, Koordinatendarstellungen, Vektorräume	29
3. Vektorräume über $\mathbb{R}$ mit Skalarprodukt; Längen und Winkel	42
4. Lineare Abbildungen und ihre Matrixdarstellungen	49
Kapitel 4. Abbildungsgeometrie und Symmetrien	55
1. Begriff der Symmetrie und der Symmetriegruppe einer Punktmenge	55
2. Beispiele für Symmetriegruppen	60
3. Symmetrien von Bandmustern	64
4. Symmetrien von Pflasterungen der Ebene	66
Kapitel 5. Inzidenzgeometrie und Eulercharakteristik	71
1. Das Königsberger Brückenproblem und Verwandtes	74
2. Vollständige Graphen und planare Graphen, Eulercharakteristik	76
3. Elementares zu Flächeninhalten und Volumina	83
Kapitel 6. Sammlung wichtiger Sätze und Formeln	85

Elementare Vorkenntnisse

1. Grundlegende Bezeichnungen, Definitionen und Resultate

1.1. Konstruktionen mit Zirkel und Lineal. Wir stellen uns vor, dass wir nur ein Lineal und einen Zirkel zur Verfügung haben. Wir können damit nur Folgendes tun: 1. Irgendeine gerade Linie ziehen, insbesondere ein beliebiges Stück der Geraden durch zwei gegebene Punkte, 2. Die Strecke zwischen zwei Punkten in den Zirkel nehmen und um irgendeinen Punkt einen Kreis mit dem damit gegebenen Radius ziehen. Es folgt, dass Schnittpunkte eines Kreises mit einer Geraden bestimmt werden können (sofern es sie gibt).

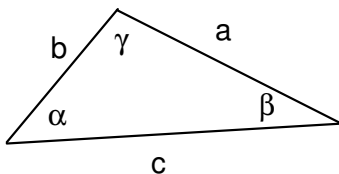
Es sei betont: Maßeinteilung und Winkelteilung eines Geodreiecks dürfen **nicht benutzt** werden!

Es zeigt sich, dass man allein mit Zirkel und Lineal Vieles konstruieren kann, aber keineswegs Alles. Ein Negativbeispiele: Man kann nicht (allein mit Zirkel und Lineal!) einen vorgegebenen Winkel in drei gleiche Winkel teilen. Man kann nicht zu einem Kreis ein flächengleiches Quadrat konstruieren (die berühmte 'Quadratur des Kreises'!) (Solche Resultate sind schwer zu gewinnen und gelangen erst E. Galois im 19. Jahrhundert mittels bahnbrechender algebraischer Resultate, für das zweite Beispiel benötigt man sogar noch, dass π transzendent ist, was erst am Ende des 19. Jahrhunderts bewiesen wurde.)

Wir erinnern uns an einige Positivbeispiele, wie sie aus der Schulmathematik bekannt sind:

- (1) Strecken und Winkel übertragen,
- (2) Strecken und Winkel halbieren, insbesondere die Mittelsenkrechte einer Strecke konstruieren, allgemeiner eine Strecke in n gleiche Teile teilen,
- (3) rechte Winkel konstruieren,
- (4) zu einem Punkt und einer Geraden g mit $P \notin g$ die parallele Gerade h zu g mit $P \in h$ konstruieren,
- (5) Dreiecke mit passenden (d.h. konsistenten und hinreichenden) Daten konstruieren, z.B. mit drei Seitenlängen oder zwei Seitenlängen und dem zugehörigen eingeschlossenen Winkel oder mit einer Seitenlänge und den benachbarten Winkeln),
- (6) zu einem Dreieck Winkelhalbierende, Höhen, Seitenhalbierende und die Mittelsenkrechten der Seiten konstruieren (und damit den Mittelpunkt des Inkreises, den Höhenschnittpunkt, den Schwerpunkt und den Mittelpunkt des Umkreises konstruieren).

Zunächst noch einmal die üblichen (mit Orientierung versehenen) Bezeichnungen von Seiten(längen) und Winkeln in einem Dreieck:



1.2. Eine kleine Liste grundlegender Bezeichnungen

P, Q, R, P_1 usw.	Punkte in der Ebene oder im dreidimensionalen Anschauungsraum
g, h usw.	Geraden in der Ebene oder im Raum
$g \parallel h$	g liegt parallel zu h
$g \perp h$	g steht senkrecht auf h
PQ	die Strecke von P nach Q
$d(P, Q)^{*1)}$	der Abstand zwischen P und Q , d.h. die Länge der Strecke von P nach Q
$D(P, M)$	Abstand zwischen P und der Punktmenge M (= Minimalabstand!)
ΔPQR oder nur PQR	das Dreieck mit den Eckpunkten P, Q, R (P, Q, R nicht kollinear) Punkte heißen kollinear , wenn sie auf einer Geraden liegen.
$\overrightarrow{PQ}$	der Vektor von P nach Q (Pfeilanzfang in P , Pfeilspitze in Q)
$\left \overrightarrow{PQ} \right $	die Länge dieses Vektors , d.h. wieder: $d(P, Q) = \left \overrightarrow{PQ} \right $.
$\vec{a}, \vec{b}$ usw.	freie Vektoren
$\angle (\vec{a}, \vec{b})$	der Winkel zwischen den Vektoren $\vec{a}, \vec{b}$ (Voraussetzung: $\vec{a}, \vec{b} \neq \vec{0}$)
$\angle (P; Q, R)^{*2)}$	der Winkel zwischen den Schenkeln PQ und PR , deren Längen > 0 sind.
$\vec{x}_P^O$ usw.	Ortsvektor von P bezüglich des Ursprungspunktes O , also $\vec{x}_P^O = \overrightarrow{OP}$
$\overrightarrow{PQ}$	der Vektor von P nach Q (Pfeilanzfang in P , Pfeilspitze in Q)
(a, b)	das Zahlenpaar a, b , analog (a, b, c) Zahlentripel, für $a, b, c \in \mathbb{R}$ auch allgemeiner für andere Objekttypen, auch (P, g) das Paar aus dem Punkt P und der Geraden g
$P(a, b), P(a, b, c)$, auch $P(a b)$ usw.	der Punkt P zum Koordinatenpaar (a, b) bezüglich eines vorausgesetzten Koordinatensystems , bzw. der Punkt P im Raum zum Koordinatentripel (a, b, c)
$F_{\Delta PQR}$ oder F_{PQR}	der Flächeninhalt des Dreiecks ΔPQR
$U_{\Delta PQR}$	der Umfang des Dreiecks PQR , also $U_{\Delta PQR} = \overline{PQ} + \overline{QR} + \overline{RP} = d(P, Q) + d(Q, R) + d(R, P)$ wohl klar
Rechteck, Quadrat	ein Viereck mit jeweils parallelen sich gegenüberliegenden Kanten
Parallelogramm	ein Parallelogramm mit gleichen Seitenlängen
Rhombus, Raute	ein Viereck mit zwei parallelen Kanten
Trapez	heißt eine Menge von Punkten, wenn sie mit zwei Punkten stets auch deren Verbindungsstrecke enthält
konvex	

Anmerkungen:

*1) Manchmal wird auch $\overline{PQ}$ für die Länge der Strecke von P nach Q verwandt, allerdings auch mit ' $\overline{PQ}$ ' die Strecke selbst bezeichnet. Daher wollen wir diese Bezeichnung überhaupt nicht verwenden.

*2) Oft wird mit ' $\angle (QPR)$ ' bezeichnet, was wir hier mit ' $\angle (P; Q, R)$ ' bezeichneten, in der anderen Notation steht der Scheitelpunkt des Winkels also in der Mitte.

Generelle Bemerkung: Die Bezeichnungen variieren, aber aus dem Kontext und den begleitenden Worten sollte man die Bedeutung immer verhältnismäßig leicht herausbekommen.

1.3. Einige notwendige begriffliche Präzisierungen. $(2, 3, 3) \neq (3, 3, 2)$, bei Paaren, Tripeln usw. kommt es auf die Reihenfolge an, ebenso ist $(2, 2) \neq 2$. Man unterscheide also endliche Folgen von Mengen.

Mit 'Dreieck' kann Verschiedenes gemeint sein: Der Rand als Menge von Punkten (also mit bestimmter Lage im jeweiligen Raum), die Fläche als Menge von Punkten (wieder mit bestimmter Lage im Raum) - einschließlich des Randes oder ohne den Rand?. Schließlich die geometrische Figur des Randes oder der Fläche (also abgesehen von der Lage), dann sieht man von der Lage ab, ein Exemplar wird als äquivalent zu einem anderen betrachtet, wenn dies durch Verschieben, Drehen, Spiegeln in jenes überführt werden kann.

Warnung: Keinesfalls verstehe man unter 'Dreieck' dasselbe wie 'Flächeninhalt des Dreiecks'.

Die analogen Bemerkungen sind natürlich zu machen für Kreise, Vierecke usw.

Eine weitere grundlegende Präzisierung bei Aussagen:

'Wenn A , dann B ' heißt für Aussagen A, B dasselbe wie: Stets gilt: (*nicht* A) oder B , also ist ausgeschlossen, dass A gilt, B aber nicht. Kurzform: $A \implies B$.

' A genau dann, wenn B ' heißt: $(A \implies B)$ und $(B \implies A)$. Kurzform: $A \iff B$. Gegenbeispiel: Wenn ein Dreieck gleichseitig ist, so ist es gleichschenkelig, aber nicht jedes gleichschenkelige Dreieck ist gleichseitig.

Warnung also: Aus $A \implies B$ schließe man also niemals $B \implies A$. Aber eine **Konvention sollte man kennen:** In einer Definition (und nur in Definitionen!) heißt 'wenn' dasselbe wie 'genau dann, wenn'.

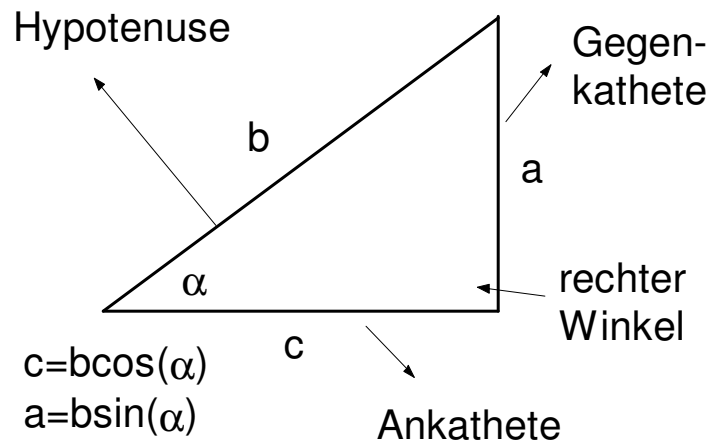
Noch ein paar Konventionen, die am Anfang verwirren können:

$a, b, c > 0$ heißt: $a > 0$ und $b > 0$ und $c > 0$, usw.

Spricht man von Punkten P, Q, R , so ist zunächst einmal keineswegs ausgeschlossen, dass $P = Q$ und $Q = R$. Will man voraussetzen, dass es sich um drei paarweise verschiedene Punkte handelt, so muss man in Worten sagen: 'Seien P, Q, R paarweise verschiedene Punkte', oder auch symbolisch $P \neq Q$ und $Q \neq R$ und $P \neq R$. Einfacher geht das mit Indexschreibweise so: Seien Punkte P_i gegeben, $1 \leq i \leq n$, mit $P_i \neq P_j$ für $1 \leq i < j \leq n$.

Merke also: Haben zwei Objekte verschiedene Bezeichnungen, so müssen die Objekte nicht verschieden sein! Das hat folgenden Sinn: 'Der Schwerpunkt eines Dreiecks bei homogener Massenbelegung S ist derselbe wie der Schnittpunkt T der Seitenhalbierenden' ist ein wahrer Satz, allgemeingültig. Um ihn zu beweisen, hat man ' $S = T$ ' zu beweisen, mit den eingeführten Beschreibungen. Man darf eben nicht für die Punkte gemäß den Beschreibungen sofort denselben Buchstaben verwenden, sie könnten ja zunächst verschieden sein. Aber man benötigt nun, dass die so mit S und T verschieden bezeichneten Punkte gleich sein können (und es im Beispiel auch zwingend sind!). Es wäre also fatal, wenn verschiedene Bezeichnungen automatisch verschiedene Objekte bezeichnen würden.

1.4. Elementares über Sinus, Cosinus, Tangens. Wir betrachten das folgende Bild eines rechtwinkligen (sonst allgemeinen) Dreiecks:



Elementare Eigenschaften von sin, cos, tan		
Elementare Definition	$\sin(\alpha) = \frac{\text{Länge der Gegenkathete}}{\text{Länge der Hypotenuse}}$	
im rechtwinkligen	$\cos(\alpha) = \frac{\text{Länge der Ankathete}}{\text{Länge der Hypotenuse}}$	
Dreieck	$\tan(\alpha) = \frac{\text{Länge der Ankathete}}{\text{Länge der Gegenkathete}}$	
Winkelmaß in Grad	Vollwinkel: 360° , rechter Winkel: 90°	
Bogenmaß für Winkel	Vollwinkel: 2π , rechter Winkel: $\pi/2$, usw.	
	Mit sin, cos, tan sollte man nur im Bogenmaß rechnen!	
Kleine Wertetabelle	x	$\sin(x)$ $\cos(x)$
	0	0 1
	$\pi/6$	$\frac{1}{2}$ $\frac{1}{2}\sqrt{3}$
	$\pi/3$	$\frac{1}{2}\sqrt{2}$ $\frac{1}{2}\sqrt{2}$
	$\pi/2$	$\frac{1}{2}\sqrt{3}$ $\frac{1}{2}$
Wichtige Formeln	$\sin^2(x) + \cos^2(x) = 1$	
zu sin, cos	$\sin(x + \frac{\pi}{2}) = \cos(x)$, $\cos(x - \frac{\pi}{2}) = \sin(x)$	
	$\sin(x + \pi) = -\sin(x)$, $\cos(x + \pi) = -\cos(x)$	
	$\sin(x + 2\pi) = \sin(x)$, $\cos(x + 2\pi) = \cos(x)$	
(Additionstheoreme)	$\sin(x + y) = \sin(x)\cos(y) + \cos(x)\sin(y)$	
	$\cos(x + y) = \cos(x)\cos(y) - \sin(x)\sin(y)$	
Zu tan	$\tan(x) = \frac{\sin(x)}{\cos(x)}$, $\cos(x) \neq 0$	
(Additionstheorem)	$\tan(x + y) = \frac{\tan(x) + \tan(y)}{1 - \tan(x)\tan(y)}$	

2. Sätze über Dreiecke im Umfeld der Pythagorasbeziehung

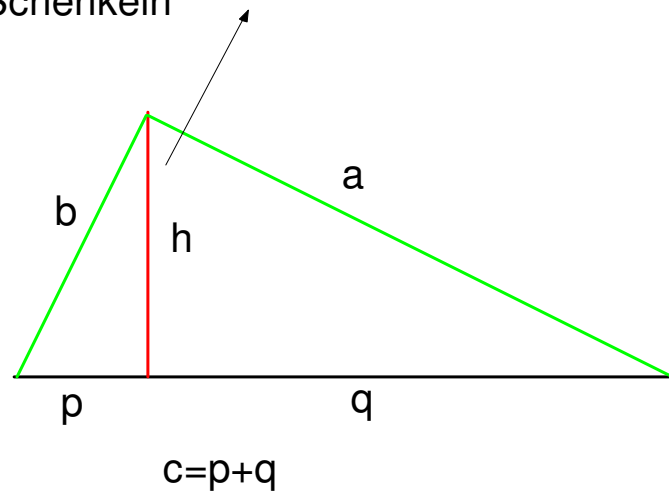
Wir wiederholen ein paar aus der Schulmathematik bekannte Sätze über Dreiecke:

SATZ 1 (Kathetensatz, 'Satz des Pythagoras', Höhensatz des Euklid). *In einem rechtwinkligen Dreieck mit rechtem Winkel γ und Hypotenuse $c = p + q$ (p, q die Abschnitte, welche auf der Seite c von der Höhe h auf c gebildet werden) gilt:*

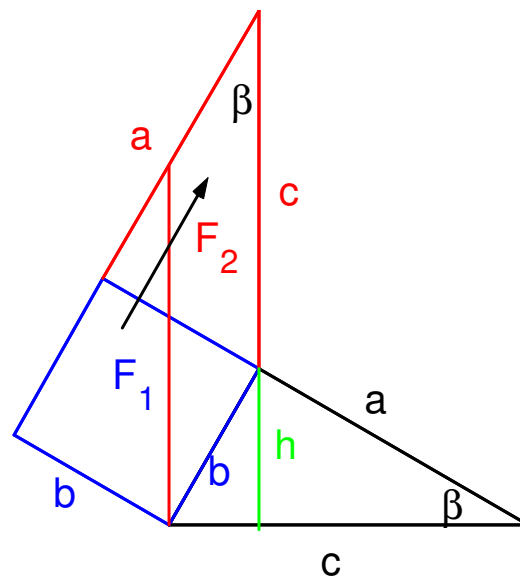
- (1) (Kathetensatz) $a^2 = qc$ und $b^2 = pc$
- (2) ('Pythagoras') $c^2 = a^2 + b^2$
- (3) (Höhensatz) $h^2 = pq$.

Zur anschaulichen Vorstellung und für die Erinnerung an die üblichen Bezeichnungen betrachte man noch einmal das Bild:

rechter Winkel zwischen den grünen Schenkeln

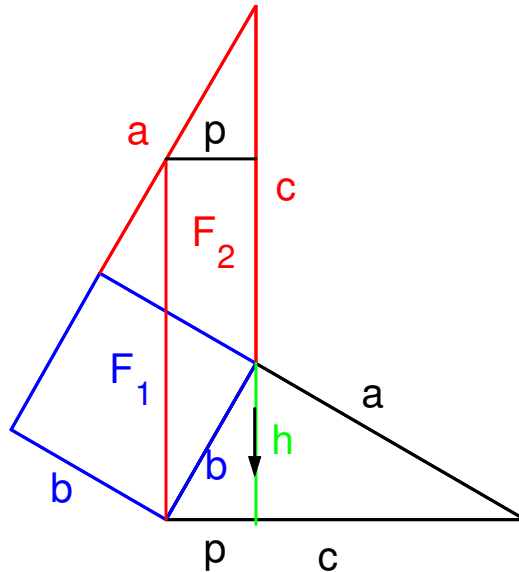


Beweis: Für (1) zeichne man zunächst über den Seiten die entsprechenden Quadrate, dann schere man das Quadrat mit Seitenlänge a zur Höhenlinie, die auf der Höhenggeraden liegende Seite des entstandenen Parallelogramms ist c (!), dann schert man dies Parallelogramm hinunter auf die Seite c , damit ist $a^2 = qc$ gezeigt. (Wir benutzen dabei, dass ein Parallelogramm bei Scherung nicht seinen Flächennhalt ändert.) Völlig analog sieht man $b^2 = pc$. Folgende Bilder illustrieren den Vorgang:



Das blaue Quadrat wird zum roten Parallelogramm (untere Seite ist blau) geschert wie der Pfeil andeutet, also $F_1 = F_2$ (damit sind die Flächeninhalte gemeint). Das obere Dreieck mit roten Seiten und einer blauen hat zwei Seiten mit Längen a, b und dazwischen einen rechten Winkel, ist also kongruent zum ursprünglichen, und so steht der rote Buchstabe c mit Recht. Nun braucht man nur noch die Scherung des roten Parallelogramms entlang h zum Rechteck, dessen Schmalseiten mit p bezeichnet sind: Dies Rechteck hat offenbar den Flächeninhalt pc , und es hat denselben Flächeninhalt wie das rote Parallelogramm, also

insgesamt $b^2 = pc$. Natürlich geht das auf der anderen Seite analog.



(1) $\implies$ (2) ergibt sich sofort mit $c = p + q$, also mit (1) $a^2 + b^2 = qc + pc = (p + q)c = c^2$.

(2) $\implies$ (3) ergibt sich mit $h^2 = a^2 - q^2 = b^2 - p^2$ (Pythagoras auf die kleinen Dreiecke angewandt), also

$$2h^2 = a^2 + b^2 - (p + q)^2 + 2pq = a^2 + b^2 - c^2 + 2pq,$$

mit $a^2 + b^2 - c^2 = 0$ (Pythagoras auf das große Dreieck angewandt) also $h^2 = pq$ und damit (3). ■

Bemerkung: Aus folgender Beobachtung ergeben sich sofort alle drei Sätze:

$$\begin{aligned} \frac{p}{b} &= \frac{b}{c} = \frac{h}{a} (= \cos(\alpha)), \\ \frac{q}{a} &= \frac{a}{c} = \frac{h}{b} (= \sin(\alpha)) \end{aligned}$$

also $pc = b^2$ und $qc = a^2$, somit (1) und (2). Aber auch: $p = \frac{hb}{a}$ und $q = \frac{ha}{b}$, also $pq = h^2 \frac{ab}{ab} = h^2$. Damit sind alle drei Sätze direkt mit der Ähnlichkeit der beiden kleinen Dreiecke und des großen Dreiecks erledigt.

Für die folgende starke Verallgemeinerung des Pythagoras-Satzes wird folgende Definition benötigt:

DEFINITION 1. Zwei Figuren heißen *ähnlich*, wenn sie durch *zentrische Streckung* auseinander hervorgehen, d.h. man bekommt die eine aus der anderen durch 'Aufblasen' bzw. 'Schrumpfen'. Genau ist eine *zentrische Streckung* mit Zentrum O und Streckungsfaktor $k > 0$ die Abbildung, welche jeden Punkt P mit Ortsvektor (bzgl. O) $\vec{x}_P$ auf den Punkt P' mit Ortsvektor $k\vec{x}_P$ abbildet.

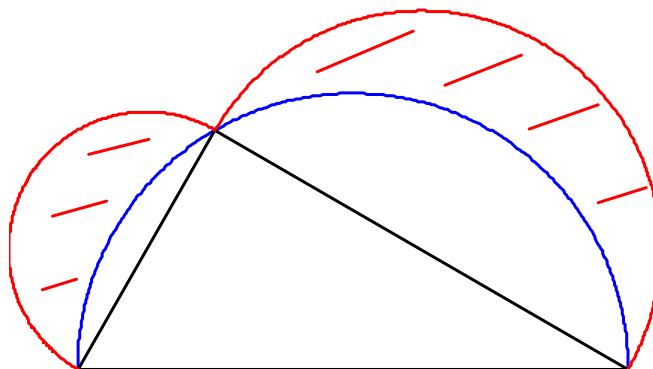
Bemerkungen: Es kommt auf das Aufblasen bzw. Schrumpfen in allen Raumrichtungen an, in der Ebene entsteht also durch zentrisches Strecken mit Streckfaktor k aus einer Figur mit Flächeninhalt F eine solche mit Flächeninhalt $k^2 F$ (im dreidimensionalen Raum wäre der neue Flächeninhalt $k^3 F$.)

Beispiele: Alle Dreieckspaare mit gleichen Winkeln sind ähnlich, alle Quadrate sind ähnlich, alle Kreise sind ähnlich zueinander, ebenso alle Kugeln und alle Würfel.

SATZ 2 (Verallgemeinerung des Pythagoras-Satzes). Errichtet man über den Seiten eines **rechtwinkligen** Dreiecks ähnliche Figuren (jeder Art), M_a über der Kathete mit Seitenlänge a , entsprechend M_b über der zweiten Kathete und M_c über der Hypotenuse, dann gilt für die zugehörigen Flächeninhalte F_a, F_b, F_c , dass

$$F_c = F_a + F_b$$

Beweis: Die Flächeninhalte der Figuren verhalten sich wie die Quadrate der Dreieckseiten, und mit Pythagoras $a^2 + b^2 = c^2$ (mit den üblichen Bezeichnungen).

Berühmtes Beispiel ('Möndchen des Hippokrates'):

Hier ist der Flächeninhalt der 'Möndchen' zusammen der des rechtwinkligen Dreiecks.

SATZ 3 (Kosinussatz). Mit den üblichen Bezeichnungen gilt

$$c^2 = a^2 + b^2 - 2ab \cos(\gamma)$$

Anschaulicher Beweis für spitzen Winkel γ : Man hat dann

$$\begin{aligned} \cos(\gamma) &= \cos(\gamma_1 + \gamma_2) = \cos(\gamma_1) \cos(\gamma_2) - \sin(\gamma_1) \sin(\gamma_2) \\ &= \frac{h}{b} \cdot \frac{h}{a} - \frac{p}{b} \cdot \frac{q}{a} \\ ab \cos(\gamma) &= h^2 - pq = b^2 - p^2 - pq = a^2 - q^2 - pq, \\ 2ab \cos(\gamma) &= a^2 + b^2 - p^2 - q^2 - 2pq = a^2 + b^2 - c^2. \end{aligned}$$

Bemerkung: Man sollte verstehen, dass der Satz den Satz des Pythagoras stark verallgemeinert, der aus dem Kosinussatz speziell mit $\cos(\pi/2) = 0$ folgt. Insbesondere sieht man weiter, dass $c^2 > a^2 + b^2$, wenn $\gamma > \pi/2$, und $c^2 < a^2 + b^2$, wenn $\gamma < \pi/2$. Denn $\cos(\gamma) < 0$ für $\pi \geq \gamma > \pi/2$, und $\cos(\gamma) > 0$ für $0 < \gamma < \pi/2$.

SATZ 4 (Sinussatz). Mit den üblichen Bezeichnungen gilt für die Seiten eines beliebigen Dreiecks der Längen $a, b, c > 0$ und die Winkel α, β, γ mit der üblichen Bezeichnung:

$$\frac{\sin \alpha}{a} = \frac{\sin(\beta)}{b} = \frac{\sin(\gamma)}{c}$$

Beweis: Einfache Übung mit Hilfe des Flächeninhalts von Dreiecken.

Dreiecke und einige besondere Punkte von ihnen werden wir später wieder aufgreifen, damit die Sache nicht zu monoton wird. Außerdem ist es günstig, weitere Hilfsmittel bereitzustellen (insbesondere Vektorrechnung).

Bemerkung: Man kann mit den hier geübten Methoden auch die Additionstheoreme für $\sin, \cos$ einsehen. Aber auch dafür werden wir später mittels vektorieller Darstellung einen eleganten Beweis geben können.

3. Konvexe Mengen

DEFINITION 2. Eine Punktmenge M heißt konvex, wenn*)

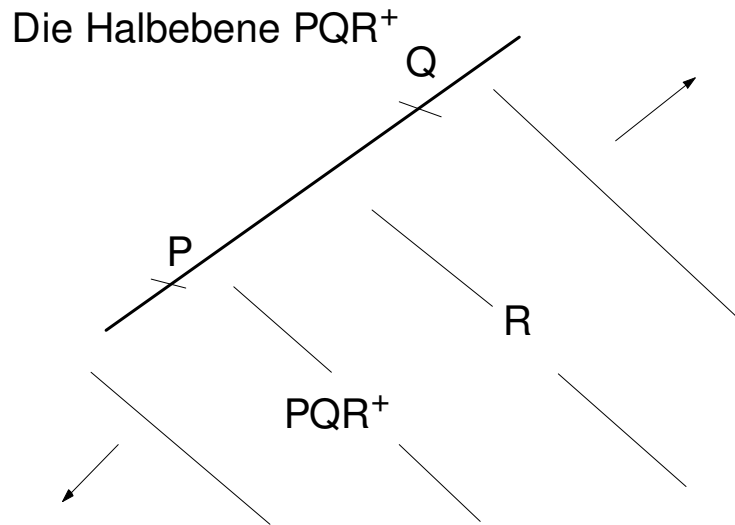
$$(A) \text{ Für alle Punkte gilt: } P, Q \in M \implies PQ \subset M.$$

Bemerkung: Insbesondere hier sollte man die früheren logischen Bemerkungen beherzigen: Der Fall $P = Q$ ist mit enthalten, wenn auch uninteressant. Das definitorische 'wenn' (mit dem Stern!) bedeutet 'genau dann, wenn'. D.h: Nur die Mengen mit der Eigenschaft (A) sind konvex. Schließlich: Das 'wenn...,dann' in der Aussage (A) (durch den Pfeil abgekürzt) ist genau dann wahr, wenn der Vordersatz falsch oder der Nachsatz wahr ist, konkret: Über Punkte, welche nicht in M liegen, wird nichts gesagt!

Sehr selbstverständliche konvexe Mengen: Halbebenen sind konvex, analog Halbräume im dreidimensionalen Raum. Dreiecke sind konvex, Rechtecke sind konvex.

DEFINITION 3. Eine Halbebene ist eine Gerade in der Ebene zusammen mit allen Punkten, die auf genau einer Seite der Geraden liegen. Seien P, Q, R drei nicht kollineare Punkte. Dann ist PQR^+ die Halbebene, welche die Punkte P, Q, R enthält und von der Geraden durch P und Q begrenzt wird. Analog ist dann PQR^- die Halbebene, welche von der Geraden durch P und Q begrenzt wird und welche den Punkt R nicht enthält.

Illustration der Bezeichnung PQR^+ (PQR^- geht dann gerade nach der anderen Seite.)



Beispiel für die Benutzung dieser Schreibweise: Wenn vom Punkt S zwei (nicht zusammenfallende) Schenkel (Halbgeraden mit Anfangspunkt S) ausgehen und $Q \neq S$ auf dem einen Schenkel und $R \neq S$ auf dem anderen Schenkel sitzt, dann hat man für den Raumwinkel W , die unendlich lange Fläche, welche von den Schenkeln eingeschlossen wird:

$$W = SQP^+ \cap SPQ^+.$$

SATZ 5. Der Durchschnitt von konvexen Mengen ist stets konvex.

(Der Beweis ist reine einfache Logik, Übungsaufgabe.)

Diese kleine Beobachtung verhilft zum nächsten Satz, für den man noch eine Definition benötigt:

DEFINITION 4. H heißt konvexe Hülle von M , wenn H die kleinste konvexe Obermenge von M ist.

Beispiele: Für nicht kollineare Punkte P, Q, R ist die Dreiecksfläche des Dreiecks PQR (einschließlich ihres Randes) die konvexe Hülle von $\{P, Q, R\}$. Besteht M aus allen Punkten eines Kreisrandes, so ist die konvexe Hülle von M die Kreisfläche einschließlich ihres Randes.

SATZ 6. Jede Punktmenge M hat eine eindeutige konvexe Hülle $H(M)$, die man definieren kann als kleinste Obermenge von M , welche konvex ist.

Beweis: Der ganze Raum (es sei die Ebene oder der dreidimensionale Anschauungsraum) ist offenbar konvex. Also ist

$$H(M) = \bigcap \{S \mid M \subset S \text{ und } S \text{ konvex}\},$$

in Worten: Der Durchschnitt aller Mengen, welche M enthalten und konvex sind, ist die konvexe Hülle von M . Denn nach dem vorigen Satz ist diese Menge $H(M)$ konvex, und sie enthält auch M , da der ganze Raum eine der Mengen im System ist. Außerdem ist jede konvexe Menge, welche M enthält, nach Definition eine Obermenge von $H(M)$. Somit ist $H(M)$ tatsächlich die kleinste Menge dieser Art. ■

Offenbar ist die Hülle einer konvexen Menge diese selbst. Das sind die typischen Eigenschaften eines Hüllenoperators, die wir einmal abstrahieren wollen:

Die drei Axiome für (beliebige) Hüllenoperatoren :

$$(i) \quad M \subset^* H(M)$$

$$(ii) \quad M \subset N \implies H(M) \subset H(N),$$

$$(iii) \quad H(H(M)) = H(M)$$

Zu *) : Das Symbol ' $\subset$ ' bedeutet hier dasselbe wie oft mit ' $\subseteq$ ' bezeichnet wird, also echte Teilmenge oder Gleichheit der Mengen. Weitere andersartige Beispiele: Lineare Hülle einer Menge M von Vektoren ist der kleinste Unter-Vektorraum, in dem M liegt. Die abgeschlossene Hülle einer Menge von Zahlen ist die kleinste Menge, welche diese enthält und aus welcher Grenzwerte nicht hinausführen. Es gibt noch viel mehr.

Wir werden bei der Vektorrechnung noch einmal auf (algebraische) Darstellungen konvexer Hüllen zurückkommen.

Zentralperspektive (Zentralprojektion)

In der Malerei entwickelte sich beim Übergang vom späten Mittelalter zur Renaissancezeit (ab etwa 1400, in Italien, dem Vorreiterland) ein deutlich realistischeres, natürlicheres Bild von Räumen, Innenräumen und Architekturen, im Rahmen einer überhaupt reichereren, realistischeren Darstellungen aus einem stark erweiterten Kreis von Objekten. Man bemerkt bereits im Mittelalter Versuche zentralperspektivischer Darstellungen, aber diese sind gewöhnlich nicht so ganz gelungen oder auch ziemlich misslungen, es fehlte die mathematische Durchdringung. Klarheit darüber wurde beispielhaft geschaffen von L. B. Alberti (1404-1472), der in einer Person Maler, Architekt, Mathematiker (und dann auch noch Sprachwissenschaftler, Verfasser von Satiren und Orgelspieler) war. Alberti schrieb ein Buch 'Della pittura' über die dem Maler notwendigen geometrischen Kenntnisse. Darin findet sich klar das Abbildungsprinzip der Zentralperspektive. Weitere Protagonisten der Zentralperspektive waren die berühmten Künstler Brunelleschi und Piero della Francesca.

1. Das Abbildungsprinzip der Zentralprojektion

Wie kommt es zum 'natürlichen' Bild, wie es dem Sehen dreidimensionaler Szenerien entspricht? Das hat Alberti klar ausgesprochen: Ein Auge (abstrahiert zu einem 'Augpunkt' O befindet sich in einiger Entfernung d von der Zeichenebene ZE , und jeder Punkt P jenseits von ZE wird durch einen 'Sehstrahl' mit dem Augpunkt verbunden, der Schnitt des Sehstrahls mit der Zeichenebene ist der Bildpunkt P' . (Man sieht mit zwei Augen, aber das ist für den Eindruck eines Bildes nicht so wichtig, stereographisches Sehen ist mit sehr viel Gehirnarbeit verbunden, von der wir das meiste nicht bemerken, und die Verhältnisse sind äußerst kompliziert.)

Wir definieren etwas formaler eine beliebige Zentralprojektion als Abbildung:

DEFINITION 5 (Zentralprojektion von O auf ZE). *Es sei ein Punkt O ('Augpunkt') und eine Ebene ZE ('Zeichenebene') gegeben mit $O \notin ZE$. Dann ist für alle Punkte P des dreidimensionalen Anschauungsraums (so dass OP nicht parallel zu ZE liegt) definiert:*

$$z_{O,ZE}(P) \quad : \quad = \text{der Schnittpunkt der Geraden durch } O \text{ und } P \text{ mit } ZE.$$

Wir schreiben auch einfach P' für $z_{O,ZE}(P)$.

Weiter ist $d = d(O, ZE) > 0$ die sogenannte Distanz, die bei vielen Konstruktionen wichtig wird. Mit H (genannt Hauptpunkt) wird der Lotfußpunkt des Lotes von O auf ZE bezeichnet.

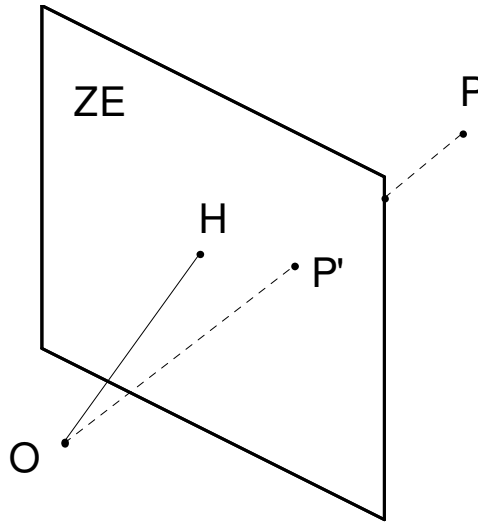
Bemerkungen: Die Bedingung ist selbstverständlich erfüllt für alle Punkte P , welche von O aus gesehen jenseits von ZE liegen. Nur solche möchte man abbilden. Wenn man nach dieser Vorschrift Abbilder herstellt, kann man eine böse Überraschung erleben: Vieles wird dann nicht nur unnatürlich, sondern geradezu falsch erscheinen: Der Grund dafür ist aber einfach der, dass wir nur einen begrenzten Blickwinkel von etwa 60° haben und eben nicht daran gewöhnt sind, das zu sehen, was außerhalb dessen liegt. Um solche Dinge im Zentralprojektionsbild zu vermeiden, kann man sich einfach an folgende Regel halten:

Praktische Regel : Das Zentralbild sollte nur Punkte enthalten,
welche in einem Kreis um H mit dem Radius von etwa $0.6 \cdot d$ liegen.

Begründung: $\tan(\pi/6)$ ist etwas weniger als 0.6. Es folgt, dass schöne zentralperspektivische Bilder im Allgemeinen ziemlich groß sein müssen, wenn man etwa $d = 1[\text{m}]$ oder noch mehr angenehm wählt. Für unsere kleinen Zeichenblätter werden wir uns auf $d = 0.1[\text{m}]$ beschränken müssen, Fokussierung ist dann jugendlichen Menschen noch gerade möglich.

Für praktisches Zeichnen stellen wir uns noch vor: Das Auge gehört zu einer menschlichen Figur, die vor der Zeichenebene steht und mit dem Auge senkrecht auf ZE schaut. Quer zu den gerade auf ZE ausgerichteten Füßen liegt die sogenannte 'Standlinie', diese ist also parallel zu ZE auf der sogenannten 'Bodenebene' - diese steht also senkrecht auf ZE . Allerdings kann der Betrachtende auch erhöht stehen oder auf einem tieferen Punkt. Dann rückt die 'Bodenebene' weiter nach unten bzw. nach oben. Für Vieles werden diese Verschiebungen gleichgültig sein, weil es oft nur auf Parallelität ankommt. Solche orientierenden Vorstellungen sind günstig für Lagebeschreibungen von Objekten. In der Zeichenebene bildet man noch die Gerade h , welche parallel zur Standlinie dur H geht. Das ist der 'Horizont', welche Bedeutung noch klar wird (alle Geraden parallel zur Bodenebene und nicht parallel zu ZE 'fluchten' auf diesen Horizont).

Wir illustrieren das Abbildungsprinzip noch mit folgender Skizze:



Da die Parallelprojektion viel vertrauter ist (und ihre eigenen Vorteile hat etwa bei technischen Zeichnungen usw.) und ihr Prinzip ähnlich dem der Zentralprojektion, wollen wir auch sie definieren und als Grenzfall der Zentralprojektion verstehen:

DEFINITION 6 (Parallelprojektion in Richtung von $\vec{a}$ auf die Ebene ZE). *Vorausgesetzt wird $\vec{a} \neq \vec{0}$, $\vec{a}$ nicht parallel zu ZE . Dann ist definiert für alle Punkte P des dreidimensionalen Anschauungsraums:*

$$par_{\vec{a}, ZE}(P) := \text{der Schnittpunkt der Geraden } g \text{ (} g \parallel \vec{a} \text{ mit } P \in g \text{) mit } ZE.$$

Bemerkung: Bildet man ein fest begrenztes Raumgebiet mit einer Zentralprojektion ab und lässt dann d gegen Unendlich gehen, so erhält man eine Parallelprojektion als Grenzfall der Zentralprojektion. Auch weit entfernte Gegenstände hinter ZE erscheinen im Zentralprojektionsbild wie parallel projiziert. Die spektakulären Effekte der Zentralprojektion ergeben sich gerade bei mäßiger Distanz und nahe gelegenen Gegenständen.

Wir kommen auf die Parallelprojektion noch einmal genauer zurück im Rahmen der Analytischen Geometrie. Einfache Würfelbilder usw. sind wir durchaus zu zeichnen gewöhnt.

Nunmehr widmen wir uns einigen einfachen theoretischen Grundlagen der Zentralperspektive, aus denen praktische Konstruktionen resultieren, die das korrekte zentralperspektivische Zeichnen klären und stark vereinfachen.

Das Abbildungsprinzip der Zentralprojektion ist denkbar einfach, hat aber den Nachteil, dass es sich um eine räumliche Konfiguration handelt. Man möchte dafür Konstruktionen haben, welche in der Zeichenebene selbst ausführbar und daher viel einfacher sind. Zugleich werden wir die wichtigsten Eigenschaften dieser Abbildungen kennenlernen.

2. Inzidenzerhaltung, Längenverzerrung, Fluchtpunkt von Geraden

'Inzidenz' ist einfach die Beziehung '...liegt auf...'. Das benutzen wir vor allem in der Form: ' P liegt auf AB ' oder ' P liegt auf g ', allgemein ' P liegt auf der Punktmenge $\mathcal{M}$ ', d.h. einfach: $P \in AB$, $P \in g$, $P \in \mathcal{M}$.

SATZ 7. *Inzidenz bleibt bei Zentralprojektion erhalten, d.h.*

Für alle Punkte P und Punktmenge $\mathcal{M}$ gilt: $P \in \mathcal{M} \implies P' \in \mathcal{M}'$.

Dabei ist $\mathcal{M}' := \{Q' \mid Q \in \mathcal{M}\}$.

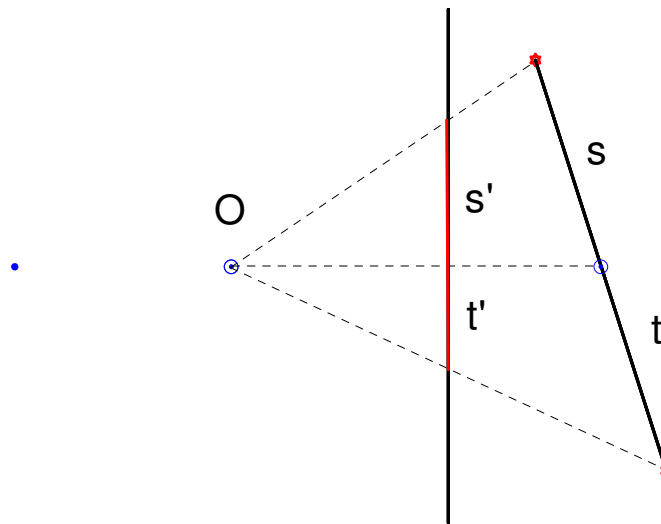
Beweis: Wenn $P \in \mathcal{M}$, dann gehört P' nach Definition zu $\mathcal{M}'$. ■

Folgerung: $(\mathcal{M} \cap \mathcal{N})' = \mathcal{M}' \cap \mathcal{N}'$. (Zeigen Sie das selber!) D.h. Das Bild eines Schnittes zweier Punktmenge ist die Schnittmenge der Bilder. Bilder von Strecken z.B. kann man korrekt 'naiv' miteinander schneiden! Nun werden wir sehen, dass Bilder von Strecken wiederum Strecken sind.

Offenbar werden alle Punkte auf einem Sehstrahl auf denselben Punkt abgebildet. Also bleiben Längen von Strecken nicht erhalten. Aber es gilt:

SATZ 8. *Strecken werden durch Zentralprojektion auf Strecken abgebildet (nur im Falle einer Strecke auf einer Sehstrahlgeraden auf einen Punkt, sonst auf Strecken einer Länge > 0).*

Beweis: Folgende Skizze zeigt die Behauptung:



Bemerkung und Anwendungen: Der Satz erlaubt also, nach Konstruktion von P', Q' das Abbild der Strecke PQ einfach als geradlinige Verbindung von P' nach Q' zu zeichnen. Mit der Folgerung aus dem vorigen Satz hat man also z.B.: Im Bild eines Rechtecks (das ein unregelmäßiges konvexes Viereck sein kann) kann man einfach gegenüberliegende Ecken geradlinig verbinden, der resultierende Schnittpunkt ist das Bild des Mittelpunktes vom Rechteck.

SATZ 9. *Längen bleiben bei Zentralprojektion nicht einmal auf derselben Geraden erhalten, sondern es entsteht der natürlich Seheindruck, dass Strecken derselben Länge sich verkürzen mit wachsender Entfernung vom Betrachter.*

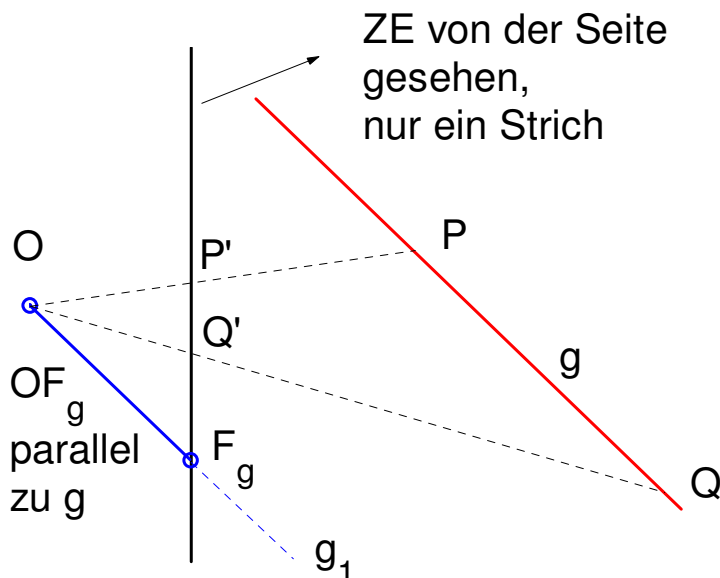
Beweis: Die letzte Skizze zeigt auch dies. ■

DEFINITION 7 (Fluchtpunkt F_g einer Geraden g). *Setzen wir zunächst voraus, dass $g \nparallel ZE$. Dann ist*

$F_g :=$ *der Schnittpunkt der Geraden g_1 ($g_1 \parallel g$ mit $O \in g_1$) mit ZE .*

Wenn $g \parallel ZE$, dann sagt man: 'Der Fluchtpunkt von g liegt im Unendlichen' und führt das mathematisch präzise so aus, dass man der Ebene ZE (welche aus allen Fluchtpunkten von Geraden $g \nparallel ZE$, also Fluchtpunkten im Endlichen besteht) als **neue Objekte** für jede Schar aller Geraden $g \parallel g_1$ mit $g_1 \parallel ZE$ noch einen 'Fluchtpunkt' F_{g_1} hinzufügt, so dass also $F_g = F_{g_1}$ für $g \parallel g_1$ mit $g_1 \parallel ZE$. Bequemer kann man auch formulieren: F_g ist die Klasse aller Geraden, welche parallel zu g liegen, wenn $g \parallel ZE$.

Bemerkungen: Die nachstehende Skizze zeigt, dass der so definierte Punkt den Namen verdient, weil er das durchaus gewohnte Phänomen beschreibt, dass das Bild einer Halbgeraden $\subset g \nparallel ZE$, deren Punkte sich vom Betrachter unendlich weit entfernen, nur eine Strecke ist (ohne den Endpunkt), deren Endpunkt gerade F_g ist. Man sagt auch: ' g fluchtet auf F_g '. Folgende Skizze illustriert diesen Sachverhalt:



Aus der Definition (für Fluchtpunkte im Endlichen und im Unendlichen) und der Skizze entnimmt man folgenden

SATZ 10. Jede Gerade hat genau einen Fluchtpunkt, und es gilt für alle Geraden g_1, g_2 (des dreidimensionalen Anschauungsraums):

$$g_1 \parallel g_2 \iff F_{g_1} = F_{g_2}.$$

Also gehört auch genau ein Fluchtpunkt zu jeder Parallelschar von Geraden.

Bemerkung zur Anwendung: Die wichtige praktische Konsequenz des Satzes ist diese: Will man zu einer gegebenen Bildstrecke von einem beliebigen Punkt aus eine Parallele ziehen, so zieht man zum Fluchtpunkt der Geraden, auf welcher die ursprüngliche Strecke liegt. Liegt dieser im Unendlichen (aber auch nur dann!), so heißt das: 'Naiv' (und korrekt!) parallel ziehen. Viele wichtige Fluchtpunkte bekommt man durch einfache Betrachtungen oder auch durch Konstruktion aus schon Bekanntem. Ein paar einfache

Beispiele: Alle Geraden, welche senkrecht auf ZE stehen, haben den Fluchtpunkt H (Hauptpunkt). Alle Geraden nicht parallel zu ZE , welche parallel zur 'Bodenebene' liegen, haben ihren Fluchtpunkt auf dem Horizont h . Weiteres folgt im nächsten Abschnitt.

Nun führen die Fluchtpunkte von Geraden aber auch zur äußerst wichtigen korrekten Bestimmung von Bildwinkeln: Dass Winkel durch Zentralprojektion völlig verzerrt werden, ist klar: Rechte Winkel können als beliebig spitze oder stumpfe abgebildet werden, was man am Zentralbild einer schräg nach hinten verlaufenden Wand sieht, die senkrecht auf dem Boden steht. (Skizzieren Sie das.)

3. Fluchtgeraden von Ebenen, Winkel und Winkelmesspunkte

3.1. Fluchtgeraden von Ebenen und deren Zusammenhang mit Fluchtpunkten von Geraden. Um beliebige Winkel zu beherrschen, benötigt man Fluchtgeraden von Ebenen (analog den Fluchtpunkten von Geraden).

DEFINITION 8 (Fluchtgerade g_E einer Ebene E). Sei die Ebene E (im dreidimensionalen Anschauungsraum) nicht parallel zu ZE . Dann ist die Fluchtgerade von E so definiert:

$$g_E := E_1 \cap ZE, \text{ wobei } E_1 \parallel E \text{ mit } O \in E_1.$$

(Da $E \not\parallel ZE$, ist g_E eine Gerade auf ZE .) Es verbleiben die Ebenen $E \parallel ZE$, dann definiert man (wieder analog zu den Fluchtpunkten im Unendlichen): Alle diese Ebenen haben die Fluchtgerade im Unendlichen, die aus allen Fluchtpunkten im Unendlichen besteht. Diese 'Gerade' heißt 'unendlich ferne Gerade'.

Bemerkung: Tatsächlich veranschaulicht das Bild oben zu den Fluchtpunkten auch die Fluchtebenen: Man stelle sich einfach alle Geraden dort als Ebenen vor, die senkrecht zum Skriptum-Blatt stehen, eben analog zu 'ZE von oben gesehen': Dann entsteht aus g eine Ebene und aus der Geraden g_1 eine Ebene und aus F_g eine Gerade, alle senkrecht zum Blatt.

Beispiele:

- 1) Eine Ebene parallel zur Bodenebene hat die Fluchtgerade h (Horizont).
- 2) Eine Ebene, welche senkrecht auf dem Boden steht, hat offenbar eine Fluchtgerade senkrecht zu h .
- 3) Insbesondere hat jede Ebene, welche senkrecht auf dem Boden steht und senkrecht zu ZE liegt, die Senkrechte zu h durch den Hauptpunkt H zur Fluchtgeraden.

SATZ 11. Es gilt für alle Ebenen E_1, E_2 :

$$(1) E_1 \parallel E_2 \iff g_{E_1} = g_{E_2}$$

(so dass also wieder genau eine Fluchtgerade zu einer gesamten Schar paralleler Ebenen gehört). Weiter besteht die Fluchtgerade von E genau aus allen Fluchtpunkten der Geraden auf E :

$$(2) g_E = \{F_k | k \text{ eine Gerade auf } E\} = \{F_k | k \text{ eine Gerade parallel zu } E\}$$

Bemerkungen: Aus (2) folgert man sofort: Es gilt für alle Geraden k :

$$(3) k \subset E \implies F_k \in g_E.$$

Beweis: (1) folgt unmittelbar aus der Definition wie bei den Fluchtpunkten von Geraden. Zu (2): Die zweite Gleichung ist klar, weil parallele Geraden denselben Fluchtpunkt haben. Wenn g_E die unendlich ferne Gerade ist, so gilt die Behauptung unmittelbar laut Definition der unendlich fernen Geraden. Sei also $E \not\parallel ZE$. Sei $Q \in g_E$. Dann bilde die Gerade k durch O und Q . (Man hat $Q \neq O$, weil $Q \in ZE$.) Dann ist $F_k = Q$ (nach Definition der Fluchtpunkte). Außerdem ist $k \parallel E$. Also $Q \in \{F_k | k \text{ eine Gerade parallel zu } E\}$. Sei umgekehrt F_k ein Fluchtpunkt einer Geraden k auf E . Sei k_1 die Gerade parallel zu k durch O , und sei E_1 die Ebene parallel zu E durch O . Dann ist $k_1 \subset E_1$. Es folgt $\{F_k\} = k_1 \cap ZE \subset E_1 \cap ZE = g_{E_1}$, also $F_k \in g_E$. ■

Anwendungen:

- 1) Kennt man von zwei Geraden auf E (oder auch nur parallel zu E) die Fluchtpunkte, so ist g_E die Gerade, welche durch diese beiden Fluchtpunkte geht.
- 2) Weiß man von einer Geraden k im Bild die Richtung der Bildgeraden und kennt eine Ebene E , auf der k liegt, samt der Fluchtgeraden g_E , dann kann man den Fluchtpunkt von k so bestimmen: Eine Strecke auf dem Bild von k wird verlängert, bis sie sich mit g_E schneidet: F_k ist dann dieser Schnittpunkt.
- 3) Kennt man von zwei Ebenen E_1, E_2 die (nicht parallelen) Fluchtgeraden, so ist der Schnittpunkt der Fluchtgeraden der Fluchtpunkt der Schnittgeraden von E_1 mit E_2 , also:

$$\{F_{E_1 \cap E_2}\} = g_{E_1} \cap g_{E_2}.$$

Konkrete Anwendung z.B.: Zwei Dächer (etwa eines L-förmigen Hauses, auch Dächer mit verschiedenen Anstellwinkeln) schneiden ineinander. Man kennt schon die Fluchtgeraden (im Endlichen) der Dächer (genauer der durch sie bestimmten Ebenen). Dann kann man den Fluchtpunkt der Schnittkante der Dächer ermitteln als Schnittpunkt der beiden Fluchtgeraden.

3.2. Das Einmessen gewünschter Winkel (oder auch das Nachmessen willkürlich gezeichneter Winkel). Vorab eine Warnung: Im Zentralbild darf man niemals einen Winkel 'naiv' antragen oder 'naiv' nachmessen, weil durch die Zentralprojektion alle Winkel völlig durcheinander geraten, außer in einem einzigen Falle: Winkel auf Ebenen parallel zu ZE können korrekt 'naiv' gezeichnet und gemessen werden.

Grundlage des korrekten Umgangs mit Winkeln ist folgender

SATZ 12. Vom Auge aus gesehen erscheinen die Fluchtpunkte der Geraden g_1 und g_2 im wahren Winkel, den die Geraden g_1, g_2 (vom Betrachtenden aus gesehen in die Ferne) bilden.

Beweis: Sehr einfach für den Fall, dass F_{g_1} und F_{g_2} im Endlichen liegen, weil $OF_{g_1} \parallel g_1$ und $OF_{g_2} \parallel g_2$. Die Behauptung stimmt aber auch noch, wenn einer oder beide Fluchtpunkte im Unendlichen liegen, weil 'zum Fluchtpunkt im Unendlichen Ziehen' dasselbe ist wie 'wörtlich parallel Ziehen'. ■

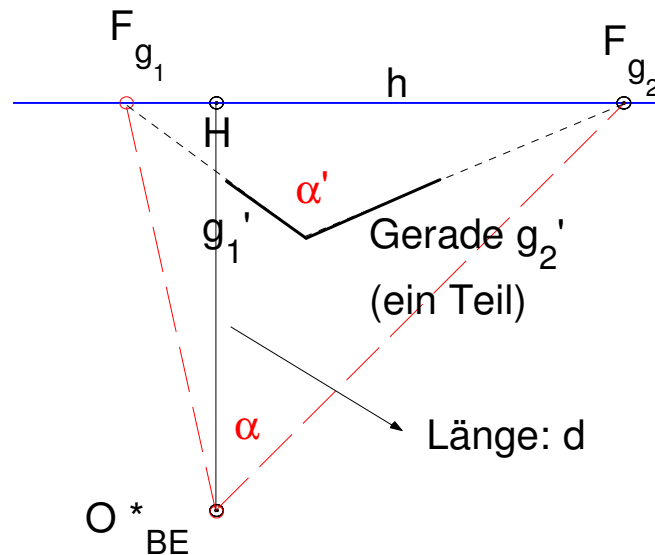
Etwas schwieriger, aber glücklicherweise auch noch ziemlich einfach, ist es, aus dieser Konfiguration im Raum (O liegt nicht in ZE) eine Konstruktion zu machen, die man (recht bequem) auf der Zeichenebene ausführen kann.

Dazu betrachten wir zunächst zwei einfache

Beispiele:

1) Liegen $g_1, g_2 \parallel ZE$ auf der Bodenebene (oder einer dazu parallelen Ebene), dann denke man sich das Dreieck $OF_{g_1}F_{g_2}$ um h in die Zeichenebene gedreht (nach unten oder nach oben): Dann landet O bei der Drehung auf einem Punkt O_{BE}^* auf einem Punkt der Zeichenebene, der auf der Lotgeraden zu h durch H liegt (oberhalb oder unterhalb des Horizonts) und den Abstand d zu H hat. Ist also bereits eine Geradenrichtung mit F_{g_1} (auf dem Horizont) bekannt, so kann man den unbekannten Fluchtpunkt F_{g_2} (mit vorgeschriebenem Winkel α von g_2 mit g_1 (in die Ferne gesehen!)) bestimmen: Man trägt in O_{BE}^* zum Schenkel $O_{BE}^*F_{g_1}$ den Winkel α an (auf der klar einsichtigen Seite), schneidet den freien Schenkel mit h , der Schnittpunkt ist F_{g_2} . Oder zum Nachmessen des tatsächlichen Winkels zwischen g_1 und g_2 bei bekannten F_{g_1}, F_{g_2} : Man misst einfach $\angle (O_{BE}^*, F_{g_1}, F_{g_2})$.

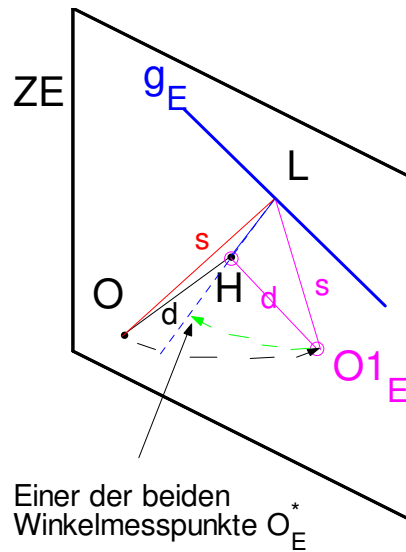
Hier ist eine Illustration - man beachte, dass die Striche bedeuten: 'Bild von...', und man beobachte, dass bei andern Anfangspunkt völlig andere Bildwinkel α' herauskämen, auch spitze:



2) Ebenso einfach ist es, entsprechende Winkelmesspunkte O_T^* (zwei gleich taugliche) für die Tiefenebenen zu bekommen, welche senkrecht sowohl auf dem Boden als auf ZE stehen: Diese Winkelmesspunkte liegen auf dem Horizont jeweils im Abstand d zu H . Damit kann man also alle Winkel auf Tiefenebenen korrekt einmessen, an welcher Stelle man auch will.

Nunmehr besprechen wir eine Konstruktion von Winkelmesspunkten, welche allgemein für Winkel auf beliebiger Ebene E taugt, wenn man g_E (im Endlichen) kennt. Um die Konstruktion zu finden und

ihre Korrektheit einzusehen, betrachten wir folgendes Bild der dreidimensionalen Situation:

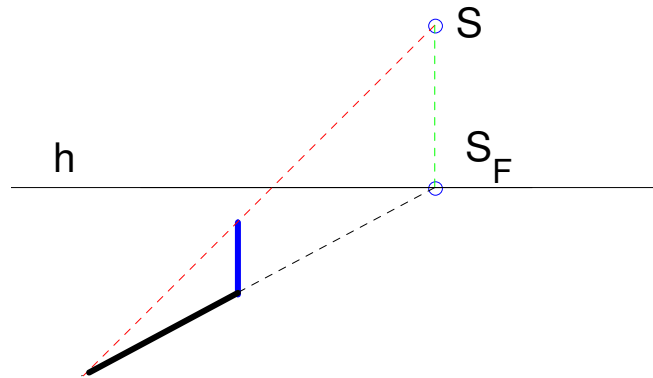


Wir lesen das Bild: O liegt im Raum vor ZE . Der Stab OH der Länge d steht senkrecht auf ZE . Das Dreieck OHL wird um das Scharnier der Lotgeraden auf g_E durch H und L (senkrecht auf g_E in der Zeichenebene) um 90° nach rechts in die Zeichenebene geklappt (die untere gestrichelte schwarze Linie deutet diese Drehung an). Dabei wird der Punkt O in den Punkt O_E^1 gedreht, der schwarze Stab der Länge d in den magentafarbenen Stab der Länge d , der auf der Ebene ZE liegt. Bei dieser Drehung entsteht auch der magentafarbene Stab der Länge s aus dem roten der Länge s . Nun drehen wir um L (grüne gestrichelte Linie, in der Zeichenebene!) den magentafarbenen Stab der Länge s auf die Lotgerade durch H und L . Der Punkt O_E^1 wird dabei in den Punkt O_E^* gedreht, und daher hat die Strecke O_E^*L wieder die Länge s . Insgesamt haben wir damit nur jedes Dreieck $OF_{g_1}F_{g_2}$ um g_E als Scharnier in die Zeichenebene geklappt, wobei O auf O_E^* geht. Also erscheinen nicht nur von O aus, sondern auch von O^* aus, die Fluchtpunkte F_{g_1} und F_{g_2} von Geraden parallel zu E im wahren Winkel der Geraden g_1 und g_2 . Was wir gewonnen haben, indem wir diese einfache letztere Drehung so scheinbar unnötig kompliziert beschrieben haben: Wir besitzen eine Konstruktion für O_E^* , welche sich alle in ZE abspielt und im nächsten Satz beschrieben wird, der mit der Erläuterung des vorigen Bildes bereits bewiesen ist:

SATZ 13. Wenn E mit g_E im Endlichen gegeben ist, d.h. $E \nparallel ZE$, dann kann man die beiden tauglichen Winkelmesspunkte für Winkel auf E so in ZE konstruieren:

1. Fülle das Lot von H auf g_E , Lotfußpunkt ist L .
2. Trage von H aus die Länge d parallel zu g_E ab (also senkrecht zum Lot).
Deren Endpunkt sei O_E^1 .
3. Stich den Zirkel in L ein und schlage einen Kreis um L mit dem Radius $d(L, O_E^1)$,
die beiden Schnittpunkte mit der Lotgeraden l sind die Winkelmesspunkte O_E^* .

und benutzt die im folgenden Abschnitt zu besprechende rechnerische Methode zur Bildkonstruktion. Im einfachen Falle jedoch wird der Schatten auf unsere idealisierte „Bodenebene“ geworfen, und dann ist folgende einfache zeichnerische Konstruktion möglich: Über den Fußpunkt F_P von P auf der Bodenebene, genauer dessen Bild $F_{P'}$, ermittelt man mittels des Sonnenfußpunktes S_F das Bild der senkrechten Projektion („senkrecht“ zur Bodenebene) des Sonnenstrahls durch P auf die Bodenebene. Damit schneidet man einfach das Bild des Sonnenstrahls durch P , wie folgende Skizzen für beide wesentlichen Fälle von Sonnenständen zeigen:



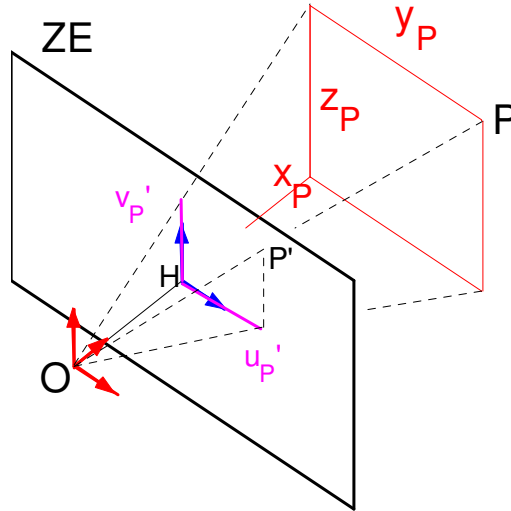
In diesem Bild steht die Sonne 'rechts oben' **vor dem Beobachtenden**, der Sonnenpunkt S ist der Fluchtpunkt der (parallelen) Sonnenstrahlen (diese können bei derartigen Entfernung zwischen Erde und Sonne als solche mit größter Genauigkeit angenommen werden). Auf dem als ideal eben gedachten 'Boden' steht senkrecht der blaue Stab, und das Ende seines Schattens kann man als Schnitt des (roten) Sonnenstrahls mit der schwarz gepunkteten Geraden vom Fußpunkt des Stabs zum Sonnenfußpunkt S_F konstruieren. Der Schatten des blauen Stabs auf dem Boden ist also der schwarze Stab (genauer sind das natürlich die Zentralbilder von diesen Objekten). Wir können uns auch leicht vorstellen, wie das auszusehen hat, wenn etwa die Sonne 'links oben' **hinter dem Beobachtenden** steht: Hier ist der Sonnenpunkt unterhalb des Horizonts, dahin fluchten die Sonnenstrahlen (wieder als parallel idealisiert), aber diesmal von der Sonne weg. Wieder ist die schwarze gepunktete Linie das Bild der Projektion des Sonnenstrahls

Um die Achse a (grün) senkrecht zur Bodenebene wird die schwarz umrahmt gezeichnete Wand $W \parallel ZE$ in die blau umrahmt gezeichnete gedreht oder umgekehrt - den Drehwinkel α kann man in O_{BE}^* nachmessen, offenbar ist er hier über 60° . Die 'Drehsehnen' sind rot gestrichelt, ihr Fluchtpunkt D ist tatsächlich einer der Winkelmesspunkte von W , O_W^* , allerdings der 'richtige' von den beiden. Der andere ergibt die Drehsehnen für die Drehung der schwarzen Wand nach vorn, in dieselbe Ebene von W . Man überzeuge sich durch Nachmessen und Einzeichnen davon.

Man kann sich durch eine Grundrisszeichnung auch klar machen, dass tatsächlich die beiden Drehsehnenfluchtpunkte (für beide beschriebenen Drehungen) die beiden Winkelmesspunkte sind. (Alternativ könnte man auch ausnutzen, dass das Dreieck, das von einer Kante parallel zur Bodenebene und der gedrehten gebildet wird, ein gleichschenkliges ist und die verbleibenden Winkel $(\pi - \alpha)/2$ konstruieren.

4. Die analytische Methode

Mit den Mitteln der Analytischen Geometrie kann man bei Kenntnis der Lage von ZE und H im Raum für einen beliebigen Raumpunkt P (nicht auf der Verschwindungsebene) den Bildpunkt in der Zeichenebene einfach ausrechnen, genauer: Aus dem Raumkoordinatentripel (x_P, y_P, z_P) von P kann man das Zeichenebenen-Koordinatenpaar $(u_{P'}, v_{P'})$ gewinnen. Das geht sogar sehr einfach, wenn man die zugehörigen Koordinatensysteme geeignet wählt, dass sie gut zusammen und zur Situation passen. Folgendes Bild illustriert die ganze Situation:



Dabei haben ZE, O, H, P, P' die bisher üblichen Bedeutungen. Ferner ist das rote kartesische **Augpunktsystem** K^A das für den dreidimensionalen Anschauungsraum, x_P ist die Tiefen-Koordinate (vom Auge aus gemessen) von P , y_P die 'Breitenkoordinate' (entlang dem Horizont gemessen, positiv nach rechts, negativ nach links), z_P die 'Höhenkoordinate' (senkrecht zum Horizont). Das passende (blaue) **Hauptpunktsystem** K^H hat die u -Achse parallel zur y -Achse und die v -Achse parallel zur z -Achse, jeweils mit denselben Einheiten versehen. Für die Koordinaten $u_{P'}$ und $v_{P'}$ des Bildpunktes P' erhält man mit Anwendung des Strahlensatzes unmittelbar die Beziehungen

$$\frac{u_{P'}}{y_P} = \frac{d}{x_P} = \frac{v_{P'}}{z_P},$$

also den

SATZ 14. *Man hat folgende Berechnung der Bildkoordinaten aus den Urbildkoordinaten:*

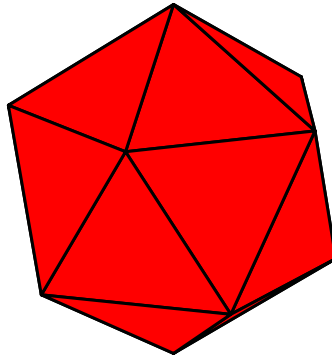
$$\begin{aligned} u_{P'} &= y_P \cdot \frac{d}{x_P}, \\ v_{P'} &= z_P \cdot \frac{d}{x_P}. \end{aligned}$$

Anwendungsbeispiel: Auf dem 'Boden' liegt 9.9[m] von ZE entfernt eine Strecke der Länge 10[m], $d = 0.1$ [m], 'Aughöhe' 2[m] (d.h. das Auge liegt 2[m] über dem 'Boden'). Der linke Endpunkt der Strecke liege 3[m] links von der Achse 'Auge-Hauptpunkt'. Dann hat man für den linken Endpunkt P der Strecke:

$$\vec{x}_P^{K^A} = \begin{pmatrix} 10 \\ -3 \\ -2 \end{pmatrix} [\text{m}], \text{ also } \vec{x}_{P'}^{K^H} = \begin{pmatrix} -3 \cdot \frac{0.1}{10} \\ -2 \cdot \frac{0.1}{10} \end{pmatrix} [\text{m}] = \begin{pmatrix} -3 \\ -2 \end{pmatrix} [\text{cm}].$$

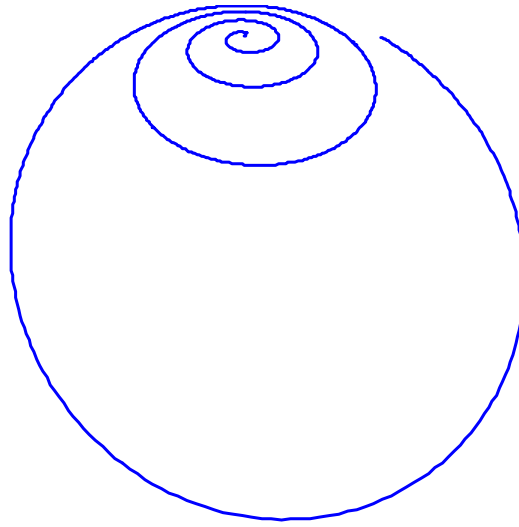
Man rechne nach, dass der rechte Endpunkt der Strecke die Bildkoordinaten $\begin{pmatrix} 7 \\ -2 \end{pmatrix} [\text{cm}]$ besitzt. Allgemeiner kann man feststellen, dass man für jede Ebene $E \parallel ZE$ einen festen Maßstab hat. Alle Bild-Längen auf einer solchen Ebene werden mit demselben Faktor an der Originallänge gebildet; der Faktor ist gerade d/x_E , wobei x_E die Entfernung zwischen Augpunkt und E ist.

Eine allgemeine wichtige Anwendung: Man kann (vgl. dazu den Abschnitt über Analytische Geometrie) recht komplizierte Punktmengen analytisch (als Menge von Koordinatentripeln) beschreiben, Kurven mit einem freien Parameter und Flächen mit zwei freien Parametern. Dann macht man aus den zugehörigen Parametermengen endliche Mengen, indem man nur ein hinreichend feines 'Gitter' bildet. Anschließend kann man mit dem vorigen Satz die Bildkoordinaten der so entstehenden endlichen Punktmenge ausrechnen und auf den Bildschirm 'plotten', mit einem sehr einfachen kurzen Computerprogramm. Zuvor kann man noch Drehwinkel angeben, um das Objekt in der gewünschten Weise zu drehen - wir werden auch die geeigneten Drehmatrizen kennenlernen, mit denen die 3d-Koordinaten entsprechend umgerechnet werden. Hier zeigen wir zwei Beispiele dafür:



Man sieht hier ein zentralperspektivisches Bild eines Ikosaeders (mit zwanzig Flächen, die alle gleichseitige Dreiecke sind), das ist einer der fünf 'vollkommenen' Körper.

Das zweite Beispiel zeigt ein zentralperspektivisches Bild einer Kurve - es handelt sich dabei um eine archimedische Spirale:



Grundzüge Analytischer Geometrie

Wir beginnen mit einigen noch ausstehenden Beobachtungen bei Dreiecken, die aus der Schulmathematik bekannt sind, und behandeln diese zunächst (im ersten Abschnitt) mit den Mitteln der Synthetischen Geometrie. Dann werden wir zur Methode der Analytischen Geometrie übergehen und sehen, dass sich starke Vereinfachungen ergeben.

1. Weiteres über Dreiecke: Umkreis, Inkreis, Schwerpunkt, Ankreise

Ein Dreieck besitzt ausgezeichnete Geraden bzw. Strecken, diese natürlich je dreifach, und drei ausgezeichnete Geraden eines Dreiecks schneiden sich gewöhnlich in einem einzigen Punkt.

1. Die Mittelsenkrechten (Geraden) der Seiten eines Dreiecks schneiden sich in einem einzigen Punkt M , und dieser Punkt ist der Mittelpunkt des Umkreises: Von ihm sind alle Punkte gleich weit entfernt, und man kann also um ihn den Umkreis des Dreiecks schlagen, der alle Ecken des Dreiecks enthält.

2. Die Winkelhalbierenden (Geraden) der Dreieckswinkel gehen durch einen einzigen Punkt, und dieser ist der Mittelpunkt des Inkreises, der alle Dreieckseiten berührt. Die zugehörigen winkelhalbierenden Strecken gehen jeweils von einer Ecke bis zur gegenüberliegenden Dreieckseite.

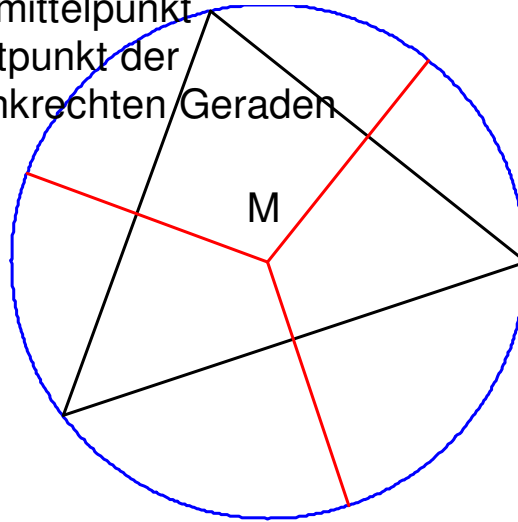
3. Die Höhen(geraden) eines Dreiecks schneiden sich in einem einzigen Punkt, dem **Höhenschnittpunkt** H . Die zugehörigen Höhenstrecken sind jeweils die Lote von einer Ecke zur gegenüberliegenden Seite. Man beachte, dass bei angrenzendem stumpfen Winkel der Höhenfußpunkt außerhalb des Dreiecks liegt.

4. Die Seitenhalbierenden Strecken verlaufen je von einer Ecke bis zum Mittelpunkt der gegenüberliegenden Seite. Sie schneiden sich in einem einzigen Punkt S , und dieser Punkt hat eine besondere Bedeutung, auch physikalisch: S ist der Schwerpunkt des Dreiecks bei homogener Massenverteilung. Das bedeutet: Wenn man sich die Dreiecksfläche aus Metall vorstellt, überall gleich dick, so muss man das Dreieck in diesem Punkt unterstützen, um Gleichgewicht zu haben. Oder auch: Man betrachtet die Ecken als Massenpunkte, in denen gleich große Massen liegen. Auch dann ist S wieder der Schwerpunkt. Aus den Ortsvektoren der Eckpunkte P, Q, R lässt sich der Schwerpunkt des Dreiecks sofort ausrechnen als $\vec{x}_S = \frac{1}{3} (\vec{x}_P + \vec{x}_Q + \vec{x}_R)$, ohne dass man den Schnittpunkt von zwei Seitenhalbierenden bildet. (Diese einfache Darstellung liefert die Analytische Geometrie, dazu später mehr.)

Die folgenden Bilder sollen diese Punkte in ihrer Bedeutung illustrieren:

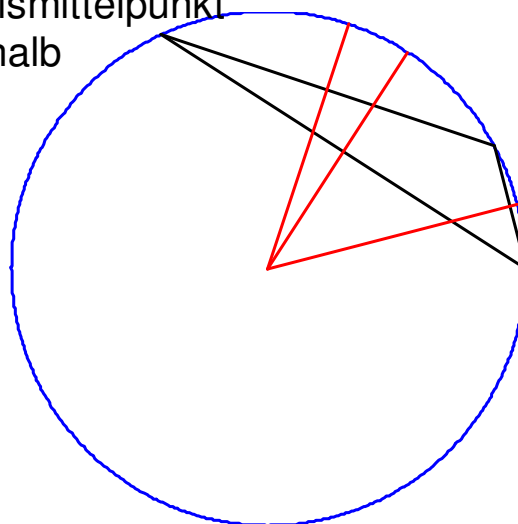
1. Umkreismittelpunkt als Schnittpunkt der Mittelsenkrechten auf den Dreieckseiten:

Umkreismittelpunkt
= Schnittpunkt der
mittelsenkrechten Geraden

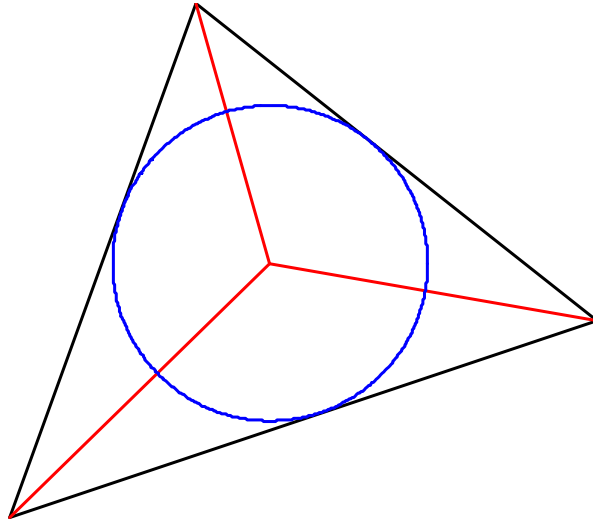


Selbstverständlich kann der Umkreismittelpunkt außerhalb des Dreiecks liegen, ein stumpfer Dreieckswinkel sorgt dafür:

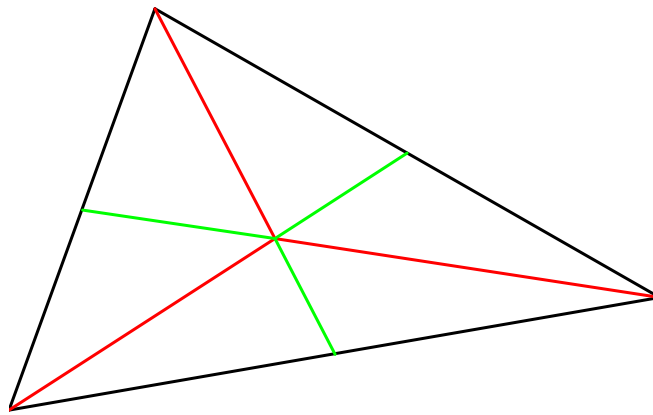
Umkreismittelpunkt
außerhalb



2. Mittelpunkt des Inkreises (blau) als Schnittpunkt der Winkelhalbierenden (rot):



3. Schwerpunkt des Dreiecks S bei homogener Massenverteilung = Schnittpunkt der Seitenhalbierenden (rot/grün): Das rote Stück verhält sich zum grünen wie 2 : 1 (auf jeder der Seitenhalbierenden).



Die Begründungen für Umkreis und Inkreis geben wir mit den Mitteln der bisher ausgeübten Synthetischen Geometrie, zum Schwerpunkt nehmen wir lieber später etwas Analytische Geometrie zu Hilfe.

1.) Zum Umkreis-Mittelpunkt: Man stellt fest, dass die mittelsenkrechte Gerade g_a auf der Dreiecksseite a die Menge aller Punkte ist, welche von den Ecken B, C gleich weit entfernt sind. Analog ist g_b die Menge aller Punkte, welche von A und B gleich weit entfernt sind. Folglich ist der Schnittpunkt von g_a mit g_b gleich weit von allen Dreieckspunkten entfernt und daher der Mittelpunkt des Umkreises für das Dreieck ABC . Dasselbe könnte man mit g_a und g_c machen, so dass also auch g_a und g_c sich wieder im Mittelpunkt des Umkreises schneiden. Daher schneiden sich auch alle drei mittelsenkrechte Geraden auf den Dreiecksseiten im selben Punkt.

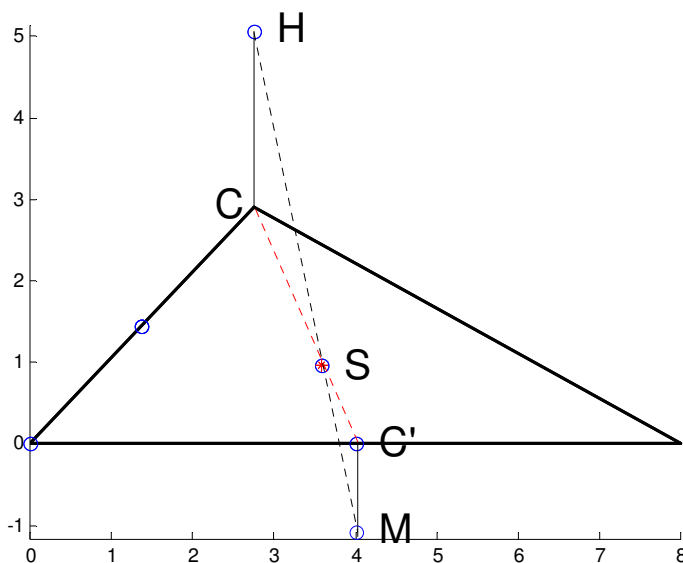
2.) Zum Inkreis-Mittelpunkt: Wir stellen uns das Dreieck ABC mit seinem Inkreis vor, der alle Dreiecksseiten berührt, mit den Berührungspunkten B_a auf der Seite a , B_b auf b und B_c auf c . Dann muss das Dreieck CB_aB_b gleichschenkelig sein und daher der Mittelpunkt M des Inkreises auf der Winkelhalbierenden durch C liegen, welche den Winkel γ halbiert. Mit dem gleichen Argument für die anderen Ecken A, B muss M auch auf den beiden anderen Winkelhalbierenden von ABC liegen. Aber wir müssen noch zeigen, dass es überhaupt einen Inkreis für ABC gibt. Dazu argumentieren wir analog zum Umkreismittelpunkt so: Die Winkelhalbierende durch C ist die Menge aller Punkte, welche Mittelpunkte eines Kreises sind, welcher die Geraden durch CA und CB berührt. Analog hat man das für die Winkelhalbierende durch B . Also ist der Schnittpunkt dieser beiden Winkelhalbierenden ein Punkt, zu dem es einen Kreis gibt, welcher a, b und c berührt. ■

Es gibt noch einen schönen Satz über die Punkte M, S, H (Umkreismittelpunkt M , Schwerpunkt S (bei homogener Massenverteilung) und Höhenschnittpunkt H): Klar ist: Wenn M und S zusammenfallen, dann handelt es sich um ein gleichseitiges Dreieck, und auch $H = M = S$. Aber wenn M und S nicht zusammenfallen, dann bestimmen sie eine Gerade, die sogenannte 'Eulersche' Gerade, auf der auch H liegen muss:

SATZ 15 (Satz von der Eulerschen Geraden). *Für $M \neq S$ gilt: Die Punkte M, S bestimmen eindeutig eine Gerade, und auf dieser liegt dann auch H .*

Bemerkung zu Leonhard Euler (1707-1783): Er stammte aus Basel und arbeitete zunächst dort, wurde einer der größten Mathematiker (und Physiker!) überhaupt und der Begründer der modernen Mathematik, der insbesondere auf den Gebieten der Analysis und Zahlentheorie ebenso wie der Grundlagen der Mathematik und der Angewandten Mathematik Bahnbrechendes hervorbrachte. Die Zahl e ist auch nach ihm benannt. Es ist unmöglich, hier darauf näher einzugehen, erwähnt werden soll hier nur, dass alle Mathematiker nach Euler sich dessen bewusst sind, Schüler von Euler zu sein, dass aber der Preußenkönig Friedrich 'der Große' seine Leistungen nicht zu schätzen wusste, weshalb er nach St. Petersburg ging. Wir werden Euler auch noch an einer tieferen Stelle der Geometrie begegnen.

Beweis des Satzes : Damit man sich diesen schönen Beweis bei Coxeter besser vorstellen kann, betrachten wir dabei folgende Skizze:



Wir definieren den Punkt H durch $\overrightarrow{MH} = 3\overrightarrow{MS}$. Damit ist $\overrightarrow{SH} = 2\overrightarrow{MS}$ und also folgt mit $\overrightarrow{SC} = 2\overrightarrow{C'S}$ (dies zeigen wir weiter unten als erste Anwendung der Vektorrechnung), dass $\overrightarrow{CH} \parallel \overrightarrow{MC'}$ (Strahlensatz-Umkehrung!). Daher liegt H auf der Höhengerechten durch C . Völlig analog folgt, dass H auch auf den Höhengerechten durch A und durch B liegt. ■

Die **Ankreise** eines Dreiecks sind die Kreise, welche das Dreieck an seinen Seiten außen derart berühren, dass sie auch die Verlängerungen der beiden anderen Dreiecksseiten berühren. Man braucht für

den Mittelpunkt des Ankreises an a nur die Seite c nach außen zu verlängern und den Winkel zwischen diesem Schenkel und a zu halbieren, dasselbe führt man mit b, c aus (b nach außen verlängert), der Schnittpunkt dieser Winkelhalbierenden ist der Mittelpunkt des Ankreises an die Seite a . Analoges gilt für die anderen Ankreise.

Nunmehr stellen wir die nützlichen einfachsten Elemente der Analytischen Geometrie bereit und wenden sie auf geometrische Fragestellungen an.

2. Punkte, Ortsvektoren, freie Vektoren, Koordinatendarstellungen, Vektorräume

Wir zeichnen willkürlich (in konkreten Situationen werden wir bedacht wählen) einen Punkt O im dreidimensionalen Anschauungsraum E^3 aus (alle Betrachtungen hier gelten auch sinngemäß für E^2 , die Anschauungsebene. Dann heißt der Pfeil von O nach P , also $\overrightarrow{OP}$, Ortsvektor von P , den wir auch mit $\vec{x}_P$ bezeichnen. Stets muss dabei O klar sein, will man von einem Ursprungspunkt O zu einem neuen O_1 übergehen, so hat man $\vec{x}_P^O$ bzw. $\vec{x}_P^{O_1}$ unterscheidend zu bezeichnen. Offenbar hat man eine bijektive Abbildung

$$\begin{aligned} E^3 &\rightarrow V_O^3 \\ P &\mapsto \vec{x}_P \end{aligned}$$

Bemerkung: Der Pfeil mit Länge Null heißt Nullvektor und tritt offenbar als Ortsvektor $\vec{x}_O^O$ auf.

Entsprechend wird aus einer Menge von Punkten eine Menge von Ortsvektoren. Was soll das? Nun, mit Vektoren lassen sich nützliche geometrische Abbildungen erklären, die man mit Punkten so nicht hinbekommt. Das liegt daran, dass man für Vektoren (anders als für Punkte) nützliche Rechenoperationen hat: Addition und Multiplikation mit einer reellen Zahl. Diese Operationen erklären wir zunächst für freie Vektoren, sie ist dann auch für Ortsvektoren klar.

Ein freier Vektor $\vec{a}$ ist die Äquivalenzklasse aller Pfeile, die zu einem gegebenen konkreten Pfeil kongruent sind (also gleich lang und in dieselbe Richtungweisend).

Wenn wir ein Koordinatensystem K^O auszeichnen (drei Achsen, nicht in einer Ebene, mit Einheiten, die nicht gleich sein müssen), mit Ursprung in O , so können wir jedem Punkt P (und damit seinem Ortsvektor $\vec{x}_P^O$) eindeutig ein Koordinatentripel zuordnen:

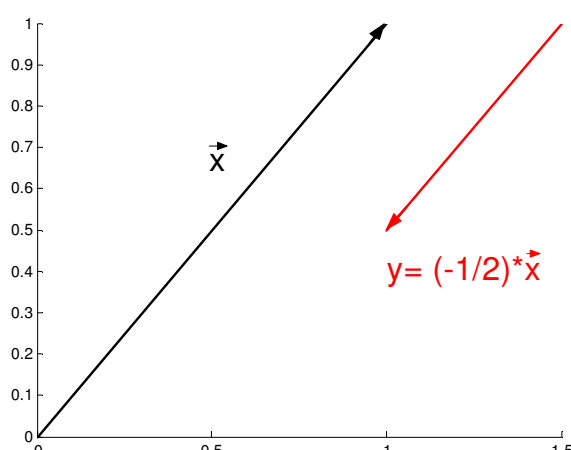
$$\begin{aligned} E^3 &\rightarrow \mathbb{R}^3 \\ P &\mapsto \vec{x}_P^{K^O} = \begin{pmatrix} x_P^{K^O} \\ y_P^{K^O} \\ z_P^{K^O} \end{pmatrix}, \end{aligned}$$

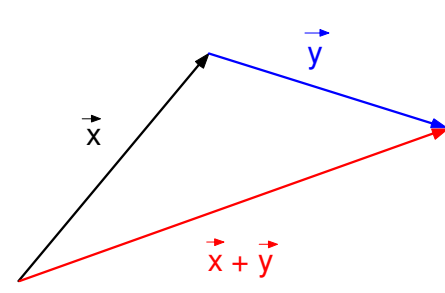
weil eindeutig bestimmt ist, wie weit und in welcher Richtung man von O aus entlang den Koordinatenachsen laufen muss, um P zu erreichen.

Addition von Vektoren und Multiplikation eines Vektors mit einer Zahl hat man nun in zwei Ausfertigungen, einer geometrischen (Pfeile aneinandersetzen bzw. strecken) und einer rechnerischen (mit den Koordinatentripeln ausgeführt): Die linearen Vektorraumoperationen

Wichtig ist es immer wieder, aus einfacheren Dingen kompliziertere zusammenzusetzen. Gerade dazu leisten die linearen Operationen Beträchtliches. Man stellt etwa zwei Kräfte durch Vektoren dar: Die resultierende Kraft ergibt sich durch Addition der Vektoren. Analog für Geschwindigkeiten, einander überlagernde elektrische Feldstärken, usw. Besonders nützlich ist es, dass man eine geometrische Version der Vektoroperationen sowie eine rechnerische auf der Seite der Zahlentripel hat. (Wir erwähnen nicht jedes Mal eigens die Zahlenpaare.) Vor allem kann man die Vektorraumoperationen dazu benutzen, aus einfachen Elementen eine kompliziertere Situation aufzubauen, etwa aus dem Ortsvektor eines Kreispunktes den Ortsvektor eines Punktes eines weiteren Kreises gewinnen, der auf ersterem abrollt. Oder man möchte von einem Punkt aus zum neuen Punkt, der aus ersterem durch eine Projektion oder eine Spiegelung entsteht. Später werden Sie noch lernen, dass auch Funktionen Vektorräume bilden und dass der Addition von Funktionen überragende Bedeutung zukommt.

2.1. Vektoraddition und Multiplikation eines Vektors mit einer Zahl. Das sind die sogenannten linearen Vektorraumoperationen. Wir stellen sie in beiderlei Form gegenüber:

Multiplikation eines Vektors mit einer Zahl (einem Skalar)	
geometrisch ($\lambda \in \mathbb{R}, \vec{a} \in V^3$)	rechnerisch ($\lambda \in \mathbb{R}, \begin{pmatrix} x \\ y \\ z \end{pmatrix} \in \mathbb{R}^3$)
$\lambda \vec{a}$ entsteht aus $\vec{a}$ durch Streckung mit λ , wobei $\lambda < 0$ zusätzlich ein Umdrehen der Pfeilspitze bewirkt. Zum Beispiel:	$\lambda \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} \lambda x \\ \lambda y \\ \lambda z \end{pmatrix}$, z.B.:
	$-\frac{1}{2} \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} -1/2 \\ -1/2 \\ 0 \end{pmatrix}$

Addition zweier Vektoren:	
geometrisch ($\vec{a}, \vec{b} \in V^3$)	rechnerisch $\left(\begin{pmatrix} u \\ v \\ w \end{pmatrix}, \begin{pmatrix} x \\ y \\ z \end{pmatrix} \in \mathbb{R}^3 \right)$
$\vec{a} + \vec{b}$ entsteht durch Hintereinandersetzen bzw. Parallelogramm- ergänzung. Zum Beispiel:	$\begin{pmatrix} u \\ v \\ w \end{pmatrix} + \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} u+x \\ v+y \\ w+z \end{pmatrix}$
	

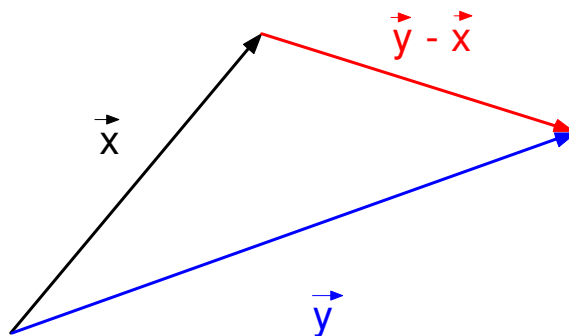
Entscheidend ist nun der Zusammenhang zwischen den rechnerischen Operationen einerseits und den geometrischen andererseits: Man beobachtet sogleich im Beispiel, wie sich bei geometrischer Addition auch die Koordinaten addieren. Das ist generell so. Für die Multiplikation mit einem Skalar wird der Zusammenhang vom Strahlensatz gesichert. Erst diese Sachverhalte garantieren, dass man etwa auch Kräfte richtig rechnerisch addiert und damit die Koordinatendarstellung der tatsächlich resultierenden Kraft erhält, usw. Wir fassen sie daher noch einmal ausdrücklich zu Formeln zusammen und unterscheiden dazu an dieser Stelle notwendigerweise auch einmal ausdrücklich die rechnerische von der geometrischen Operation durch Anfügen eines Index (r bzw. g). Zusätzlich verwenden wir einen Multiplikationspunkt, den man sonst gewöhnlich weglässt.

$$\begin{aligned}(\lambda \cdot_g \vec{x})^K &= \lambda \cdot_r (\vec{x}^K) \\ (\vec{x} +_g \vec{y})^K &= \vec{x}^K +_r \vec{y}^K\end{aligned}$$

Diese allgemeingültigen Formeln besagen gerade, dass die Rechenoperationen den geometrisch-physikalischen wirklich entsprechen.

Warnung: Es gibt noch eine andere „Operation“, die dem Anfänger vielfach nahe liegt: Das komponentenweise Multiplizieren zweier Vektoren zu einem neuen Vektor. Dieser Vektor ist zwar dann aufgrund einer wohldefinierten Operation klar definiert, nur ist er zu nichts Gutem zu gebrauchen. Damit wird stets Unsinn produziert. Gefährlich ist vor allem, das später einzuführende Skalarprodukt damit zu verwechseln.

Recht wichtig dagegen ist es, die Differenz zweier Vektoren zu bilden und eine unmittelbare Anschauung dafür zu gewinnen, obwohl man einfach lapidar $\vec{x} - \vec{y} = \vec{x} + (-\vec{y}) = \vec{x} + (-1)\vec{y}$ definieren kann und diese Operation daher nichts eigentlich Neues ist. Hier die Illustration:



Zunächst mache man sich am Bild einfach klar, dass $\vec{x} - \vec{y}$ allein schon deswegen so aussehen sollte, weil sicher gelten sollte: $\vec{y} + (\vec{x} - \vec{y}) = \vec{x}$. Die Differenz zweier Vektoren bildet man also, indem man ihre Anfänge aneinanderlegt und den zwischen die Pfeilspitzen eingespannten Vektor bildet. Man hat nur zu beachten, dass die Spitze des Differenzvektors an die des Vektors stößt, von dem abgezogen wird. Das bedeutet insbesondere, dass man mit der Differenz der Ortsvektoren $\vec{x}_Q - \vec{x}_P$ den Vektor von P nach Q bekommt, den man auch mit $\overrightarrow{PQ}$ bezeichnet. Man achte auf die Reihenfolgen!

2.2. Die abstrakten Rechenregeln für das Rechnen mit Vektoren (Vektorraumaxiome).

In Worten lässt sich knapp zusammenfassen: Mit Vektoren rechnet man wie mit Zahlen, einzig und allein das Dividieren durch einen Vektor entfällt (und ist auch nicht sinnvoll im allgemeinen Fall zu definieren!). Man kann jedoch auch einen vollständig ausreichenden Satz von Regeln (alle weiteren kann man daraus herleiten) knapp aufschreiben, hier ist er. Zu verstehen sind diese Formeln als allgemeingültige, d.h. man

darf für λ, μ jeden Skalar(ausdruck) und für $\vec{x}, \vec{y}, \vec{z}$ jeden Vektor(ausdruck) einsetzen, erhält stets wieder allgemeingültige Formeln. Der Bereich der Zahlen muss nicht $\mathbb{R}$ sein, es kann jeder Körper sein, der Fall der komplexen Zahlen ist besonders wichtig. Man spricht dann allgemein von einem Vektorraum über einem Körper K . Nun nennt man definitionsgemäß jede Struktur mit einer inneren Verknüpfung und einer äußeren mit den Elementen eines Körpers K genau dann einen Vektorraum über K , wenn diese Regeln erfüllt sind. Was hat man davon? Man verifiziert diese Regeln und weiß dann, dass auch alle Folgerungen gelten und man also „wie üblich“ bequem und korrekt rechnen kann.

DEFINITION 9. Eine Menge V mit einer zweistelligen Verknüpfung $+$ zusammen mit einem Körper (vgl. die Definition dafür durch die Körperaxiome in Kap. 1) K und einer äußeren Verknüpfung $\cdot$, welche jedem Paar $(\lambda, \vec{x})$, $\lambda \in K, \vec{x} \in V$, genau ein Element $\lambda \cdot \vec{x} \in V$ zuordnet, heißt Vektorraum über K , wenn folgende Rechengesetze (die Vektorraumaxiome) allgemein erfüllt sind:

$$\begin{aligned}\vec{x} + (\vec{y} + \vec{z}) &= (\vec{x} + \vec{y}) + \vec{z} \text{ (Assoziativität der Addition)} \\ \vec{0} + \vec{x} &= \vec{x} \text{ (neutrales Element, Nullvektor)} \\ -\vec{x} + \vec{x} &= \vec{0} \text{ (inverse Elemente bzgl. Addition)} \\ \vec{x} + \vec{y} &= \vec{y} + \vec{x} \text{ (Kommutativität)} \\ (\lambda\mu)\vec{x} &= \lambda(\mu\vec{x}) \text{ (Klammern überflüssig)} \\ (\lambda + \mu)\vec{x} &= \lambda\vec{x} + \mu\vec{x} \text{ (erstes Distributivgesetz)} \\ \lambda(\vec{x} + \vec{y}) &= \lambda\vec{x} + \lambda\vec{y} \text{ (zweites Distributivgesetz)} \\ 1 \cdot \vec{x} &= \vec{x} \text{ (Normierung der Mult. mit Zahl)}\end{aligned}$$

Die Elemente von V nennt man dann Vektoren, die von K Skalare (für: Zahlen).

Man sieht die zwei Gruppen: Erstere handelt nur von $+$, letztere auch von der Multiplikation mit einem Skalar, insbesondere von der Verbindung beider in den Distributivgesetzen. Ersetzen Sie einmal die Vektoren durch Zahlen, so sehen Sie Ihnen längst bekannte Rechengesetze.

Eigentlich müssten wir jetzt für unsere eingeführten Strukturen V^3 (analog V_O^3) sowie $\mathbb{R}^3$ zeigen, dass diese Axiome alle erfüllt sind. Damit erst ist klar, dass wir mit Recht von Vektoren gesprochen haben. Aber man beachte, dass diese Nachprüfung etwa für $\mathbb{R}^3$ sehr selbstverständlich ist, da sich jedes Axiom auf Entsprechendes innerhalb der Zahlen sofort zurückführt. Ein Beispiel:

$$\begin{aligned}\lambda \left(\begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} u \\ v \\ w \end{pmatrix} \right) &= \lambda \begin{pmatrix} x+u \\ y+v \\ z+w \end{pmatrix} = \begin{pmatrix} \lambda(x+u) \\ \lambda(y+v) \\ \lambda(z+w) \end{pmatrix} \\ &= \begin{pmatrix} \lambda x + \lambda u \\ \lambda y + \lambda v \\ \lambda z + \lambda w \end{pmatrix} = \begin{pmatrix} \lambda x \\ \lambda y \\ \lambda z \end{pmatrix} + \begin{pmatrix} \lambda u \\ \lambda v \\ \lambda w \end{pmatrix} \\ &= \lambda \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \lambda \begin{pmatrix} u \\ v \\ w \end{pmatrix}.\end{aligned}$$

Dabei wurden nur die Definitionen der linearen Operationen im $\mathbb{R}^3$ sowie das Distributivgesetz der reellen Zahlen verwandt. Analog funktioniert die Sache mit den andern Gesetzen. Wir erwähnen nur, dass wir im $\mathbb{R}^3$ haben:

$$\vec{0} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \quad - \begin{pmatrix} x \\ y \\ z \end{pmatrix} = \begin{pmatrix} -x \\ -y \\ -z \end{pmatrix}.$$

Wenn wir nunmehr die Rechengesetze für den Vektorraum $\mathbb{R}^3$ verifiziert haben, so können wir sofort schließen, dass sie auch für V^3 gelten: Denn die obenstehenden Formeln der Entsprechung zwischen rechnerischer und geometrischer Operation machen klar, dass man über K in gewisser Weise V^3 mit $\mathbb{R}^3$ identifizieren kann. (Entsprechen sich Elemente umkehrbar eindeutig und ebenfalls die Operationen, dann spricht man von isomorphen Strukturen, und man kann ganz allgemein zeigen, dass dann stets dieselben Rechengesetze gelten.) Dasselbe Argument gilt dann auch für V_O^3 . Andererseits ist es interessant, zu beobachten, dass ein Gesetz wie $\lambda(\vec{x} + \vec{y}) = \lambda\vec{x} + \lambda\vec{y}$ geometrisch immerhin so viel wie den

Strahlensatz bedeutet, das heisst: Wendet man dies Rechengesetz einfach formal an, so hat man eine sonst im allgemeinen viel kompliziertere geometrische Überlegung gespart. Insgesamt werden wir sehen: Kompliziertere geometrische Probleme lassen sich mit Vektorrechnung viel leichter lösen als ohne sie. Halten wir noch einmal unser Hauptresultat fest:

SATZ 1. $V^3, V_O^3, \mathbb{R}^3$ sind paarweise isomorphe Vektorräume über dem Körper der reellen Zahlen. Das entsprechende Resultat gilt für Dimension 2. Analog hat man auch die Räume $\mathbb{R}^n$ der sogenannten „ n -Tupel“ (also Folgen der Länge n , $n \in \mathbb{N}$, $n \geq 1$) von reellen Zahlen, in denen ebenfalls komponentenweise gerechnet wird.

Allerdings fehlen für $n \geq 4$ die zugehörigen anschaulichen Räume, wenn auch nicht der unmittelbare Bezug zur Physik (Raumzeit, aber auch mechanische Systeme mit riesigen Freiheitsgraden usw.)

2.3. Einige praktisch wichtige Folgerungen. Folgende Regeln sind herleitbar aus den Axiomen, man hat sie aber zweckmäßig im Kopf und benutzt sie ohne weiteres:

$$\begin{aligned} 0 \cdot \vec{x} &= \vec{0} \\ -(\lambda \vec{x}) &= (-\lambda) \vec{x} \\ \lambda(\vec{x} - \vec{y}) &= \lambda \vec{x} - \lambda \vec{y} \end{aligned}$$

Man beachte: Die zweite Gleichung ist nicht von vornherein selbstverständlich, sie besagt, dass der additiv inverse (negative) Vektor zu $\lambda \vec{x}$ dasselbe ist wie der Vektor $\vec{x}$ mit $(-\lambda)$ multipliziert (Inversenbildung im Körper der reellen Zahlen).

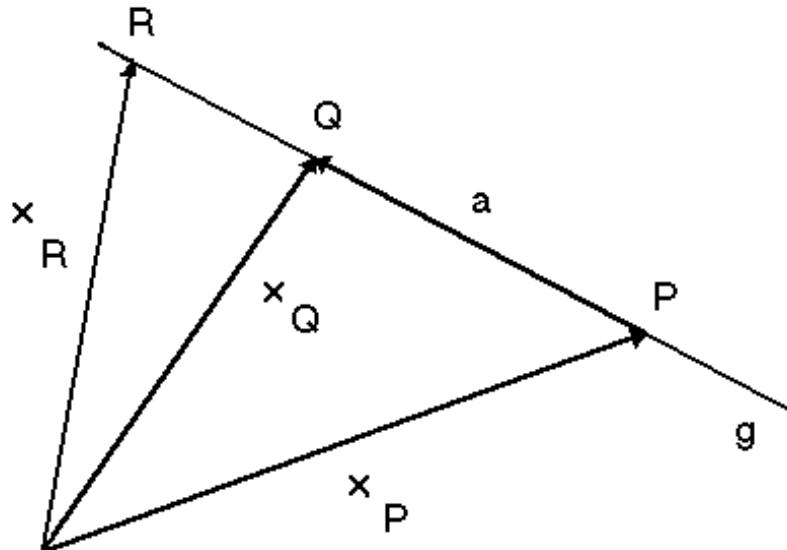
Eine immer wieder nützliche Anwendung der Rechenregeln ist diese Folgerung: Jeder Rechenausdruck der Vektorraumstruktur in den Vektoren $\vec{a}_1, \dots, \vec{a}_n$ lässt sich auf die Form $\sum_{i=1}^n \lambda_i \vec{a}_i = \lambda_1 \vec{a}_1 + \dots + \lambda_n \vec{a}_n$ bringen. Das ist dann auch die ordentliche Endform eines solchen Ausdrucks, die man für viele Zwecke benötigt und stets produzieren sollte. Man nennt diesen Ausdruck eine **Linearkombination der Vektoren** $\vec{a}_1, \dots, \vec{a}_n$. Beispiel:

$$2(\vec{x} + \lambda \vec{y}) - 3\vec{x} + \mu \frac{\vec{y}}{3} = -\vec{x} + \left(2\lambda + \frac{\mu}{3}\right) \vec{y}.$$

Hier wären also $n = 2$, $\vec{a}_1 = \vec{x}$, $\vec{a}_2 = \vec{y}$, $\lambda_1 = -1$, $\lambda_2 = 2\lambda + \frac{\mu}{3}$. Beachten Sie, wie man hierzu insbesondere die beiden Distributivgesetze benötigt (und zwar jeweils in beiden Richtungen gelesen, das eine Mal Ausmultiplizieren, das andere Mal Ausklammern!).

2.4. Parameterdarstellungen für Geraden und Ebenen. Mittels der Vektorrechnung gelingt es leicht, die oben bereits angesprochene Parameterdarstellung für Geraden und Ebenen geometrisch zu verstehen.

Sei g eine Gerade im E^3 . Dann kann man g festlegen durch einen Punkt und eine Richtung. Die Richtung von g kann wiederum durch einen freien Vektor $\vec{a} \neq \vec{0}$ parallel zu g bestimmt werden. Sei also P ein fester Punkt auf g und $\vec{a}$ ein solcher Richtungsvektor für g . Betrachten wir folgende Skizze:



Daraus wird klar: Wenn man $\vec{x}_P + \lambda \vec{a}$ bildet mit dem Ortsvektor $\vec{x}_P$ von P und einer beliebigen Zahl $\lambda \in \mathbb{R}$, so erhält man stets einen Ortsvektor eines Punktes der Geraden. (Beachte Konvention: Ortsvektor plus freier (Richtungs-) Vektor gleich Ortsvektor!) Zum Beispiel ist in der Skizze $\vec{x}_P + 1 \cdot \vec{a} = \vec{x}_Q$ und etwa $\vec{x}_P + \frac{3}{2} \cdot \vec{a} = \vec{x}_R$. Umgekehrt erhält man zu jedem Punkt S der Geraden eine eindeutig bestimmte Zahl $\lambda(S)$, so dass $\vec{x}_S = \vec{x}_P + \lambda(S) \vec{a}$ ist. Die Menge aller Ortsvektoren von Geradenpunkten wird also durchlaufen mit

$$\vec{x}_g(\lambda) = \vec{x}_P + \lambda \vec{a}, \quad \lambda \in \mathbb{R}. \quad (\lambda: \text{Freier Parameter}),$$

(Parameterdarstellung einer Geraden g , geometrische Form.)

$$\vec{x}_g^K(\lambda) = \vec{x}_P^K + \lambda \vec{a}^K, \quad \lambda \in \mathbb{R}. \quad (\text{Koordinatenform.})$$

Dabei sind : $\vec{x}_P$ Aufpunktvektor,
 $\vec{a} \neq \vec{0}$ Richtungsvektor.

Man beachte hier die beiden Formen. $\vec{x}_g(\lambda)$ ist der geometrische Pfeil „Ortsvektor des Geradenpunktes zum Parameterwert von λ “. Er ergibt sich durch die geometrische Vektoroperationen als $\vec{x}_P + \lambda \vec{a}$. Diese Form hatten wir für unsere allgemeine anschaulich-geometrische Einsicht benutzt. Dagegen werden wir für Fragen wie nach dem Schnittpunkt einer Geraden mit der xy -Ebene usw. die Koordinatenform benutzen. Dabei rechnet man rein formal mit Zahlentripeln.

Ein Beispiel dazu: Geben wir die Gerade g mit folgender Parameterdarstellung in Koordinatenform vor:

$$\vec{x}_g^K(\lambda) = \begin{pmatrix} 2 \\ -1 \\ 2 \end{pmatrix} + \lambda \begin{pmatrix} -3 \\ 2 \\ 2 \end{pmatrix}, \quad \lambda \in \mathbb{R}.$$

Suchen wir den Schnitt mit der xy -Ebene, so ergibt sich die Bedingung

$$\begin{pmatrix} 2 \\ -1 \\ 2 \end{pmatrix} + \lambda \begin{pmatrix} -3 \\ 2 \\ 2 \end{pmatrix} = \begin{pmatrix} x \\ y \\ 0 \end{pmatrix}.$$

Man beachte: Zwei Vektoren aus $\mathbb{R}^3$ sind genau dann gleich, wenn sie in allen drei Komponenten übereinstimmen. Unsere Vektorgleichung läuft also auf folgendes System von drei Zahlengleichungen hinaus:

$$\begin{aligned} 2 - 3\lambda &= x \\ -1 + 2\lambda &= y \\ 2 + 2\lambda &= 0 \end{aligned}$$

Das ist ein lineares Gleichungssystem in den Unbestimmten λ, x, y . Aber es ist besonders einfach (zur Systematik mehr im nächsten Abschnitt). Aus der letzten Gleichung liest man sofort $\lambda = -1$ ab, und mit beiden andern Gleichungen sind dann x, y auch eindeutig bestimmt. Wir erhalten also einen einzigen Schnittpunkt, es ist der Geradenpunkt zum Parameterwert $\lambda = -1$. Der gesuchte Schnitt besteht also aus dem einzigen Punkt S mit

$$\vec{x}_S^K = \begin{pmatrix} 2 \\ -1 \\ 2 \end{pmatrix} + (-1) \begin{pmatrix} -3 \\ 2 \\ 2 \end{pmatrix} = \begin{pmatrix} 5 \\ -3 \\ 0 \end{pmatrix}.$$

Man wird in solchen Zusammenhängen später den Index K weglassen - man sieht ja, dass es sich um eine Koordinatendarstellung handelt.

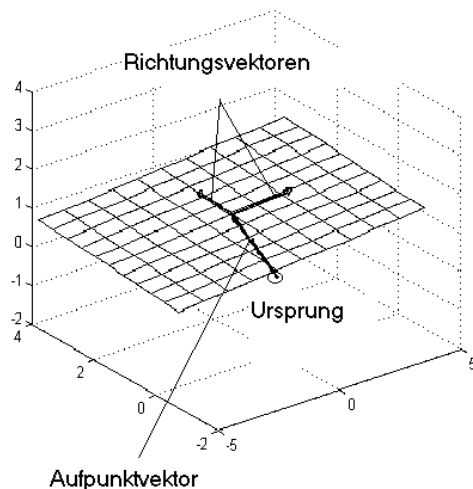
Es sollte deutlich geworden sein, dass die Spaltenschreibweise für die Zahlentripel für solche Rechnungen besonders übersichtlich ist. (Noch einmal: Wenn wir hier Zeilenvektoren schreiben, dann nur aus Gründen der Platzersparnis.)

Mit diesen Mitteln können wir auch sofort einsehen, dass so etwas wie

$$\vec{x}(\lambda) = \begin{pmatrix} 2 - 3\lambda \\ -1 + 2\lambda \\ -3\lambda \end{pmatrix} = \begin{pmatrix} 2 \\ -1 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -3 \\ 2 \\ -3 \end{pmatrix}, \lambda \in \mathbb{R},$$

eine Parameterdarstellung einer Geraden ist. Die ausgeführte Umformung macht die Aufspaltung der Standardform deutlich, und wir können nunmehr Aufpunkt- und Richtungsvektor identifizieren und die Gerade uns vorstellen sowie in ein Koordinatensystem einzeichnen.

Nun das entsprechende Programm für Ebenen. Wieder startet die Argumentation mit der geometrischen Vektorrechnung: Geben wir einen Punkt P einer Ebene E im E^3 vor sowie zwei unabhängige (d.h. hier: beide ungleich Null und nicht parallel) Richtungsvektoren $\vec{a}, \vec{b}$ parallel zur Ebene. Folgende Skizze zeigt dann, dass man mit $\vec{x}_P + \lambda \vec{a} + \mu \vec{b}$, $\lambda, \mu \in \mathbb{R}$, stets einen Ortsvektor eines Punktes von E erhält und dass man umgekehrt zu jedem Punkt Q aus E auf diese Weise (sogar eindeutig) seinen Ortsvektor bekommt. Denn man kann stets auf der Ebene ein passendes Parallelogramm mit Seiten parallel zu $\vec{a}, \vec{b}$ ansetzen, so dass eine Ecke in P , die andere im gewünschten Punkt Q landet.



Also hat man ganz analog zu den Geraden folgende Parameterform für eine Ebene mit zwei freien Parametern:

$$\vec{x}_E(\lambda, \mu) = \vec{x}_P + \lambda \vec{a} + \mu \vec{b}, \quad \lambda, \mu \in \mathbb{R}, \quad \vec{a}, \vec{b} \text{ linear unabhängige Richtungsvektoren, } \vec{x}_P \text{ Aufpunktvektor.}$$

Entsprechend hat man wiederum die Koordinatenform mit Zahlentripeln. Zwei Vektoren $\vec{a}, \vec{b}$ sind definitionsgemäß genau dann linear unabhängig, wenn keiner ein Vielfaches des andern ist, also wenn

sie beide nicht Null sind und nicht parallel sind, also verschiedene Richtungen angeben. (Dies ist ein Spezialfall des später einzuführenden allgemeinen Begriffes der linearen Unabhängigkeit.) Nun sollte klar sein, dass man eine Parameterdarstellung für eine gewünschte Ebene (analog: Gerade) einfach aufstellt, indem man dies Schema ausfüllt mit einem Aufpunktvektor und einem Richtungsvektor.

Beispiele: Eine Parameterdarstellung ist gesucht für die Ebene E , welche durch P geht, $\vec{x}_P^K = (1, 2, -2)$, und parallel zur yz -Ebene liegt. Offenbar sind $(0, 1, 0)$ und $(0, 0, 1)$ passende Richtungsvektoren, und damit können wir eine Parameterdarstellung für E in Koordinatenform sofort hinschreiben:

$$\vec{x}_E(\lambda, \mu) = (1, 2, -2) + \lambda(0, 1, 0) + \mu(0, 0, 1), \quad \lambda, \mu \in \mathbb{R}.$$

Natürlich gibt es unendlich viele Parameterdarstellungen für *dieselbe* Gerade/Ebene, in unserm Beispiel wäre auch

$$\vec{y}_E(\lambda, \mu) = (1, 3, -2) + \lambda(0, 1, 1) + \mu(0, -1, 1), \quad \lambda, \mu \in \mathbb{R},$$

eine andere Parameterdarstellung derselben Ebene. (Wie könnte man das nachprüfen?)

Weiteres Beispiel: Wir suchen eine Parameterdarstellung der Ebene H , welche durch die drei Punkte P, Q, R geht mit $\vec{x}_P^K = (1, 2, -2)$, $\vec{x}_Q^K = (1, 4, 2)$, $\vec{x}_R^K = (-2, 1, 3)$. Wiederum finden wir leicht einen Aufpunktvektor. Wir wählen den ersten dazu. Richtungsvektoren erkennen wir leicht als Differenzen zwischen den Ortsvektoren. Im Rahmen der Dreiecksituation liegt es nahe, dazu die Vektoren $\vec{PQ} = \vec{x}_Q - \vec{x}_P$ und $\vec{PR} = \vec{x}_R - \vec{x}_P$ zu wählen. Damit erhalten wir folgende Parameterdarstellung für H :

$$\vec{x}_H(\lambda, \mu) = (1, 2, -2) + \lambda(0, 2, 4) + \mu(-3, -1, 5), \quad \lambda, \mu \in \mathbb{R}.$$

Rechentchnik: Im Zusammenhang mit den Parameterdarstellungen von Geraden und Ebenen ist folgende typische Rechnung *in beiden Richtungen* wichtig:

$$\begin{pmatrix} 1 + \lambda - 2\mu \\ 2\lambda + 3\mu \\ 2 + 3\lambda \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 2 \end{pmatrix} + \lambda \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + \mu \begin{pmatrix} -2 \\ 3 \\ 0 \end{pmatrix}.$$

Es handelt sich lediglich (von rechts nach links) um die Ausführung der linearen Operationen mit den Zahlentripeln. In der andern Richtung ist es das „Auseinanderziehen“, und das macht man mit dem allenthalben wichtigen Vergleichen von Koeffizienten. Der Plan ist eine Zerlegung wie auf der rechten Seite, also fasst man alle konstanten Glieder zu einem Vektor (Fehlen: Koeffizient Null!), dann alle Koeffizienten bei λ usw. Nunmehr können wir auch die linke Seite, die aus einer Rechnung resultieren könnte, sofort als Ausdruck einer Parameterdarstellung einer Ebene verstehen und mit der rechten Seite diese Ebene geometrisch in ihrer Lage verstehen. Die andere Richtung braucht man etwa beim Aufstellen eines linearen Gleichungssystems, um den Schnitt zweier linearer Gebilde zu berechnen, welche durch Parameterdarstellungen gegeben sind. Schneiden wir g mit E , wobei $\vec{x}_g(\alpha) = (1, 1, 1) + \alpha(2, 1, 1)$, $\alpha \in \mathbb{R}$, $\vec{y}_E(\lambda, \mu) = \lambda(2, -1, 1) + \mu(1, 2, 2)$, $\lambda, \mu \in \mathbb{R}$, so haben wir folgende Gleichung zu lösen:

$$\vec{x}_g(\alpha) = \vec{y}_E(\lambda, \mu)$$

(„Gleichsetzen“ - dabei ist auf Trennung der Parameter aus den verschiedenen Parameterdarstellungen zu achten, ein Parameter λ in $\vec{x}_g$ hätte umbenannt werden müssen). Im Beispiel sollte man sofort das daraus entstehende lineare Gleichungssystem überblicken:

$$\begin{aligned} 2\lambda + \mu - 2\alpha &= 1 \quad (I) \\ -\lambda + 2\mu - \alpha &= 1 \quad (II) \\ \lambda + 2\mu - \alpha &= 1 \quad (III). \end{aligned}$$

Nunmehr ist dies zu lösen. Man achte auf die *im allgemeinen* sinnvolle Technik: Es wird das lineare Gleichungssystem durch Kombination der Zeilen in jedem Schritt um eine Gleichung und eine Unbekannt „kürzer“ gemacht. Wir werden hier λ, μ hinauswerfen (warum?). Zunächst λ , man erhält:

$$\begin{aligned} 5\mu - 4\alpha &= 3, \quad (I') = (I) + 2(II) \\ 2\mu - \alpha &= 1, \quad (II') = ((II) + (III))/2. \end{aligned}$$

Nunmehr kann man μ hinauswerfen:

$$3\alpha = -1, \quad 5(II') - 2(I').$$

Das ergibt $\alpha = -1/3$, und das reicht bereits aus, die gestellte Aufgabe zu lösen: Einsetzen in die Parameterdarstellung der Geraden ergibt den einzigen Schnittpunkt $\vec{x}_S = \vec{x}_g(-1/3) = (1/3, 2/3, 2/3)$. Man kann und sollte allerdings zur Kontrolle noch durch Einsetzen die Werte der andern Parameter ausrechnen: $\lambda = 0, \mu = 1/3$. Die einzige Lösung des Gleichungssystems lautet also $(0, 1/3, -1/3)$ (zur Reihenfolge μ, λ, α). Das ist nicht etwa die Koordinatendarstellung des Schnittpunktes, sondern man muss diese gewinnen, indem man $\alpha = -1/3$ in die Parameterdarstellung der Geraden oder aber $\lambda = 0, \mu = 1/3$ in die Parameterdarstellung der Ebene einsetzt. Beide Male kommt dasselbe heraus, wie man überprüfe. Damit ist eine wirksame Kontrolle gemacht. Wir werden im nächsten Abschnitt 5. über Schnittpunktaufgaben und lineare Gleichungssysteme allerdings noch sehen, dass die Gleichungsform für solche Aufgaben wesentlich günstiger ist, auch eine Kombination Gleichung/Parameterdarstellung.

2.5. Behandlung des Dreiecksschwerpunktes mit koordinatenfreier Vektorrechnung. Das rein geometrische Rechnen mit den geometrischen Vektorraumoperationen insbesondere in V^2, V^3 hat für gewisse Aufgaben eine Eleganz und Einfachheit, die sich mit dem Rechnen mittels Koordinatendarstellungen bei weitem nicht erreichen lässt. Dazu zählt insbesondere das Herausarbeiten bzw. Beweisen allgemeiner Sachverhalte, die jedoch für Anwendungen durchaus konkrete Bedeutung haben. (Es ist ein Irrglaube, unter „theoretisch“ dasselbe wie „uninteressant, da für praktische Anwendungen irrelevant“ zu verstehen.) Ein erstes Beispiel haben wir schon mit der allgemeinen Überlegung gegeben, dass man eine Parameterdarstellung für die Ebene E durch drei nicht auf einer Geraden liegende Punkte P, Q, R sofort angeben kann mit $\vec{x}_E(\lambda, \mu) = \vec{x}_P + \lambda(\vec{x}_Q - \vec{x}_P) + \mu(\vec{x}_R - \vec{x}_P)$, $\lambda, \mu \in \mathbb{R}$. Dagegen ist die konkrete Ausführung für Beispiele reichlich langweilig (!). Wir geben ein prägnantes weiteres Beispiel, dazu das ungemein wichtige Prinzip des Koeffizientenvergleichs, das sich in unendlich vielen Aufgaben nutzen lässt:

Die Aufgabe. Zu drei Punkten P, Q, R , welche nicht auf einer Geraden liegen, ist der Schnittpunkt S der Seitenhalbierenden des zugehörigen Dreiecks allgemein zu berechnen. Wir werden dabei auch das Teilungsverhältnis der Strecken sehen, in welchem S die seitenhalbierenden Strecken teilt, und wir werden auch zu der bereits vorgestellten Darstellung von S als Schwerpunkt kommen. Dazu bilden wir zunächst die *freien* Vektoren $\vec{a} = \overrightarrow{PQ}$ und $\vec{b} = \overrightarrow{PR}$. Diese Vektoren sind linear unabhängig nach unserer Voraussetzung eines nicht ausgearteten Dreiecks. Das hat eine Konsequenz von allgemeiner Bedeutung: Man kann Koeffizientenvergleich machen, d.h.

$$\begin{aligned} \text{Aus } \alpha \vec{a} + \beta \vec{b} &= \lambda \vec{a} + \mu \vec{b} \text{ folgt stets („Koeffizientenvergleich“)} \\ \alpha &= \lambda, \beta = \mu, \text{ wenn } \vec{a}, \vec{b} \text{ linear unabhängig sind.} \end{aligned}$$

Das können wir sehr leicht begründen: Sei etwa $\alpha \neq \lambda$, aber $\alpha \vec{a} + \beta \vec{b} = \lambda \vec{a} + \mu \vec{b}$. Dann ist $(\alpha - \lambda) \vec{a} = (\mu - \beta) \vec{b}$, also mit $\alpha \neq \lambda$ $\vec{a} = \vec{b}(\mu - \beta)/(\alpha - \lambda)$. Somit wäre $\vec{a}$ ein Vielfaches von $\vec{b}$, also $\vec{a}, \vec{b}$ nicht linear unabhängig.

Nunmehr nennen wir den Schwerpunkt S , arbeiten jedoch mit dem *freien* Vektor $\vec{s} = \overrightarrow{PS}$. Diesen Vektor stellen wir auf zwei Arten dar, wir haben, wenn wir die seitenhalbierenden Geraden durch $\overrightarrow{PQ}$ und $\overrightarrow{PR}$ betrachten:

$$\vec{s} = \frac{1}{2} \vec{a} + \lambda \left(\vec{b} - \frac{1}{2} \vec{a} \right) = \frac{1}{2} \vec{b} + \mu \left(\vec{a} - \frac{1}{2} \vec{b} \right).$$

Schreiben wir beide Seiten als Linearkombinationen ordentlich hin:

$$\left(\frac{1}{2} - \frac{\lambda}{2} \right) \vec{a} + \lambda \vec{b} = \mu \vec{a} + \left(\frac{1}{2} - \frac{\mu}{2} \right) \vec{b},$$

so erhalten wir mit Koeffizientenvergleich sofort das einfache lineare Gleichungssystem

$$\begin{aligned} \left(\frac{1}{2} - \frac{\lambda}{2} \right) &= \mu \\ \lambda &= \left(\frac{1}{2} - \frac{\mu}{2} \right). \end{aligned}$$

Die eindeutige Lösung ist $\lambda = \mu = 1/3$. Also haben wir die Lösung des Problems, letztlich auch mittels der ursprünglichen Ortsvektoren ausgedrückt:

$$\vec{s} = \frac{1}{3} \vec{a} + \frac{1}{3} \vec{b}, \text{ und } \vec{x}_S = \vec{x}_P + \vec{s} = \frac{1}{3} (\vec{x}_P + \vec{x}_Q + \vec{x}_R).$$

(Überprüfen Sie, dass dies Resultat mittels Einführung von Koordinatendarstellungen bei weitem nicht so leicht zu gewinnen wäre.) Unser Resultat für den Schwerpunkt sollte auch intuitiv einleuchten: Der Schwerpunkt ergibt sich einfach rechnerisch über arithmetische Mittelbildung. Wir werden die Verallgemeinerung erwarten, dass der Schwerpunkt eines Systems von n Massenpunkten P_i , $1 \leq i \leq n$, in denen jeweils dieselbe Masse sitzt, ganz entsprechend zu bilden ist:

$$\vec{x}_S = \frac{1}{n} \sum_{i=1}^n \vec{x}_{P_i}.$$

2.6. Verschieben und Strecken mittels Vektorrechnung. Wir nennen sogleich die Hauptsache:

SATZ 16. 1.) Eine Menge von Punkten $\mathcal{M}$ wird mit dem Vektor $\vec{a}$ (fest) verschoben, indem man

$$T_{\vec{a}}(\mathcal{M}) = \{P \mid \vec{x}_P = \vec{a} + \vec{x}_Q \text{ für ein } Q \in \mathcal{M}\}$$

bildet. Die Operation der Verschiebung mit $\vec{a}$ nennen wir $T_{\vec{a}}$ (in Worten: Translation mit dem Vektor $\vec{a}$). Besonders einfach wird das mit parametrisierter Menge $\mathcal{M}$: Wenn

$$\vec{x}(\lambda, \dots) \text{ eine Parametrisierung für } \mathcal{M} \text{ ist,}$$

dann ist

$$\vec{y}(\lambda, \dots) := \vec{a} + \vec{x}(\lambda, \dots) \text{ (mit demselben Bereich für den oder die Parameter)}$$

eine Parametrisierung für $T_{\vec{a}}(\mathcal{M})$.

2.) Eine Menge von Punkten $\mathcal{M}$ wird mit Zentrum O und Faktor κ gestreckt ($\kappa > 0$), indem man bildet:

$$Z_{O;\kappa}(\mathcal{M}) := \{P \mid \vec{x}_P = \kappa \vec{x}_Q \text{ für ein } Q \in \mathcal{M}\}.$$

Wieder sieht man mit Parametrisierung $\vec{x}(\lambda, \dots)$ für $\mathcal{M}$ besonders einfach aus:

$$\vec{u}(\lambda, \dots) := \kappa \vec{x}(\lambda, \dots)$$

ist dann eine Parametrisierung der gestreckten Menge.

Beispiele (Anwendungen auf Kreise):

Wir wissen dass

$$\vec{x}(t) = \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}, \quad 0 \leq t < 2\pi,$$

den Kreis (**genauer die Kreiskurvenbahn**) um O mit Radius 1 parametrisiert (vorauszusetzen ist ein kartesisches System!). Genau genommen haben wir es hier mit den Koordinatendarstellungen der zu gehörigen Ortsvektoren zu tun.

Nun wollen wir einen Kreis um den Punkt $\begin{pmatrix} 2 \\ 5 \end{pmatrix}$ mit dem Radius 5 parametrisieren: Nach dem Satz bilden wir

$$\vec{y}(t) = \begin{pmatrix} 2 \\ 5 \end{pmatrix} + 5 \vec{x}(t) = \begin{pmatrix} 2 \\ 5 \end{pmatrix} + 5 \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}, \quad 0 \leq t < 2\pi.$$

(Zuerst müssen wir den Kreis um O aufblasen und dann die Verschiebung anbringen. Wenn wir die Reihenfolge ändern, bekommen wir:

$$\begin{aligned} \vec{u}(t) &= 5 \left(\begin{pmatrix} 2 \\ 5 \end{pmatrix} + \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix} \right) \\ &= \begin{pmatrix} 10 \\ 25 \end{pmatrix} + 5 \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}, \quad 0 \leq t < 2\pi, \end{aligned}$$

also den Kreis mit Mittelpunkt $\begin{pmatrix} 10 \\ 25 \end{pmatrix}$ und dem Radius 5.)

Wenn wir um das Zentrum P strecken wollen mit Faktor κ , so legen wir uns das so zurecht: Wir verschieben zuerst mit dem Vektor $-\vec{x}_P$, strecken dann mit Zentrum O , und Faktor κ anschließend verschieben wir wieder mit dem Vektor $\vec{x}_P$. Führen wir das aus mit dem Kreis um M mit Radius R , der parametrisiert ist mit

$$\vec{x}(t) = \vec{x}_M + R \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}, \quad 0 \leq t < 2\pi.$$

Damit führen wir die beschriebenen Operationen aus und bekommen:

$$\begin{aligned}\vec{y}(t) &= \kappa \left(\vec{x}_M - \vec{x}_P + R \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix} \right) + \vec{x}_P \\ &= \kappa \vec{x}_M + (1 - \kappa) \vec{x}_P + \kappa R \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}, \quad 0 \leq t < 2\pi.\end{aligned}$$

Daraus können wir das Resultat ablesen: Streckt man den Kreis um M mit Zentrum P und Faktor κ , so kommt der Kreis heraus mit Radius κR und Mittelpunkt $\vec{x}_M + (1 - \kappa) \vec{x}_P$.

Wir legen uns zum besseren Verständnis die Sache noch auf andere Weise zurecht: Streckt man mit κ um das Zentrum P , so werden die Vektoren $\vec{x}(t) - \vec{x}_P$ gestreckt, zu addieren ist dann $\vec{x}_P$. Wieder kommt:

$$\vec{y}(t) = \vec{x}_P + \kappa (\vec{x}(t) - \vec{x}_P) = (1 - \kappa) \vec{x}_P + \kappa \vec{x}_M + \kappa R \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}.$$

Kreisfläche statt nur Randkurve: Man parametrisiert die Fläche des Einheitskreises mit Rand um O :

$$\vec{x}(r, t) = r \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}, \quad 0 \leq r \leq 1, \quad 0 \leq t < 2\pi.$$

(Klar, zwei freie Parameter werden benötigt.) Nun kann man völlig analog damit operieren wie oben mit der Kurve. Man beachte aber auch, dass es nichts ausmacht, ob es sich um Parabeln oder Ellipsen handelt oder um Hyperboloide, die Operationen sind immer dieselben.

Bemerkung: Ebenso leicht kann man Viertelkreise usw. nehmen, dazu braucht man nur den Parameterbereich für den Winkel (oben: t) entsprechend einzuschränken.

2.7. Schnitte linearer Gebilde und lineare Gleichungssysteme. Im Titel sind die Zusätze „linear“ zu beachten - nur dafür kann man ein so einfaches geschlossenes Bild sowie völlig problemlose Lösungsmethoden erhalten, nicht so im nichtlinearen Fall. Wir beginnen mit einer Reihe konkreter Beispiele. Darin finden sich auch nichtlineare, die man in besonders einfachen Fällen durchaus auch von Hand rechnend bewältigt. Insbesondere soll hervorgehoben werden, dass die *Logik* der Schnittbildung im linearen wie nichtlinearen Fall ganz dieselbe ist. Nur werden im allgemeinen Fall sowohl Rechnung als auch theoretischer Überblick über Lösungsmengen extrem viel schwieriger. Vorab sei einmal klargestellt, was genau unter dem Schnitt zweier geometrischer Figuren, also Punktmengen F und G zu verstehen ist:

DEFINITION 10. *Seien A, B beliebige Mengen. Dann ist der Durchschnitt (oder die Schnittmenge) von A und B die folgende Menge (gelesen: „ A geschnitten mit B “): $A \cap B := \{x \mid x \in A \text{ und } x \in B\}$. Der Schnitt zweier geometrischer Figuren ist einfach der Durchschnitt der Figuren als Punktmengen.*

Beispiele: $\{1, 2, 3, 4\} \cap \{-1, 2, 3, 5, 6\} = \{2, 3\}$. Der Durchschnitt eines Kegelmantels und einer Ebene ist einer der Kegelschnitte: Ellipse, Parabel, Hyperbel.

2.8. Beispiele (ausschließlich linear). Hier sollte man zunächst die Logik wahrnehmen, die differenziert ist je nach dem, ob beide Gebilde in Parameterdarstellung, beide in Form von Bestimmungsgleichungen (bzw. bereits Systemen davon) oder das eine in der einen, das andere in der andern Form gegeben ist. Zu dieser Logik gehört auch, wie man von der technischen Lösung eines Gleichungssystems zur Darstellung der Schnittmenge gelangt. Schließlich ist das Verständnis wichtig, dass der Schnitt zweier linearer Gebilde stets ein lineares Gleichungssystem hervorbringt, dessen Dimensionierung (Zahl der Bedingungen und der Unbekannten) man aus der Aufgabenstellung vorhersagen können sollte. Gerade an der Dimensionierung sollte man feststellen, dass Gleichungsform oder Mischform für Schnittaufgaben wesentlich günstiger sind als Parameterformen. Der zweite Hauptpunkt des Interesses ist dann natürlich auch die Technik des Lösen der resultierenden Gleichungssysteme.

2.8.1. Beide (linearen) Gebilde in Parameterform. (Vorbemerkung: Es handelt sich hier um eine sehr unpraktische Art, zwei Ebenen miteinander zu schneiden, es geht jedoch um die Systematik des Lösen linearer Gleichungssysteme und um die Logik der Schnittaufgaben.) Wir schneiden die Ebenen E und F , beide in Parameterform gegeben durch

$$\begin{aligned}\vec{x}_E(\alpha, \beta) &= (1, 1, 1) + \alpha(1, 2, -1) + \beta(2, 1, 3), \quad \alpha, \beta \in \mathbb{R}, \\ \vec{x}_F(\lambda, \mu) &= (2, 1, 2) + \lambda(-1, 2, 2) + \mu(2, -2, 3), \quad \lambda, \mu \in \mathbb{R}.\end{aligned}$$

Dabei ist zunächst zu beachten, dass die Parameter beider Darstellungen getrennt sind. Andernfalls muss man umbenennen. Nun zur Logik: Nehmen wir einen beliebigen Schnittpunkt $S \in E \cap F$, so haben wir folgende Aussagen:

1. Es gibt $\alpha, \beta \in \mathbb{R}$, so dass $\vec{x}_S = \vec{x}_E(\alpha, \beta)$.
2. Es gibt $\lambda, \mu \in \mathbb{R}$, so dass $\vec{x}_S = \vec{x}_F(\lambda, \mu)$.

Das bedeutet zusammen, dass wir eine Lösung der Gleichung

$$(3) \quad \vec{x}_E(\alpha, \beta) = \vec{x}_F(\lambda, \mu)$$

haben. Das besagt aber nichts etwa darüber, dass z.B. $\alpha = \lambda$ sein sollte. (Dies war der logische Grund für die Parametertrennung.) Umgekehrt gehört zu jeder solchen Lösung ein Schnittpunkt, den wir dann durch einen der Rechenausdrücke dieser Gleichung darstellen können. Nun lösen wir Gleichung (3) im Beispiel. Sie lautet konkret

$$\begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} + \alpha \begin{pmatrix} 1 \\ 2 \\ -1 \end{pmatrix} + \beta \begin{pmatrix} 2 \\ 1 \\ 3 \end{pmatrix} = \begin{pmatrix} 2 \\ 1 \\ 2 \end{pmatrix} + \lambda \begin{pmatrix} -1 \\ 2 \\ 2 \end{pmatrix} + \mu \begin{pmatrix} 2 \\ -2 \\ 3 \end{pmatrix}$$

Dies sollte man *unmittelbar* zu folgendem linearen Gleichungssystem ummünzen können (gleiche Unbekannte geordnet, untereinander, die Konstanten isoliert auf der anderen Seite):

$$\begin{array}{rrrrrr} (I) & \alpha & +2\beta & +\lambda & -2\mu & = 1 \\ (II) & 2\alpha & +\beta & -2\lambda & +2\mu & = 0 \\ (III) & -\alpha & +3\beta & -2\lambda & -3\mu & = 1 \end{array}$$

Dies System hat drei Bedingungen für vier Unbestimmte, es ist also im Normalfall eine unendliche Lösungsmenge mit einem freien Parameter zu erwarten. Strategisch verläuft die Lösung so, dass man jeweils eine Unbekannte hinauswirft und eine Gleichung weniger bekommt. Hinsichtlich unserer Aufgabenstellung genügt es, auf eine Gleichung zu kommen, in der nur noch λ, μ vorkommen. (α, β wäre ebenso gut, nicht etwa α, λ .) Wir werfen α hinaus und erhalten:

$$\begin{array}{rrrrrr} (I') & = 2(I) - (II) & 3\beta & +4\lambda & -6\mu & = 2 \\ (II') & = (I) + (III) & 5\beta & -\lambda & -5\mu & = 2 \end{array}$$

Nunmehr β hinaus:

$$(I'') = 5(I') - 3(II') \quad 23\lambda \quad -15\mu \quad = 4$$

Der zweite Schritt im Lösen eines solchen Gleichungssystems besteht nach dieser Arbeit darin, die Lösungsmenge aufzuschreiben. Wir beobachten: Zu jedem beliebig vorgegebenen Wert von λ können wir eindeutig μ, β, α ausrechnen. Unsere Lösungsmenge hat also hier einen freien Parameter. (Warum ist es günstiger, λ dafür zu wählen als μ ?) Wir haben also alle andern Unbestimmten durch λ auszudrücken und erhalten aus $(I''), (I'), (I)$ der Reihe nach:

$$\mu = \frac{23}{15}\lambda - \frac{4}{15}, \quad \alpha = -\frac{7}{5}\lambda + \frac{1}{5}, \quad \beta = \frac{26}{15}\lambda + \frac{2}{15}.$$

Die Lösungsmenge in Parameterform aufgeschrieben lautet damit (unter Beibehaltung der Reihenfolge $\alpha, \beta, \lambda, \mu$):

$$L(\lambda) = \frac{1}{5}(1, 2, 0, -4) + \lambda \left(-\frac{7}{5}, \frac{26}{15}, 1, \frac{23}{15} \right), \quad \lambda \in \mathbb{R}.$$

Dies kann man geometrisch interpretieren als eine Gerade im $\mathbb{R}^4$, und wir werden sehen, dass jede Lösungsmenge eines linearen Gleichungssystems entweder leer ist oder aber geometrisch als lineares Gebilde (Verallgemeinerung von Punkt, Gerade, Ebene,...) zu deuten ist und eine entsprechende Parameterdarstellung besitzt. Aber man verwechsle das in unserem Fall nicht mit der gesuchten Schnittmenge! Um diese zu erhalten, hätten wir mit (I'') und der Beziehung $\mu = (23\lambda - 4)/15$ aufhören können. Das haben wir einzusetzen (so stets der letzte Schritt bei Schnittpunktaufgaben mit Parameterformen) in den Ausdruck

$\vec{x}_F(\lambda, \mu)$. Also ist die Schnittmenge, schon in Parameterdarstellung:

$$\begin{aligned}\vec{x}_{g_S}(\lambda) &= (2, 1, 2) + \lambda(-1, 2, 2) + \left(\frac{23}{15}\lambda - \frac{4}{15}\right)(2, -2, 3) \\ &= \left(\frac{22}{15}, \frac{23}{15}, \frac{6}{5}\right) + \lambda\left(\frac{31}{15}, -\frac{16}{15}, \frac{33}{5}\right).\end{aligned}$$

Es resultiert die geometrisch zu erwartende Schnittgerade g_S der beiden Ebenen im dreidimensionalen Raum. Man sollte noch daran denken, eine umparametrisierung vorzunehmen mit dem einfacheren Richtungsvektor $(31, -16, 99)$, der ohne Brüche auskommt. Fassen wir das ganze Verfahren zusammen:

- (1) Parameter trennen
- (2) Aufstellen eines linearen Gleichungssystems nach Gleichsetzen der Terme (d.h. Rechenausdrücke) der Parameterdarstellungen
- (3) Lösen des Gleichungssystems
- (4) Einsetzen der Lösungsmenge in der Parameterfolge nur *einer* der Parameterdarstellungen in die zugehörige Parameterdarstellung ergibt eine Parameterdarstellung des Schnittgebildes. (Auf Endform achten.)

2.8.2. *Beide (linearen) Gebilde in Gleichungsform.* Betrachten wir dieselbe Aufgabe, zwei Ebenen zu schneiden, die jedoch in Gleichungsform gegeben sind:

$$\begin{array}{rclcl} E: & 2x & -y & +3z & = 1 \\ F: & x & +3y & -4z & = 1 \end{array}.$$

Viel günstiger, nur zwei Gleichungen mit drei Parametern. Zur Logik: Sei $\vec{x}_S = (x_S, y_S, z_S)$ der Ortsvektor eines beliebigen Schnittpunktes. Dann bedeutet $S \in E$ gerade $2x_S - y_S + 3z_S = 1$, analog bedeutet $S \in F$ die Erfüllung der zweiten Gleichung, $S \in E \cap F$ bedeutet also, dass *die Koordinaten des Schnittpunktes das Gleichungssystem erfüllen*. Somit ist die Lösungsmenge des Gleichungssystems *unmittelbar* die Schnittmenge! Das Gleichungssystem wird nun nach demselben Muster wie oben gelöst: Hinauswerfen von x bringt

$$7y - 11z = 1.$$

Wir erhalten den erwarteten einen freien Parameter, wählen dazu z . Das ergibt $y = 1/7 + 11/7z$, $x = 1 + 4z - 3y = 1 + 4z - 3/7 - 33/7z = 4/7 - 5/7z$ (aus der zweiten Gleichung). Die Lösungsmenge in ordentlicher Parameterform:

$$\vec{x}_{g_S}(z) = \frac{1}{7}(4, 1, 0) + z\left(-\frac{5}{7}, \frac{11}{7}, 1\right), \quad z \in \mathbb{R}.$$

Damit ist die Schnittgerade ausgerechnet. Später werden wir lernen, Parameterdarstellungen von Ebenen systematisch in Gleichungsform zu bringen, im Zusammenhang mit Skalarprodukt und Vektorprodukt.

Das Verfahren zusammengefasst:

- (1) Die Gleichungen (oder auch schon Systeme) beider Gebilde zu einem Gleichungssystem zusammenfassen. (Die Koordinaten-Unbestimmten müssen überall *gleich* benannt sein.)
- (2) Lösen des entstehenden linearen Gleichungssystems (Unbekannte sind die Koordinaten!).
- (3) Die Lösungsmenge in parametrisierter Form ist unmittelbar eine Parameterdarstellung des Schnittgebildes.

2.8.3. *Mischform: Eine Parameterdarstellung und eine Gleichung, linearer Fall.* Wir schneiden g mit E , gegeben durch $\vec{x}_g(\lambda) = (1, 2, 2) + \lambda(2, -1, 3)$, $\lambda \in \mathbb{R}$, E mit der Gleichung $2x - y - z = 0$. Logik: Jeder Punkt $S \in g \cap E$ muss auf der Geraden liegen, also mit einem Wert λ_S muss $\vec{x}_S = \vec{x}_g(\lambda_S)$ gelten. Dessen Koordinaten müssen nun auch die Ebenengleichung erfüllen. Dies ergibt folgende Gleichung für λ_S :

$$2 + 4\lambda - 2 + \lambda - 2 - 3\lambda = 0, \text{ also } \lambda = 1.$$

Dies ist nun wieder in die Parameterdarstellung einzusetzen. Es gibt also genau einen Schnittpunkt S , und es ist $\vec{x}_S = (3, 1, 5)$. Man sollte sich nun vorstellen können, dass der Schnitt zweier Ebenen bei Mischform darauf hinausläuft, eine lineare Parametergleichung mit zwei Parametern als Unbestimmten zu lösen. Es bleibt ein freier Parameter, Einsetzen funktioniert dann wie im ersten Beispiel (zwei Ebenen, beide in Parameterform). Das Verfahren ist also:

- (1) Einsetzen der Koordinaten des Terms der Parameterdarstellung des einen Gebildes in die Gleichung(en) des andern Gebildes.
- (2) Lösen des entstehenden Gleichungssystems. (Die Unbekannten sind die freien Parameter.)
- (3) Einsetzen der Lösungsmenge in den Term der Parameterdarstellung ergibt eine Parameterdarstellung des Schnittgebildes.

2.9. Zwei nichtlineare Beispiele. Wir schneiden die Parabel, die durch $\vec{x}(t) = t(1, 1, 2) + t^2(1, 2, 1)$, $t \in \mathbb{R}$, gegeben ist, mit der Ebene, deren Gleichung $2x - y + z = 1$ lautet. Die Logik funktioniert offenbar genau wie im vorigen Beispiel: Die Koordinaten von $\vec{x}(t)$ sind in die Ebenengleichung einzusetzen, das ergibt die quadratische Gleichung $2t + 2t^2 - t - 2t^2 + 2t + t^2 = 1$, und man erhält die Lösungen $t_{1,2} = -\frac{3}{2} \pm \frac{1}{2}\sqrt{13}$, die wiederum in $\vec{x}(t)$ einzusetzen sind. Entsprechend dem Lösungsverhalten quadratischer Gleichungen verhalten sich die geometrischen Möglichkeiten beim Schnitt einer Parabel mit einer Ebene: Kein Schnittpunkt, ein Schnittpunkt (Berührungspunkt) oder zwei. Dazu kommt noch der Fall, dass die Parabel ganz in der Ebene liegt, was sich beim beschriebenen Rechenvorgang darin äußern würde, dass die Gleichung $0 = 0$ für t entstünde. Dann ist *jede* Zahl $t \in \mathbb{R}$ Lösung, somit der Schnitt die gesamte Parabel.

Auch in den anderen Fällen bleibt die Logik dieselbe, das Unangenehme sind nur die entstehenden nichtlinearen Gleichungen. Aber in einfachen Fällen geht das. Schneiden wir etwa das Paraboloid $z = x^2 + y^2$ mit der Ebene $x + y - z = 1$, so erhalten wir für die Koordinaten der Schnittpunkte die Gleichung $x^2 + y^2 - x - y + 1 = 0$. Fassen wir y als freien Parameter auf, so haben wir eine quadratische Gleichung mit den Lösungen

$$x_{1,2}(y) = \frac{1}{2} \pm \sqrt{y - y^2 - \frac{3}{4}},$$

wenn denn der Term unter der Wurzel positiv ist. Man sieht leicht, dass er stets < 0 ist, also gibt es keine Lösung, die Schnittmenge ist leer, Ebene und Paraboloid laufen knapp aneinander vorbei. An diesem Beispiel sollte man auch erkennen, dass Unlösbarkeit einer Gleichung durchaus einen Sinn haben kann. Die Verfahren wären wie oben in den linearen Fällen zu nennen, allerdings mit dem gewaltigen Unterschied, dass die entstehenden Gleichungen bzw. Gleichungssysteme nichtlinear sind, was das *exakte* Lösen im Normalfall sogar *prinzipiell* unmöglich macht, nicht jedoch das Finden von Näherungslösungen.

3. Vektorräume über $\mathbb{R}$ mit Skalarprodukt; Längen und Winkel

Vorbemerkung: Wenn hier Zahlenpaare bzw. Zahlentripel geometrisch gedeutet werden, so wird dabei stets ein kartesisches Koordinatensystem vorausgesetzt. (Sonst werden die Deutungen dieses Abschnitts falsch, im Gegensatz zum Bisherigen.)

3.1. Der Betrag (oder die Länge) eines Vektors im $\mathbb{R}^{2,3}$ oder auch allgemein $\mathbb{R}^n$. Wir besprechen zunächst die konkrete Fassung für die Dimensionen 2 und 3, dann die abstrakte.

Aus einer einfachen Anwendung des Satzes von Pythagoras erkennt man, dass die Längen des Vektors $\begin{pmatrix} a \\ b \end{pmatrix}$ (bei Deutung in kartesischem System) sich berechnet zu

$$(1) \quad \left| \begin{pmatrix} a \\ b \end{pmatrix} \right| = \sqrt{a^2 + b^2}.$$

(Wir bezeichnen den Betrag bzw. die Länge von $\begin{pmatrix} a \\ b \end{pmatrix}$ mit $\left| \begin{pmatrix} a \\ b \end{pmatrix} \right|$. Das folgende Bild verdeutlicht diesen Sachverhalt:

Wendet man den Pythagoras erneut an, so bekommt man auch: Die Länge des Vektors $\begin{pmatrix} a \\ b \\ c \end{pmatrix}$ ist

$$\left| \begin{pmatrix} a \\ b \\ c \end{pmatrix} \right| = \sqrt{a^2 + b^2 + c^2}.$$

Begründung: Nach Pythagoras hat man

$$\left| \begin{pmatrix} a \\ b \\ c \end{pmatrix} \right|^2 = c^2 + \left| \begin{pmatrix} a \\ b \end{pmatrix} \right|^2 = c^2 + a^2 + b^2, \text{ also (2).}$$

Es liegt nunmehr nahe, diese Bildung auf $\mathbb{R}^n$ zu verallgemeinern mit

$$(3) \quad \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{pmatrix} : = \sqrt{a_1^2 + a_2^2 + \dots + a_n^2} \text{ (abstrakte Definition) der Länge im } \mathbb{R}^n,$$

$$(4) \quad d(\vec{x}, \vec{y}) : = |\vec{x} - \vec{y}| \quad (\vec{x}, \vec{y} \in \mathbb{R}^n) \text{ (zugehöriger induzierter Abstandsbegriff).}$$

Allerdings wird man sich zunächst einmal fragen, wozu das gut sein soll. Das kann man aber leicht beantworten: Es ist doch offensichtlich sogar sinnvoll, danach zu fragen, wie weit eine Näherung einer Funktion von der Funktion entfernt ist, die dadurch genähert werden soll, oder danach zu fragen, wie stark die Umlaufbahn eines Mondes um einen Planeten variiert. Dann hat man es mit noch viel Schlimmerem zu tun, einem Abstandsbegriff in einem unendlichdimensionalen Vektorraum. Dagegen ist die Bildung für $\mathbb{R}^n$ harmlos. Aber diese ist immerhin schon anwendbar auf den Abstand von zwei Funktionen, welche nur durch ihre Werte an endlich vielen Stellen (auf einem hinreichend feinen Gitter in einem endlichen Intervall). Hier ist noch eine konkrete Anwendung 'aus dem Leben', von der sehr häufig in verschiedenen Wissenschaften Gebrauch gemacht wird: Man stelle sich vor, dass ein Profil (eines Menschen, eines Staates, eines Wirtschaftsunternehmens) durch eine Folge von Messwerten gegeben wird. Dann kann man mit dem Betrag im $\mathbb{R}^n$ den Abstand zwischen zwei solchen Profilen beschreiben - ein sehr natürliches Anliegen. Aber für den Sinn einer solchen Verallgemeinerung sind nicht nur entsprechende Bedürfnisse notwendig, sondern es muss auch rein mathematisch geklärt werden, dass die Absicht tatsächlich durchführbar ist. Wir fragen also: Ist mit (3) tatsächlich so etwas wie ein Längenbegriff und entsprechend mit (4) ein vernünftiger Abstandsbegriff definiert? Dazu macht man sich zunächst einmal klar, was die wesentlichen Eigenschaften eines Längenbegriffes sind (dies genau ist das Vorgehen der axiomatischen Mathematik), und dann erkennt man folgende Eigenschaften, durch welche der Begriff eines Betrages oder einer Länge auf einem Vektorraum abstrakt definiert wird:

DEFINITION 11 (abstrakte Definition des Begriffs eines Betrages auf einem Vektorraum). *Sei V ein Vektorraum über $\mathbb{R}$ und*

$$\| : \quad \begin{array}{ccc} V & \rightarrow & \mathbb{R}_{\geq 0} \\ \vec{x} & \mapsto & |\vec{x}| \end{array}.$$

Dann ist $\|$ genau dann ein Betrag auf V (man sagt auch: eine Norm auf V), wenn folgende Eigenschaften erfüllt sind, für alle $\lambda \in \mathbb{R}$, $\vec{x}, \vec{y} \in V$:

$$\begin{aligned} (1) \quad |\lambda \vec{x}| &= |\lambda| |\vec{x}|, \\ (2) \quad |\vec{x} + \vec{y}| &\leq |\vec{x}| + |\vec{y}| \\ (3) \quad |\vec{x}| &= 0 \iff \vec{x} = \vec{0}. \end{aligned}$$

Zusammen mit einer solchen Norm heißt V dann ein normierter Vektorraum.

Bemerkungen: Beträge sind also stets Zahlen ≥ 0 , und nur der Nullvektor hat den Betrag Null. (2) ist die berühmte Dreiecksungleichung: läuft man über zwei Vektoren, dann ist das mindestens so lang wie über die Summe zu laufen. (Nicht etwa kommt Gleichheit, man kann zeigen, dass die nur bei $\vec{x} = \lambda \vec{y}$ mit $\lambda \geq 0$ auftritt). (1) ist dann noch eine gewisse Form der Verträglichkeit mit der Multiplikation eines Vektors mit einem Skalar.

Wir bemerken, dass (2) durchaus nicht leicht für unsere Beispiele oben zu zeigen ist. Das gelingt erst über senkrechte Projektion (s.u.).

3.2. Das Skalarprodukt auf dem $\mathbb{R}^n$. Vorbemerkung: Uns interessieren vor allem die Fälle $n = 2, 3$. Motivieren kann man das Skalarprodukt so: Wir wissen, dass der Winkel zwischen $\vec{x}$ und $\vec{y}$ genau

dann ein rechter ist, wenn $|\vec{x} - \vec{y}|^2 = |\vec{x}|^2 + |\vec{y}|^2$. Mit der analytischen Beschreibung des Betrages heißt das:

$$\begin{aligned} \vec{x} \perp \vec{y} &\iff \sum_{i=1}^n (x_i - y_i)^2 = \sum_{i=1}^n x_i^2 + \sum_{i=1}^n y_i^2 - 2 \sum_{i=1}^n x_i y_i = \sum_{i=1}^n x_i^2 + \sum_{i=1}^n y_i^2 \\ &\iff \sum_{i=1}^n x_i y_i = 0. \end{aligned}$$

Damit erweist sich die folgende Größe als wichtig, das Skalarprodukt:

DEFINITION 12 (Skalarprodukt zweier Vektoren im $\mathbb{R}^n$). *Das Skalarprodukt auf $\mathbb{R}^n$ ist definiert durch*

$$\cdot : \left(\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \right) \mapsto \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \cdot \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} := \sum_{i=1}^n x_i y_i$$

Das Skalarprodukt heißt so, weil Skalare (d.h. Zahlen) herauskommen. Wir stellen zunächst die selbstverständlich zu folgernden algebraischen Eigenschaften zusammen:

SATZ 17. *Das Skalarprodukt hat folgende algebraische Eigenschaften:*

- (i) $(\lambda \vec{x}) \cdot \vec{y} = \lambda (\vec{x} \cdot \vec{y})$,
- $(\vec{x} + \vec{y}) \cdot \vec{u} = \vec{x} \cdot \vec{u} + \vec{y} \cdot \vec{u}$ (Bilinearität zusammen mit der Symmetrie (ii))
- (ii) $\vec{x} \cdot \vec{y} = \vec{y} \cdot \vec{x}$ (Symmetrie)
- (iii) $\vec{x} \cdot \vec{x} \geq 0$, und $\vec{x} \cdot \vec{x} = 0 \iff \vec{x} = \vec{0}$ (positive Definitheit)

Diese Eigenschaften rechnet man sofort nach mit der Definition oben. Man beachte noch, dass $\vec{x} \cdot \vec{x} = |\vec{x}|^2$

Wie oben bereits festgestellt, hat es mit Winkeln zu tun, und man kann damit nicht nur rechte Winkel, sondern alle quantifizieren:

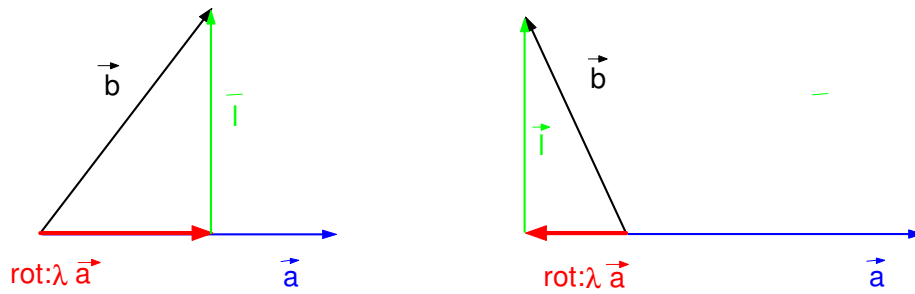
SATZ 18. *Die senkrechte Projektion eines Vektors $\vec{b}$ auf den Vektor $\vec{a} \neq \vec{0}$ ist der Vektor*

$$\frac{\vec{b} \cdot \vec{a}}{|\vec{a}|^2} \vec{a}.$$

Der Kosinus des Winkels zwischen den Vektoren $\vec{a}$ und $\vec{b}$ (beide $\neq \vec{0}$), kurz $\cos(\angle(\vec{a}, \vec{b}))$, ist:

$$\cos(\angle(\vec{a}, \vec{b})) = \frac{\vec{a} \cdot \vec{b}}{|\vec{a}| |\vec{b}|}.$$

Beweis: Zunächst wollen wir den Begriff der senkrechten Projektion erklären und illustrieren mit folgenden Bildern:



Man hat: Die senkrechte Projektion von $\vec{b}$ auf $\vec{a}$ ist einerseits von der Form $\lambda \vec{a}$, andererseits gilt die Gleichung:

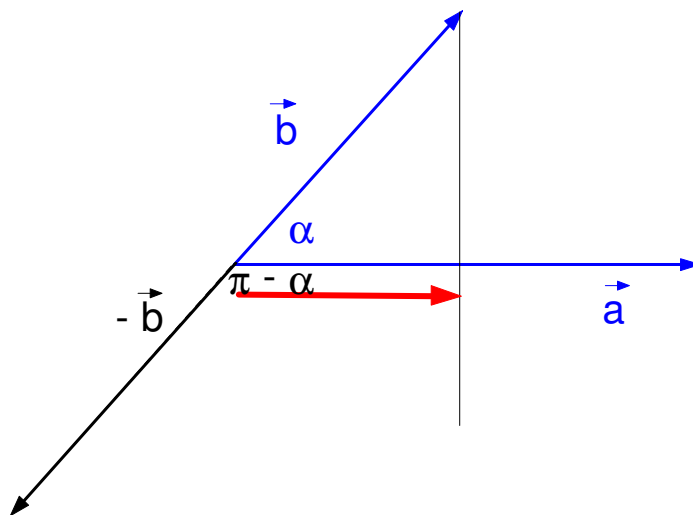
$$\lambda \vec{a} + \vec{I} = \vec{b}.$$

Wenn man diese Gleichung skalar mit $\vec{a}$ multipliziert, so kommt, da $\vec{I} \cdot \vec{a} = 0$, heraus:

$$\lambda = \frac{\vec{b} \cdot \vec{a}}{\vec{a} \cdot \vec{a}} = \frac{\vec{b} \cdot \vec{a}}{|\vec{a}|^2}.$$

Also ist die senkrechte Projektion von $\vec{b}$ auf $\vec{a}$ gleich dem Vektor $\frac{\vec{b} \cdot \vec{a}}{|\vec{a}|^2} \vec{a}$.

Für die Kosinusdarstellung betrachten wir folgendes Bild, zunächst setzen wir den Winkel zwischen den Vektoren als spitz voraus - rot gezeichnet ist der Vektor: Senkrechte Projektion von $\vec{b}$ auf $\vec{a}$:



Nach der elementaren Definition von $\cos$ hat man:

$$\begin{aligned}\cos\left(\angle\left(\vec{a}, \vec{b}\right)\right) &= \cos(\alpha) = \frac{\text{Länge der Ankathete}}{\text{Länge der Hypotenuse}} \\ &= \frac{\left|\frac{\vec{b} \cdot \vec{a}}{|\vec{a}|^2} \vec{a}\right|}{|\vec{b}|} = \frac{1}{|\vec{b}|} \cdot \left|\frac{\vec{b} \cdot \vec{a}}{|\vec{a}|^2}\right| |\vec{a}| \\ &= \frac{1}{|\vec{b}|} \frac{|\vec{b} \cdot \vec{a}|}{|\vec{a}|^2} |\vec{a}| = \frac{|\vec{b} \cdot \vec{a}|}{|\vec{b}| |\vec{a}|} =^* \frac{\vec{a} \cdot \vec{b}}{|\vec{a}| |\vec{b}|}\end{aligned}$$

*) Das letzte Gleichheitszeichen ergibt sich mit dem spitzen Winkel, woraus folgt, dass das Skalarprodukt positiv ist. Das wissen wir daher, dass **im Falle eines spitzen Winkels** gilt

$$|\vec{a} - \vec{b}|^2 = |\vec{a}|^2 + |\vec{b}|^2 - 2\vec{a} \cdot \vec{b} < |\vec{a}|^2 + |\vec{b}|^2, \text{ also } \vec{a} \cdot \vec{b} > 0.$$

Nunmehr sehen wir am Winkel $\pi - \alpha$ zwischen $-\vec{b}$ und $\vec{a}$ sofort ein:

$$\begin{aligned}\cos\left(\angle\left(\vec{a}, -\vec{b}\right)\right) &= \cos(\pi - \alpha) = -\cos(\alpha) \\ &= -\frac{\vec{a} \cdot \vec{b}}{|\vec{a}| |\vec{b}|} = \frac{\vec{a} \cdot (-\vec{b})}{|\vec{a}| |\vec{b}|}.\end{aligned}$$

■

Anwendungsbeispiele:

1) Welchen Winkel α bilden die Vektoren $\begin{pmatrix} 2 \\ 1 \end{pmatrix}$ und $\begin{pmatrix} -3 \\ 2 \end{pmatrix}$ miteinander? Antwort:

$$\begin{aligned}\cos(\alpha) &= \frac{\begin{pmatrix} 2 \\ 1 \end{pmatrix} \cdot \begin{pmatrix} -3 \\ 2 \end{pmatrix}}{\left|\begin{pmatrix} 2 \\ 1 \end{pmatrix}\right| \cdot \left|\begin{pmatrix} -3 \\ 2 \end{pmatrix}\right|} = \frac{-4}{\sqrt{5}\sqrt{13}}, \text{ also} \\ \alpha &= \arccos\left(\frac{-4}{\sqrt{5}\sqrt{13}}\right) \approx 2.09, \text{ das sind etwa 120 Grad.}\end{aligned}$$

Am Vorzeichen sah man schon, dass es sich um einen stumpfen Winkel handeln musste.

2) Welche geometrische Bedeutung hat unmittelbar die Bestimmungsgleichung $\vec{a} \cdot \vec{x} = d$ mit festem Vektor $\vec{a} \neq \vec{0}$, $\vec{a} \in \mathbb{R}^3$, und fester Zahl $d \in \mathbb{R}$? Wir schreiben die Gleichung auch noch einmal in Koordinaten auf:

$$a_1x + a_2y + a_3z = d.$$

Wir wissen schon, dass dies eine Ebene E ergibt. Aber nun können wir mehr sagen: Nehmen wir zwei Punkte P, Q der Ebene, dann haben wir

$$\vec{a} \cdot \vec{x}_P = \vec{a} \cdot \vec{x}_Q = d,$$

also

$$\vec{a} \cdot (\vec{x}_P - \vec{x}_Q) = 0.$$

Nun ist $\vec{x}_P - \vec{x}_Q$ ein beliebiger Vektor parallel zur beschriebenen Ebene E . Also steht $\vec{a}$ senkrecht zu E . Wir sagen auch: $\vec{a}$ ist ein Normalenvektor zu E . Außerdem können wir auch die Zahl d genau deuten: Es gibt genau eine Zahl λ , so dass $\lambda \vec{a}$ Ortsvektor eines Punktes auf E ist, also

$$\vec{a} \cdot (\lambda \vec{a}) = d \text{ gilt,}$$

und damit hat man

$$\lambda = \frac{d}{|\vec{a}|^2}.$$

Daher ist der Abstand zwischen dem Ursprung O und E :

$$d(O, E) = |\lambda \vec{a}| = \frac{|d|}{|\vec{a}|}.$$

3) Welchen (kleineren) Winkel α bilden die Ebenen E und F miteinander, welche beschrieben sind durch die Koordinatengleichungen

$$E : 2x - 3y + 4z = 0$$

$$F : -x + 2y + z = 1?$$

Offenbar bilden die Ebenen miteinander dieselben Winkel wie ihre Normalenvektoren, also (für den kleineren Winkel haben wir das Skalarprodukt positiv zu nehmen (!):

$$\cos(\alpha) = -\frac{\begin{pmatrix} 2 \\ -3 \\ 4 \end{pmatrix} \cdot \begin{pmatrix} -1 \\ 2 \\ 1 \end{pmatrix}}{\left| \begin{pmatrix} 2 \\ -3 \\ 4 \end{pmatrix} \right| \cdot \left| \begin{pmatrix} -1 \\ 2 \\ 1 \end{pmatrix} \right|} = \frac{4}{\sqrt{29}\sqrt{6}}, \text{ also}$$

$$\alpha \approx 1.26 \text{ (im Bogenmaß)}, \text{ etwas mehr als } 72 \text{ Grad.}$$

4) Wie kann man den (kleineren) Winkel α zwischen einer Geraden mit Richtungsvektor $\vec{a}$ und einer Ebene mit Normalenvektor $\vec{b}$ ausrechnen? Mit einer Skizze macht man sich leicht klar, dass

$$\sin(\alpha) = \frac{|\vec{a} \cdot \vec{b}|}{|\vec{a}| |\vec{b}|}.$$

5) Wir zeigen mit dem Skalarprodukt und der senkrechten Projektion, dass für den oben definierten Betrag tatsächlich die Dreiecksungleichung gilt: Wenn $\vec{a} \neq \vec{0}$, dann hat man für einen beliebigen Vektor $\vec{b}$:

$$\vec{b} = \lambda \vec{a} + \vec{l}, \text{ mit } \vec{l} \perp \vec{a},$$

dabei ist $\lambda \vec{a}$ gerade die oben auch ausgerechnete senkrechte Projektion von $\vec{b}$ auf $\vec{a}$. Es folgt:

$$\vec{a} \cdot \vec{b} = \vec{a} \cdot (\lambda \vec{a} + \vec{l}) = \lambda \vec{a} \cdot \vec{a} + \vec{a} \cdot \vec{l} = \lambda \vec{a} \cdot \vec{a}.$$

Ferner

$$|\vec{b}|^2 = |\lambda \vec{a} + \vec{l}|^2 = (\lambda \vec{a} + \vec{l}) \cdot (\lambda \vec{a} + \vec{l}) = \lambda^2 \vec{a} \cdot \vec{a} + 2\lambda \vec{a} \cdot \vec{l} + \vec{l} \cdot \vec{l} = \lambda^2 \vec{a} \cdot \vec{a} + \vec{l} \cdot \vec{l},$$

also

$$|\vec{a}|^2 |\vec{b}|^2 \geq \lambda^2 |\vec{a}|^4 \geq |\vec{a} \cdot \vec{b}|^2.$$

Das ergibt die berühmte

SATZ 19 (Schwarzsche Ungleichung).

$$|\vec{a} \cdot \vec{b}| \leq |\vec{a}| \cdot |\vec{b}|,$$

und man kann oben noch entnehmen, dass das Gleichheitszeichen genau dann eintritt, wenn $\vec{a}$ parallel zu $\vec{b}$ ist (gleichwertig $\vec{l} = \vec{0}$).

FOLGERUNG 1 (Dreiecksungleichung für den Betrag).

$$|\vec{x} + \vec{y}| \leq |\vec{x}| + |\vec{y}|.$$

Beweis:

$$\begin{aligned} |\vec{x} + \vec{y}|^2 &= (\vec{x} + \vec{y}) \cdot (\vec{x} + \vec{y}) = \vec{x} \cdot \vec{x} + 2\vec{x} \cdot \vec{y} + \vec{y} \cdot \vec{y} \\ &\leq |\vec{x}|^2 + |\vec{y}|^2 + 2|\vec{x}| |\vec{y}| = (|\vec{x}| + |\vec{y}|)^2. \end{aligned}$$

(Für die Ungleichung wurde die Schwarzsche Ungleichung verwandt.) Also

$$\begin{aligned} |\vec{x} + \vec{y}|^2 &\leq (|\vec{x}| + |\vec{y}|)^2, \text{ und nach Wurzelziehen} \\ |\vec{x} + \vec{y}| &\leq |\vec{x}| + |\vec{y}|. \blacksquare \end{aligned}$$

3.3. Anwendung des Skalarproduktes auf kürzeste Wege zwischen zwei Punkten auf einer Kugel. Wir stellen uns die folgende Aufgabe: Wie weit ist es von einem Punkt einer Kugel zu einem anderen Punkt der Kugel - gemeint ist nicht etwa die 'Luftlinie', die hier besser 'Linie durch die Erde grabend - oder unter dem Wasser tauchend', also die geradlinige Verbindung, sondern die kürzeste Verbindung auf der Kugeloberfläche fahrend. Wir setzen voraus, dass die so verstandene kürzeste Entfernung immer die Länge der Großkreisverbindung ist. Wir denken uns den Ursprung im Kugelmittelpunkt und die Punkte A und B auf der Kugeloberfläche durch ihre Ortsvektoren $\vec{x}_A$ und $\vec{x}_B$ gegeben. Dann ist klar: Wenn α der Winkel (im Bogenmaß) zwischen $\vec{x}_A$ und $\vec{x}_B$ ist, so ist die gesuchte kürzeste Entfernung:

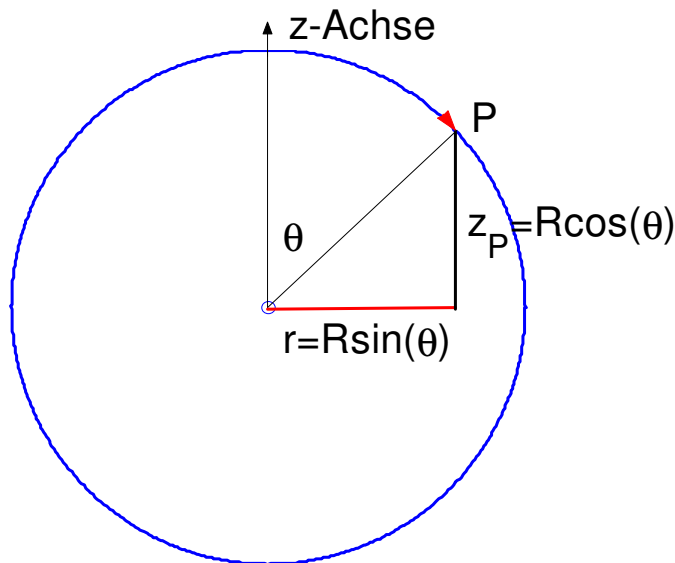
$$d_K(A, B) = \frac{\alpha}{2\pi} \cdot 2\pi R = \alpha R.$$

Denn man hat die Länge des Anteils $\frac{\alpha}{2\pi}$ an der vollen Großkreislänge $2\pi R$ zu nehmen. Wir benötigen also nur den Winkel (im Bogenmaß!) und den Kugelradius R .

Für die Berechnung von α haben wir ein wenig zu arbeiten, aber das geht ganz gut mit der Skalarproduktmethode, wenn wir einmal verstanden haben, die kartesische Koordinatendarstellung eines Punktes auf der Kugeloberfläche aus den üblichen Längen- und Breitengrad- Angaben zu produzieren. Zunächst einmal nehmen wir diesen Winkel im Bogenmaß, dabei gehen die Längengrade φ also von 0 bis unter 2π . Die Breitengrade ϑ nimmt man gern von 0 bis π (einschließlich, wobei 0 zum Nordpol und also π zum Südpol gehört). Damit hat man für den Punkt mit den Winkeln φ, ϑ folgende kartesische Koordinatendarstellung:

$$\vec{x}(\vartheta, \varphi) = \begin{pmatrix} R \cos(\varphi) \sin(\vartheta) \\ R \sin(\varphi) \sin(\vartheta) \\ R \cos(\vartheta) \end{pmatrix}.$$

Das sieht man mit diesem Bild ein:



Blau eingezeichnet ist der Großkreis auf einer Kugelschale vom Radius R , auf welcher der Punkt P liegt. Der Winkel ϑ (typographisch θ im Bild) wird vom Nordpol nach unten gemessen, im Beispiel nahe $\pi/4$. Dann ist die Entfernung des Punktes P von der z -Achse gleich $r = R \sin(\vartheta)$. Also hat die Projektion von P auf die xy -Ebene die Koordinatendarstellung $(r \cos(\varphi), r \sin(\varphi)) = (R \cos(\varphi) \sin(\vartheta), R \sin(\varphi) \sin(\vartheta))$. Das sind also die (x, y) -Koordinaten von P . Ferner hat P die z -Koordinate $z_P = R \cos(\vartheta)$, wie das

Bild ebenfalls sofort zeigt. Damit ist die Koordinatendarstellung von $\vec{x}(\vartheta, \varphi)$ gezeigt. Wir rechnen nunmehr den Winkel α aus zwischen $\vec{x}(\vartheta_1, \varphi_1)$ und $\vec{x}(\vartheta_2, \varphi_2)$. Man beachte, dass man dafür den Faktor R weglassen kann (warum?) und dass mit $R = 1$ gilt: $|\vec{x}(\vartheta, \varphi)| = 1$. Also hat man

$$\begin{aligned}\cos(\alpha) &= \begin{pmatrix} \cos(\varphi_1) \sin(\vartheta_1) \\ \sin(\varphi_1) \sin(\vartheta_1) \\ \cos(\vartheta_1) \end{pmatrix} \cdot \begin{pmatrix} \cos(\varphi_2) \sin(\vartheta_2) \\ \sin(\varphi_2) \sin(\vartheta_2) \\ \cos(\vartheta_2) \end{pmatrix} \\ &= \cos(\varphi_1) \cos(\varphi_2) \sin(\vartheta_1) \sin(\vartheta_2) \\ &\quad + \sin(\varphi_1) \sin(\varphi_2) \sin(\vartheta_1) \sin(\vartheta_2) + \cos(\vartheta_1) \cos(\vartheta_2).\end{aligned}$$

Dies können wir mit Anwendung des Additionstheorems für $\cos$ noch so vereinfachen:

$$\cos(\alpha) = \cos(\varphi_1 - \varphi_2) \sin(\vartheta_1) \sin(\vartheta_2) + \cos(\vartheta_1) \cos(\vartheta_2).$$

Damit ist die Aufgabe gelöst, die Entfernung zwischen den beiden Punkten auf der Kugeloberfläche mit den bezeichneten Winkeln ist dann:

$$d_K(A, B) = R \cdot \arccos[\cos(\varphi_1 - \varphi_2) \sin(\vartheta_1) \sin(\vartheta_2) + \cos(\vartheta_1) \cos(\vartheta_2)].$$

Ein Anwendungsbeispiel (dabei geben wir die Punkte mit den üblichen geographischen Angaben auf der Erdkugel mit Radius etwa $\frac{40000}{2\pi}$ [km]): A habe Längengrad 50° westlicher Länge und Breitengrad 30° nördlicher Breite, B habe Längengrad 50° östlicher Länge und Breitengrad 45° südlicher Breite. Damit haben wir

$$\begin{aligned}A &: \varphi_1 = \frac{-50}{180} \cdot \pi, \vartheta_1 = \frac{60}{180} \cdot \pi = \frac{\pi}{3}, \\ B &: \varphi_2 = \frac{50}{180} \cdot \pi, \vartheta_2 = \frac{135}{180} \cdot \pi = \frac{3\pi}{4}.\end{aligned}$$

Damit ist $\varphi_2 - \varphi_1 = \frac{5}{9}\pi$ und

$$\begin{aligned}d_K(A, B) &= \frac{40000}{2\pi} \cdot \arccos\left[\cos\left(\frac{5}{9}\pi\right) \sin\left(\frac{\pi}{3}\right) \sin\left(\frac{3\pi}{4}\right) + \cos\left(\frac{\pi}{3}\right) \cos\left(\frac{3\pi}{4}\right)\right] \\ &= \frac{40000}{2\pi} \cdot \arccos\left[\frac{1}{2}\sqrt{3}\frac{1}{2}\sqrt{2}\cos\left(\frac{5}{9}\pi\right) + \frac{1}{2} \cdot \left(-\frac{1}{2}\sqrt{2}\right)\right] \\ &\approx 13042[\text{km}].\end{aligned}$$

4. Lineare Abbildungen und ihre Matrixdarstellungen

Eine Grundlage der Computergrafik besteht in der Berechnung der zweidimensionalen Bildkoordinaten aus den Koordinaten von Punkten in einem dreidimensionalen Koordinatensystem für den Anschauungsraum. Wir besprachen bereits, wie das bei Zentralprojektion oder auch Parallelprojektion funktioniert. Die zweite Grundlage besteht darin, dass man bestimmen kann, in welcher Weise ein Körper für eine günstige Darstellung gedreht werden soll. Wie man analytisch Drehungen durch Drehmatrizen beschreiben kann, werden wir in diesem Abschnitt sehen.

Noch ein anderer Gesichtspunkt ist für das nächste Kapitel 'Symmetrien' wichtig: Wir wollen auch Spiegelungen durch Matrizen darstellen und anschaulich sowie analytisch sehen, was sich bei Hintereinanderschaltung von Spiegelungen und Drehungen ergibt. Das Folgende beschreiben wir allgemein, aber wir interessieren uns vor allem für $n = 2, 3$.

4.1. Lineare und orthogonale Abbildungen.

DEFINITION 13. Eine Abbildung $\vec{f}: \mathbb{R}^n \rightarrow \mathbb{R}^n$ heißt linear, wenn folgende Formeln gelten, allgemein für $\vec{x}, \vec{y} \in \mathbb{R}^n, \lambda \in \mathbb{R}$:

$$\begin{aligned}(i) \quad \vec{f}(\vec{x} + \vec{y}) &= \vec{f}(\vec{x}) + \vec{f}(\vec{y}) \\ (ii) \quad \vec{f}(\lambda \vec{x}) &= \lambda \vec{f}(\vec{x}).\end{aligned}$$

Bemerkung: Für $n = 1$ hat man nur die linearen Abbildungen $f(x) = ax$, mit fester Zahl $a \in \mathbb{R}$. (Additive Konstanten sind bei diesem engeren Begriff von 'linear' nicht zulässig, man nennt solche wie $g(x) = ax + b$ zur Unterscheidung 'affine Abbildungen'. Aber für $n \geq 2$ hat man recht wichtige interessante Abbildungen, die linear sind: Drehungen um den Ursprung und Spiegelungen an einer Geraden oder einer Ebene, welche durch den Ursprung gehen. Wir werden sehen, dass man Spiegelungen an Geraden oder Ebenen, welche nicht durch den Ursprung gehen, oder Drehungen um andere Punkte als den Ursprung dann auch durch Hinzufügen additiver Konstanten in den Griff bekommt. Das sind dann also affine Abbildungen.

Aus der Definition folgt: Wenn man z.B. im $\mathbb{R}^2$ ein Zweibein $\vec{a}, \vec{b}$ hat, also zwei nicht parallele Vektoren $\neq \vec{0}$, welche eine Basis bilden, dann ist jeder Vektor $\vec{x} \in \mathbb{R}^2$ in *eindeutiger Weise* darstellbar als

$$\begin{aligned}\vec{x} &= \lambda \vec{a} + \mu \vec{b}, \text{ und dann folgt für eine lineare Abbildung } \vec{f}: \\ \vec{f}(\vec{x}) &= \lambda \vec{f}(\vec{a}) + \mu \vec{f}(\vec{b}).\end{aligned}$$

Also braucht man nur $\vec{f}(\vec{a})$ und $\vec{f}(\vec{b})$ zu kennen, um alle Bilder zu kennen. Das werden wir im nächsten Abschnitt für die Matrizendarstellung ausnutzen. Analog ist eine lineare Abbildung $\vec{f}: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ bekannt, wenn sie auf einem Dreibein bekannt ist.

Nun ist aber eine Drehung oder eine Spiegelung nicht nur eine lineare Abbildung, sondern darüber hinaus bleiben unter ihr Längen und Winkel erhalten. Solche Abbildungen heißen dann orthogonale Abbildungen.

DEFINITION 14. Eine Abbildung $\vec{f}: \mathbb{R}^n \rightarrow \mathbb{R}^n$ heißt *orthogonal*, wenn sie linear ist und Längen und Winkel erhält, also zusätzlich gilt:

$$\vec{f}(\vec{x}) \cdot \vec{f}(\vec{y}) = \vec{x} \cdot \vec{y}, \text{ für alle } \vec{x}, \vec{y} \in \mathbb{R}^n.$$

Bemerkungen: Folgern Sie aus der angegebenen Eigenschaft, dass das Skalarprodukt erhalten bleibt, dass Längen und Winkel erhalten bleiben. Tatsächlich kann man aus der Eigenschaft der Erhaltung des Skalarproduktes auch auf die Linearität schließen, so dass diese nicht gesondert gefordert werden muss.

Wichtige Beispiele: Alle Drehungen um den Ursprung und Spiegelungen an Ursprungsgeraden bzw. Ursprungsebenen sind orthogonale Abbildungen.

4.2. Elementares über Hintereinanderschaltungen von Abbildungen, insbesondere von Drehungen, Spiegelungen und Translationen. Wir bezeichnen mit zwei Abbildungen $\vec{f}: \mathbb{R}^n \rightarrow \mathbb{R}^n$ und $\vec{g}: \mathbb{R}^n \rightarrow \mathbb{R}^n$ die Hintereinanderschaltung 'zuerst $\vec{f}$, dann $\vec{g}$ ' mit: $\vec{g} \circ \vec{f}$. Offenbar funktioniert diese stets, da die Bilder von $\vec{f}$ im Definitionsbereich von $\vec{g}$ sind. Formal ist definiert:

$$\vec{g} \circ \vec{f}(\vec{x}) := \vec{g}(\vec{f}(\vec{x})).$$

Damit kann man auch zu bijektiven Abbildungen $\vec{f}$ die Umkehrabbildungen erklären:

$$\begin{aligned}\vec{g} &\text{ heißt Umkehrabbildung zu } \vec{f} \text{ genau dann, wenn} \\ \vec{g} \circ \vec{f}(\vec{x}) &= \vec{x} \text{ also } \vec{g} \circ \vec{f} = id \text{ (Identität). Man schreibt dann} \\ \vec{g} &= \vec{f}^{-1} \text{ (Achtung: Das hat nichts mit Kehrwert zu tun!)}\end{aligned}$$

Man kann leicht verstehen:

$$(\vec{g} \circ \vec{f})^{-1} = \vec{f}^{-1} \circ \vec{g}^{-1}.$$

Begründung: Man führt zuerst $\vec{f}$ aus und dann $\vec{g}$. Um dies rückgängig zu machen, muss man offenbar zuerst $\vec{g}$ rückgängig machen (d.h. $\vec{g}^{-1}$ anwenden) und dann $\vec{f}$ rückgängig machen, (also danach $\vec{f}^{-1}$ anwenden).

In einigen Fällen kann man sich anschaulich gut klar machen, was bei einer Hintereinanderschaltung von geometrischen Abbildungen herauskommt, in komplizierteren Fällen kann man das rechnen, vgl den nächsten Unterabschnitt zu den Matrizendarstellungen.

Beispiele:

1) Dreht man in der Ebene um den Ursprung zuerst mit Winkel α entgegen dem Uhrzeigersinn und dann mit Winkel β entgegen dem Uhrzeigersinn, so hat man insgesamt mit Winkel $\alpha + \beta$ entgegen dem Uhrzeigersinn gedreht. Diese Drehungen bezeichnen wir kurz mit D_α, D_β . Symbolisch hat man dann

$$D_\beta \circ D_\alpha = D_{\alpha+\beta} = D_\alpha \circ D_\beta.$$

(Wir benutzen dies so oft, dass der Vektorpfeil lästig wird.) Man beachte noch *Hier* ist es einmal gleichgültig, in welcher Reihenfolge die Operationen ausgeführt werden, allgemein ist das nicht so!

2) Dreht man zuerst in der Ebene um den Ursprung mit Winkel α entgegen dem Uhrzeigersinn und dann mit Winkel α *im Uhrzeigersinn* (oder gleichwertig mit Winkel $-\alpha$ *entgegen dem Uhrzeigersinn*), so bekommt man die Identität. Umkehrabbildung einer Linksdrehung ist also die entsprechende Rechtsdrehung in der Ebene.

3) Die Spiegelung in der Ebene an der Achse a (sie braucht dazu nicht durch den Ursprung zu gehen) hat sich selbst als Umkehrung! Wir notieren die Spiegelung an der Achse a kurz mit S_a (hier ohne Vektorpfeil, das wird sonst zu unbequem).

4) Die Parallelverschiebung mit dem Vektor $\vec{a}$ nennen wir kurz $T_{\vec{a}}$ ('Translation mit dem Vektor $\vec{a}$ '). Man hat offenbar

$$\begin{aligned} T_{\vec{b}} \circ T_{\vec{a}} &= T_{\vec{b}+\vec{a}} = T_{\vec{a}} \circ T_{\vec{b}}, \\ (T_{\vec{a}})^{-1} &= T_{-\vec{a}}. \end{aligned}$$

Man beachte: $T_{\vec{a}}$ ist für $\vec{a} \neq \vec{0}$ keine lineare Abbildung.

5) $T_{\vec{a}} \circ D_\alpha(\vec{x}) = \vec{a} + D_\alpha(\vec{x})$, aber $D_\alpha \circ T_{\vec{a}}(\vec{x}) = D_\alpha(\vec{a} + \vec{x}) = D_\alpha(\vec{a}) + D_\alpha(\vec{x})$. Daran sieht man, dass es außer in trivialen Fällen wie $\vec{a} = \vec{0}$ oder $\alpha = 0$ hier auf die Reihenfolge ankommt: $T_{\vec{a}} \circ D_\alpha \neq D_\alpha \circ T_{\vec{a}}$.

6) Spiegelt man zuerst in der Ebene an der Ursprungsgeraden a und anschließend an der Ursprungsgeraden b und entsteht b aus a durch Drehung um α *entgegen dem Uhrzeigersinn*, so hat man:

$$S_b \circ S_a = D_{\alpha/2}.$$

Um dies zu zeigen, braucht man nur für das Zweibein von zwei gleich langen Vektoren im rechten Winkel die Abbildungen hintereinander auszuführen!

Es folgt sofort:

$$S_a \circ S_b = S_a^{-1} \circ S_b^{-1} = (S_b \circ S_a)^{-1} = D_{-\alpha/2}.$$

Also hat man auch hier im nichttrivialen Fall, dass es auf die Reihenfolge ankommt.

7) Analog macht man sich klar: Sind a, b parallele Achsen in der Ebene und entsteht b durch Verschieben von a mit dem Vektor $\vec{c}$, dann gilt:

$$\begin{aligned} S_b \circ S_a &= T_{\vec{c}/2}, \\ S_a \circ S_b &= T_{-\vec{c}/2}. \end{aligned}$$

Wieder kommt es auf die Reihenfolge an, sofern nur $a \neq b$.

4.3. Matrixdarstellung linearer Abbildungen und orthogonaler Abbildungen, insbesondere in der Ebene. Wir wissen (s.o.), dass eine lineare Abbildung bekannt ist, wenn die Bilder von einer Basis bekannt sind. Konkret für den $\mathbb{R}^2$ und die Basis $\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ und $\vec{e}_2 = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$ heißt das:

Wenn $\vec{f}$ linear ist und gilt:

$$\vec{f}(\vec{e}_1) = \begin{pmatrix} a \\ b \end{pmatrix} \text{ und } \vec{f}(\vec{e}_2) = \begin{pmatrix} c \\ d \end{pmatrix},$$

dann

$$\begin{aligned}
 \vec{f} \left(\begin{pmatrix} x \\ y \end{pmatrix} \right) &= \vec{f} (x \vec{e}_1 + y \vec{e}_2) = x \vec{f} (\vec{e}_1) + y \vec{f} (\vec{e}_2) \\
 &= x \begin{pmatrix} a \\ b \end{pmatrix} + y \begin{pmatrix} c \\ d \end{pmatrix} \\
 &= \begin{pmatrix} xa + yc \\ xb + yd \end{pmatrix} \\
 &= : \begin{pmatrix} a & c \\ b & d \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix}.
 \end{aligned}$$

Man bekommt also das Bild eines beliebigen Vektors $\begin{pmatrix} x \\ y \end{pmatrix} \in \mathbb{R}^2$, indem man die Matrix

$$A = \begin{pmatrix} a & c \\ b & d \end{pmatrix},$$

welche die Bilder der Einheitsvektoren in den Spalten enthält (Reihenfolge ist zu beachten!), in der hier definierten Weise mit dem Vektor $\begin{pmatrix} x \\ y \end{pmatrix}$ multipliziert.

Nun möchte man zwei Matrixabbildungen hintereinanderschalten: Wie muss dann die Matrix der Hintereinanderschaltung aussehen?

SATZ 20 (Matrizenprodukt). *Die Matrix BA (für die Hintereinanderschaltung der Matrixabbildungen 'zuerst A , dann B ') ist*

$$j. \text{ Spalte von } BA = B(j. \text{ Spalte von } A).$$

Beweis: $(BA) \vec{e}_j = B(A \vec{e}_j)$. ■

Berechnungsbeispiel: $\begin{pmatrix} -1 & 2 \\ 3 & -4 \end{pmatrix} \begin{pmatrix} 2 & 3 \\ 4 & 5 \end{pmatrix} = \begin{pmatrix} 6 & 7 \\ -10 & -11 \end{pmatrix}$

$$\begin{aligned}
 \begin{pmatrix} -1 & 2 \\ 3 & -4 \end{pmatrix} \begin{pmatrix} 2 & 3 \\ 4 & 5 \end{pmatrix} &= \begin{pmatrix} 6 & 7 \\ -10 & -11 \end{pmatrix}; \text{ denn} \\
 \begin{pmatrix} -1 & 2 \\ 3 & -4 \end{pmatrix} \begin{pmatrix} 2 \\ 4 \end{pmatrix} &= \begin{pmatrix} 6 \\ -10 \end{pmatrix}, \\
 \begin{pmatrix} -1 & 2 \\ 3 & -4 \end{pmatrix} \begin{pmatrix} 3 \\ 5 \end{pmatrix} &= \begin{pmatrix} 7 \\ -11 \end{pmatrix}.
 \end{aligned}$$

Beispiele zur Anwendung von Matrizen und Matrixprodukten für die Beschreibung von Drehungen und Spiegelungen:

1) Man dreht in der Ebene um den Ursprung mit dem Winkel α *entgegen* dem Uhrzeigersinn. Dann hat man für die gesuchte Matrix D_α :

$$\begin{aligned}
 D_\alpha \vec{e}_1 &= \begin{pmatrix} \cos(\alpha) \\ \sin(\alpha) \end{pmatrix}, \\
 D_\alpha \vec{e}_2 &= \begin{pmatrix} -\sin(\alpha) \\ \cos(\alpha) \end{pmatrix}.
 \end{aligned}$$

Es folgt:

$$D_\alpha = \begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix}.$$

Dies ist auch eine orthogonale Matrix: Man sieht das daran, dass die Spaltenvektoren (gleich Bilder der Einheitsvektoren) wieder senkrecht aufeinander stehen und alle beide den Betrag 1 haben.

2) Spiegelung in der Ebene an der x -Achse: Die Matrix dazu sei S_x , dann hat man

$$\begin{aligned} S_x \vec{e}_1 &= \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \\ S_x \vec{e}_2 &= \begin{pmatrix} 0 \\ -1 \end{pmatrix}, \text{ also} \\ S_x &= \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}. \end{aligned}$$

3) Nun sei a die Achse, welche aus der x -Achse durch Drehung D_α (um den Ursprung (!)) entsteht. Wie lautet die Matrix S_a der Spiegelung an der Achse a ? Dazu setzen wir S_a aus schon Bekanntem zusammen:

$$\begin{aligned} S_a &= D_\alpha S_x D_{-\alpha} = \begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} \cos(\alpha) & \sin(\alpha) \\ -\sin(\alpha) & \cos(\alpha) \end{pmatrix} \\ &= \begin{pmatrix} \cos(\alpha) & \sin(\alpha) \\ \sin(\alpha) & -\cos(\alpha) \end{pmatrix} \begin{pmatrix} \cos(\alpha) & \sin(\alpha) \\ -\sin(\alpha) & \cos(\alpha) \end{pmatrix} \\ &= \begin{pmatrix} \cos^2(\alpha) - \sin^2(\alpha) & 2\cos(\alpha)\sin(\alpha) \\ 2\cos(\alpha)\sin(\alpha) & \sin^2(\alpha) - \cos^2(\alpha) \end{pmatrix} \\ &= \begin{pmatrix} \cos(2\alpha) & \sin(2\alpha) \\ \sin(2\alpha) & -\cos(2\alpha) \end{pmatrix}. \end{aligned}$$

4) Die Spiegelung an der xy -Ebene im $\mathbb{R}^3$ hat die Matrix:

$$\begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{pmatrix}.$$

5) Wir betrachten die Drehung um die z -Achse im $\mathbb{R}^3$ mit Winkel α , entgegen dem Uhrzeigersinn, wenn man 'von oben', also von einem Punkt auf der z -Achse mit positiver z -Koordinate aus, auf die xy -Ebene schaut: Die zugehörige Matrix ist

$$\begin{pmatrix} \cos(\alpha) & -\sin(\alpha) & 0 \\ \sin(\alpha) & \cos(\alpha) & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Abbildungsgeometrie und Symmetrien

1. Begriff der Symmetrie und der Symmetriegruppe einer Punktmenge

Menschliche Gesichter sind (mit gewisser Brechung) spiegelsymmetrisch, und diese Symmetrie ist jedem Menschen aus frühester Kindheit bekannt. Fast ebenso geläufig sind Drehsymmetrien. Aber die Mathematik begnügt sich nicht damit, Beispiele für Symmetrien anzugeben, sondern sie will *alle denkbaren Symmetrien* erfassen. Sie stellt sich folgende Aufgaben:

1) Was sind alle denkbaren Symmetrien (zunächst in der Ebene, dann auch im dreidimensionalen Raum)? (Dazu muss zunächst ermittelt werden, was eine Symmetrie überhaupt ist.)

2) Wie kann man alle Symmetrien einer vorgegebenen Figur erfassen? (Um diese Frage praktisch angehen zu können, ist es wichtig, die weitere Frage zu beantworten, welche Struktur die Gesamtheit der Symmetrien einer Figur besitzt.)

Zur Vorbereitung erklären wir zuerst, was eine einzelne Symmetrie einer Figur (d.h. einer Menge von Punkten) ist:

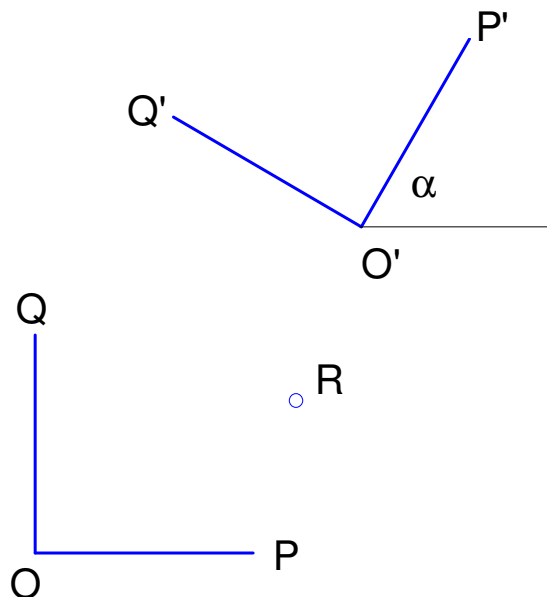
DEFINITION 15 (Begriff der einzelnen Symmetrie als Symmetrieoperation). Eine **Symmetrieoperation** ist eine Abbildung $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ (bzw. $\mathbb{R}^3 \rightarrow \mathbb{R}^3$), welche Längen erhält, für die also erfüllt: Für alle Punkte P, Q gilt: $d(P, Q) = d(P', Q')$, wobei P' das Bild von P und Q' das von Q ist. Solche Abbildungen heißen auch **Isometrien**, auch **Kongruenzabbildungen**. Eine Symmetrie einer Punktmenge $\mathcal{M}$ (im $\mathbb{R}^2$ oder im $\mathbb{R}^3$) ist eine Symmetrieoperation, welche die Menge $\mathcal{M}$ in sich überführt.

Bemerkung: Wesentlich ist also das Verständnis einer Symmetrie als einer Abbildung. Machen wir das an einem Beispiel klar: Ein Quadrat hat offenbar mehrere Spiegelsymmetrien (die zugehörigen Achsen gehen durch die Seitenmitten oder aber durch die Diagonalen). Diese Symmetrien verstehen wir nunmehr direkt als die Abbildungen 'Spiegelung an Achse a ', die jedem Punkt des $\mathbb{R}^2$ wieder einen solchen zuordnen. Eine Spiegelung an einer diagonalen Geraden eines Quadrats überführt nun offenbar das Quadrat in sich, es bildet das Quadrat bijektiv auf sich selbst ab. D.h.: Jedem Punkt des Quadrats wird ein Punkt des Quadrats zugeordnet, verschiedenen Punkten werden verschiedene zugeordnet, und jeder Punkt des Quadrats tritt auch als Bildpunkt auf. D.h. die Abbildung ist auch bijektiv. Außerdem ist eine solche Spiegelung eine Isometrie.

Beispiele: Nicht nur orthogonale Abbildungen sind Symmetrieoperationen, sondern offenbar auch Translationen. Diese und Spiegelungen an Achsen, welche nicht durch den Ursprung gehen, sowie Drehungen um andere Punkte als den Ursprung sind Isometrien, aber nicht linear und auch nicht das Skalarprodukt erhaltend. Aber sie können stets als Hintereinanderschaltungen von Translationen und orthogonalen Abbildungen dargestellt werden. Dazu haben wir den

SATZ 21. 1.) Zwei Kongruenzabbildungen (oder Isometrien) im $\mathbb{R}^2$ sind gleich genau dann, wenn sie ein Zweibein (oder die Eckpunkte eines nicht ausgearteten Dreiecks) gleich abbilden.
 2.) Alle Symmetrieoperationen (oder Isometrien oder Kongruenzabbildungen) im $\mathbb{R}^2$ sind Hintereinanderschaltungen von einer einzigen Translation hinter eine orthogonale Abbildung.
 3.) Jede Symmetrieoperation (oder Isometrie) im $\mathbb{R}^2$ lässt sich als Produkt von höchstens drei Spiegelungen schreiben (deren Achsen nicht durch den Ursprung gehen müssen).
 4.) Zwei Spiegelungen hintereinandergeschaltet ergeben stets eine Translation (wenn die Spiegelachsen parallel sind) oder eine Drehung (um den Schnittpunkt ihrer Achsen, wenn sie sich schneiden).

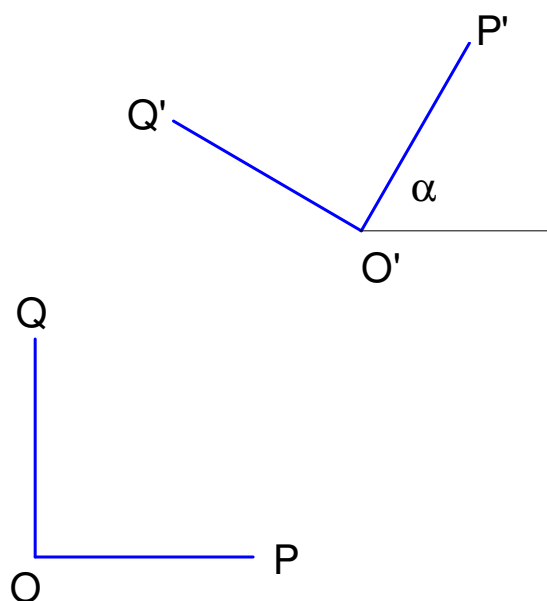
Beweis: Zu 1.): Man betrachte folgendes Bild - es zeigt, wie die Bilder von O, P, Q bei einer beliebigen Kongruenzabbildung allgemein aussehen im Fall erhaltender Orientierung:



Offenbar ist eindeutig bestimmt, wie R abzubilden ist, wenn die Bilder von O, P, Q bestimmt sind, durch die Forderung, dass R' dieselben Abstände zu O', P', Q' hat wie R zu O, P, Q . (Wenn die beliebige Kongruenzabbildung σ die Orientierung umkehrt, ist es genau so.)

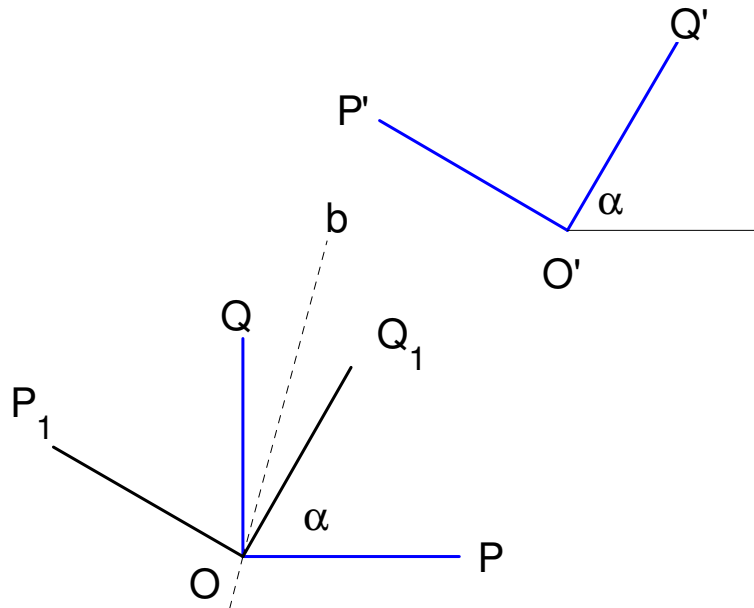
Zu 2.): Eine Isometrie muss ein rechtwinkliges Zweibein mit gleicher Beinlänge in ein solches abbilden. Dann gibt es nur die beiden allgemeinsten Möglichkeiten, die hier skizziert sind:

a) Die Orientierung bleibt gleich:



Dann ist die ganze Kongruenzabbildung σ durch $\sigma(O) = O', \sigma(P) = P', \sigma(Q) = Q'$ schon eindeutig bestimmt. Nun betrachten wir die Drehung $D_{O,\alpha}$. Das ist eine orthogonale Abbildung, weil um den Ursprung O gedreht wird. Dazu nehmen wir $T_{\overrightarrow{OO'}}$. Dann bildet $\tau = T_{\overrightarrow{OO'}} \circ D_{O,\alpha}$ die Punkte O, P, Q auch auf O', P', Q' ab, also $\sigma = \tau$.

b) Die Orientierung wird umgekehrt:



Nun sei wieder $\sigma(O) = O'$ usw. Wir betrachten die Spiegelung S_b an der Achse b , welche durch O geht und von Q und Q_1 gleich weit entfernt ist. (Ihr Winkel zu OP ist offenbar $\alpha + (\pi/2 - \alpha)/2 = \alpha/2 + \pi/4$.) Dann hat man mit $\tau = T_{\overrightarrow{OO'}} \circ S_b$, dass

$$\tau(O) = O', \tau(P) = P', \tau(Q) = Q'.$$

Also $\tau = \sigma$. S_b ist eine orthogonale Abbildung, da b durch den Ursprung geht.

Dass α oben spitz gewählt wurde, tut offenbar nichts zur Sache.

Zu 3.) Diese Aussage folgt im Falle der Orientierungsumkehr unmittelbar aus dem Beweis von 2.), da die Translation sich als Produkt zweier Spiegelungen darstellen lässt. Für den Fall der Orientierungserhaltung müssen wir geschickter vorgehen und offenbar mit nur zwei Spiegelungen auskommen: Betrachten Sie das Bild zu 2, a). Dann kann man dieselbe Kongruenzabbildung erhalten durch: Erstens: Spiegelung an der Mittelsenkrechten der Strecke OO' . Dabei geht OQ über in $O'Q_1$. Zweitens: Spiegelung an der Achse, welche den Winkel zwischen den Schenkeln $O'Q'$ und $O'Q_1$ halbiert. Überzeugen Sie sich mit einer Skizze, dass das klappt.

Zu 4.) Das wurde schon in 4.2 erwähnt und genauer ausgeführt. Man sollte das mit Skizzen bestätigen.

FOLGERUNG 2. Jede Isometrie erhält auch alle Winkel. (Denn alle Spiegelungen tun dies.) Also $\angle(M, P, Q) = \angle(M', P', Q')$ für drei nicht kollineare Punkte M, P, Q .

FOLGERUNG 3. Jede Isometrie erhält entweder den Drehsinn (man sagt auch: die Orientierung) oder kehrt ihn um. Translationen und Drehungen erhalten den Drehsinn, eine Spiegelung kehrt ihn um. Daher kehrt auch ein Produkt von drei Spiegelungen den Drehsinn um.

Für Symmetriebetrachtungen im $\mathbb{R}^2$ ist es nützlich, folgende Grundtypen herauszuheben:

DEFINITION 16. Vier Grundtypen von Symmetrieoperationen (und damit Symmetrien) im $\mathbb{R}^2$, mit geeigneten Bezeichnungen:

- (i) Translation mit Vektor $\vec{a}$, bezeichnet mit $T_{\vec{a}}$,
- (ii) Drehung um den Punkt P mit dem Winkel α entgegen dem Uhrzeigersinn, Bez. $D_{P,\alpha}$,
- (iii) Spiegelung an Achse a , bezeichnet mit S_a ,
- (iv) Gleitspiegelung (auch 'Schubspiegelung') mit Spiegelachse a und Translationsvektor $\vec{b} \parallel a$, $\vec{b} \neq \vec{0}$ (Bez. $G_{a,\vec{b}}$)

Dabei ist definiert: $G_{a,\vec{b}} := T_{\vec{b}} \circ S_a$. Bemerkung: Man hat hier $T_{\vec{b}} \circ S_a = S_a \circ T_{\vec{b}}$.

Bemerkung: Nur 'echte' Gleitspiegelungen wollen wir als solche nehmen, also $G_{a,\vec{b}}$ mit $\vec{b} \neq \vec{0}$, weil $G_{a,\vec{0}} = S_a$ wäre.

Bemerkung: Aus der zweiten Folgerung des vorigen Satzes schließt man, dass Gleitspiegelungen die Orientierung umdrehen. Echte Gleitspiegelungen erfordern die Hintereinanderschaltung von drei Spiegelungen, da sie keine Spiegelungen sind und da zwei Spiegelungen hintereinander geschaltet die Orientierung erhalten.

Die Umkehroperationen zu den Symmetrien dieser Typen sind:

SATZ 22.

$$T_{\vec{a}}^{-1} = T_{-\vec{a}}, \quad D_{P,\alpha}^{-1} = D_{P,-\alpha}, \quad S_a^{-1} = S_a, \quad G_{a,\vec{b}}^{-1} = G_{a,-\vec{b}}.$$

Damit ist die Umkehrabbildung jeder Symmetrieoperation eines Grundtyps wieder vom selben Grundtyp. Wir kennen nunmehr die einzelnen denkbaren Symmetrien recht gut, wenn wir daran denken, dass nur wenige Exemplare der Grundtypen hintereinanderschalten sind.

Nunmehr wenden wir uns der zweiten Frage zu: Welche Struktur hat die Menge aller Symmetrien einer Punktmenge $\mathcal{M}$? Dazu benötigen wir die

DEFINITION 17 (Gruppe). Eine Menge G mit einer Verknüpfung $\cdot : G \times G \rightarrow G$ und einem ausgezeichneten Element e und einer Abbildung, welche jedem Element x von G eindeutig ein Element x^{-1} zuordnet, heißt Gruppe, wenn für alle $x, y, z \in G$ gilt (wir schreiben wie vielfach üblich nur xy statt $x \cdot y$)

$$\begin{aligned} x(yz) &= (xy)z, \\ ex &= x, \\ x^{-1}x &= e. \end{aligned}$$

Dabei heißt e das neutrale Element, und x^{-1} heißt Inverses zu x .

Bemerkung: Wir kennen bereits die additive Gruppe eines Körpers und eines Vektorraums, aber diese sind kommutativ, aber die hier zu betrachtenden Gruppen sind (meist) nicht kommutativ, es werden Gruppen von Abbildungen sein, und die Verknüpfung ist dabei die Hintereinanderschaltung $\circ$.

Jede Gruppe ist nicht leer, sie enthält mindestens das neutrale Element.

DEFINITION 18. Eine Untergruppe einer Gruppe G ist eine nicht leere Teilmenge $H \subset G$, welche unter der Gruppenoperation von G abgeschlossen ist sowie unter der Inversenbildung in G . Das kann man in der Forderung zusammenfassen: Für alle $x, y \in H$ gilt $xy^{-1} \in H$. H ist dann mit den von G 'geerbten' Operationen selbst eine Gruppe. (Die Gruppenaxiome gelten dann für H automatisch, weil sie in G nach Voraussetzung gelten.)

Ein Beispiel, das auch bei Symmetriebetrachtungen eine Rolle spielen wird: Die Gruppe S_n der Permutationen einer Menge von n Elementen auf sich ($n \in \mathbb{N}$) ist mit der Hintereinanderschaltung eine Gruppe. Die Teilmenge A_n der geraden Permutationen (welche sich durch Hintereinanderschaltung einer geraden Zahl von Transpositionen darstellen lassen), ist eine Untergruppe davon.

Nun kommen wir zu der hier relevanten Anwendung:

SATZ 23. Sei $\mathcal{M}$ eine Punktmenge im $\mathbb{R}^2$. Sei $G_S(\mathcal{M})$ die Menge aller Symmetrien. Dann ist $(G_S(\mathcal{M}), \circ, id, (\cdot)^{-1})$ eine Gruppe. (Dabei ist $\circ$ die Hintereinanderschaltung, id ist die identische Abbildung $id(P) = P$ für alle Punkte, und $(\cdot)^{-1}$ ist die Operation, die aus jeder Symmetrie in $G_S(\mathcal{M})$ ihre

Umkehrabbildung macht. Die Menge aller Isometrien ist die Symmetriegruppe vom ganzen $\mathbb{R}^2$, und jede Symmetriegruppe $G_S(\mathcal{M})$ mit $\mathcal{M} \subset \mathbb{R}^2$ ist eine Untergruppe davon. Die so definierte Gruppe $G_S(\mathcal{M})$ heißt Symmetriegruppe von $\mathcal{M}$.

Beweis: Zunächst gilt jedenfalls $id \in G_S(\mathcal{M})$. Weiter hat man: Wenn $\sigma, \tau \in G_S(\mathcal{M})$, also σ, τ Abbildungen sind, welche Längen und Winkel erhalten und $\mathcal{M}$ in sich überführen, dann ist auch τ^{-1} eine solche und $\sigma \circ \tau^{-1}$ eine solche. Konkreter können wir so argumentieren: Jede Symmetrie ist eine Hintereinanderschaltung von Spiegelungen, und $(S_b S_a)^{-1} = S_a^{-1} S_b^{-1} = S_a S_b$. Mit $S_a(\mathcal{M}) = \mathcal{M}$ und $S_b(\mathcal{M}) = \mathcal{M}$ ist auch $S_a S_b(\mathcal{M}) = \mathcal{M}$. Analog ist auch für die Hintereinanderschaltung von drei Spiegelungen zu argumentieren. Ferner erhält jede Spiegelung Längen und Winkel, also bleiben diese auch nach Anwendung einer zweiten Spiegelung unverändert. ■

FOLGERUNG 4. Alle Kongruenzabbildungen $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ bilden die Symmetriegruppe des ganzen Raums $\mathbb{R}^2$ und daher eine Gruppe. Jede Teilmenge $\mathcal{M} \subset \mathbb{R}^2$ hat daher eine Symmetriegruppe, welche eine Untergruppe von $G_S(\mathbb{R}^2)$ ist.

Wenn man die Symmetriegruppe einer Figur beschreiben möchte, so ist es oft lästig, wenn nicht gar unmöglich, alle einzelnen Symmetrien aufzuzählen - insbesondere kann eine Symmetriegruppe unendlich werden. Daher stellen wir die Frage, ob man ein (im günstigen Fall kleines) System von Symmetrien einer Figur nennen kann und damit schon 'alles gesagt' ist. Dazu definieren wir den Begriff eines Erzeugendensystems einer Gruppe:

DEFINITION 19 (Erzeugendensystem einer Gruppe). Eine Teilmenge $U \subset G$ einer Gruppe G heißt Erzeugendensystem für G , wenn jedes Element σ von G sich in der Form schreiben lässt:

$$\sigma = \tau_m \circ \dots \circ \tau_1,$$

mit $\tau_k \in U$ oder $\tau_k^{-1} \in U$ für alle $k = 1, \dots, m$.

Erläuterung zur Nützlichkeit von Erzeugendensystemen: Wenn wir ein Erzeugendensystem U für die Symmetriegruppe $G_S(\mathcal{M})$ einer Punktmenge $\mathcal{M}$ besitzen, so können wir $G_S(\mathcal{M})$ einfach beschreiben als 'die Gruppe, welche von U erzeugt wird'. Daraus, dass $\mathcal{M}$ alle Symmetrien $\tau \in U$ besitzt, folgen dann alle weiteren Symmetrien von $\mathcal{M}$. Denn mit den Symmetrien in U muss $\mathcal{M}$ dann auch alle Symmetrien der Form $\sigma = \tau_m \circ \dots \circ \tau_1$ (mit $\tau_k \in U$ oder $\tau_k^{-1} \in U$ für $k = 1, \dots, m$) besitzen, und auch nur diese. Es ist also die Symmetrie von $\mathcal{M}$ mit U vollständig beschrieben.

Beispiel: Eine Parkettierung der Ebene mit lauter gleichen Parallelogrammen, auf denen noch in der Mitte stets dieselbe unsymmetrische Figur gezeichnet ist, z. B. die Ziffer 7), hat die Symmetriegruppe, welche von den beiden Translationen mit den Kantenvektoren $\vec{a}, \vec{b}$ des Parallelogramms, also von $T_{\vec{a}}, T_{\vec{b}}$ erzeugt wird. Offenbar ist dies Erzeugendensystem auch minimal. Daraus folgen alle anderen möglichen Translationssymmetrien des Musters (und damit alle Symmetrien in diesem Fall), deren Gesamtheit aus allen $T_{k\vec{a}+m\vec{b}}$ besteht mit $k, m \in \mathbb{Z}$. Die Symmetriegruppe ist also isomorph zu $\mathbb{Z} \times \mathbb{Z}$. Nimmt man dagegen nur die Parallelogramme, die keine Rauten und keine Rechtecke sind, so hat man noch die Drehungen um jeden Mittelpunkt, jeden Eckpunkt und jeden Seiten-Mittelpunkt mit Winkel π . Weitere Beispiele finden sich im nächsten Abschnitt.

Zwei Bemerkungen zur Vorsicht: 1) Man sollte nicht 'analoge' Symmetrien (es kann z.B. unendlich viele parallele Spiegelachsen geben) als eine einzige auffassen, sondern sollte nur konkret formulierte einzelne Symmetrien (also z.B. Spiegelungen an einzelnen Achsen usw.) als Erzeugende angeben, wenn man auf ein endliches Erzeugendensystem einer Symmetriegruppe hinaus will. 2) Bei Vektorräumen sind die minimalen Erzeugendensysteme Basen, und je zwei Basen haben gleich viele Elemente. Bei Gruppen ist das anders, auch minimale Erzeugendensysteme haben unter Umständen verschiedene Anzahlen.

Oft lassen sich Erzeugendensysteme angeben auch in solchen Fällen, in denen die Gesamtheit nicht mehr direkt beschreibbar ist. Dabei ist folgender Satz sehr nützlich, der das Auffinden eines Erzeugendensystems und den Nachweis dieser Eigenschaft auf anschaulich-geometrischem Wege ermöglicht. Zunächst benötigt man den Begriff eines Fundamentalgebietes:

DEFINITION 20. Ein Fundamentalgebiet einer Punktmenge $\mathcal{M}$ ist eine Teilmenge $\mathcal{A} \subset \mathcal{M}$, welche durch Aufteilen von $\mathcal{M}$ mittels der Symmetrien von $\mathcal{M}$ gewonnen werden kann, derart, dass $\mathcal{A}$ selbst nicht mehr durch Symmetrien der ganzen Menge $\mathcal{M}$ weiter zu verkleinern ist. Anders gesagt: Aus

einem solchen Fundamentalgebiet $\mathcal{A}$ von $\mathcal{M}$ kann man durch fortgesetztes Anwenden von Symmetrien von $\mathcal{M}$ die ganze Menge $\mathcal{M}$ (geometrisch!) 'erzeugen', eventuell erst mit unendliche vielen Schritten.

Beispiele: Ein einzelnes Parallelogramm der oben besprochenen Parallelogramm-Pflasterung (mit der '7' daraufgemalt) der Ebene ist ein Fundamentalgebiet dieses Musters. Wenn man nur die Parallelogramme nimmt (keine Rauten), so ist ein halbes Parallelogramm ein Fundamentalgebiet. Noch ein etwas merkwürdig anmutendes Beispiel: Wenn $\mathcal{M}$ völlig unsymmetrisch ist, dann ist ganz $\mathcal{M}$ das Fundamentalgebiet, und id ist die einzige Symmetrie, also $G_S(\mathcal{M}) = \{id\}$. Die leere Menge ist dafür ein Erzeugendensystem. Weitere Beispiele folgen im nächsten Abschnitt.

Bemerkung zur Vorsicht: Ein Fundamentalgebiet kann in sich noch symmetrisch sein und als solches damit weiter aufteilbar. Das ist aber irrelevant, wenn es sich bei diesen Symmetrien nicht mehr um Symmetrien der ursprünglichen Figur handelt. Beispiel: Ein Rechteck, das kein Quadrat ist, hat zwei offensichtliche senkrecht aufeinander stehende Spiegelachsen und wird damit in vier kleine Rechtecke aufgeteilt. Jedes davon ist ein Fundamentalgebiet, das in sich wieder die Symmetrien eines Rechtecks hat - aber die Spiegelachsen bzw. der Drehpunkt (für eine Drehung mit π) sind keine für das große Rechteck! Also ist nicht weiter zu zerteilen, und man hat mit den beiden Spiegelungen ein Erzeugendensystem für die Symmetriegruppe des Rechtecks, mit dem folgenden Satz:

SATZ 24. Sei $\mathcal{A} \subset \mathcal{M}$ ein Fundamentalgebiet von $\mathcal{M}$. Sei $U \subset G_S(\mathcal{M})$. Wenn es gelingt, durch fortgesetztes Anwenden von Symmetrien aus U sowie deren Inversen aus $\mathcal{A}$ die ganze Menge $\mathcal{M}$ wiederherzustellen ('geometrisch zu erzeugen'), dann ist U ein Erzeugendensystem für $G_S(\mathcal{M})$.

Beweis: $\mathcal{M}$ ist eine Vereinigung von Kopien von $\mathcal{A}$. Sei $\sigma \in G_S(\mathcal{M})$ beliebig. Dann $\sigma(\mathcal{A}) = \mathcal{B}$, und $\mathcal{B}$ ist eine dieser Kopien. Nach Voraussetzung gibt es also Symmetrien τ_k , $k = 1, \dots, m$, mit $\tau_k \in U$ oder $\tau_k^{-1} \in U$, so dass $\tau_m \circ \dots \circ \tau_1(\mathcal{A}) = \mathcal{B}$. Daher $\sigma(\mathcal{A}) = \tau_m \circ \dots \circ \tau_1(\mathcal{A})$, und es folgt $\sigma = \tau_m \circ \dots \circ \tau_1$, weil σ und $\tau_m \circ \dots \circ \tau_1$ mit $\mathcal{A}$ auch ein Zweibein (für Symmetrien von Punktmengen $\mathcal{M} \subset \mathbb{R}^2$, die nicht schon auf einer Geraden liegen) gleich abbildet. (Analog gilt das auch für andere Dimensionen als 2.) ■

Bemerkung: Man sollte stets damit beginnen, einzelne wesentlich verschiedene Symmetrien einer Figur aufzuzählen und dabei systematisch Exemplare der genannten vier Grundtypen getrennt suchen. Wenn ein Grundtyp nicht vorkommt, ist das auch festzustellen. Dann hat man gute Aussichten, nichts zu übersehen und zu einem korrekten Fundamentalgebiet zu kommen.

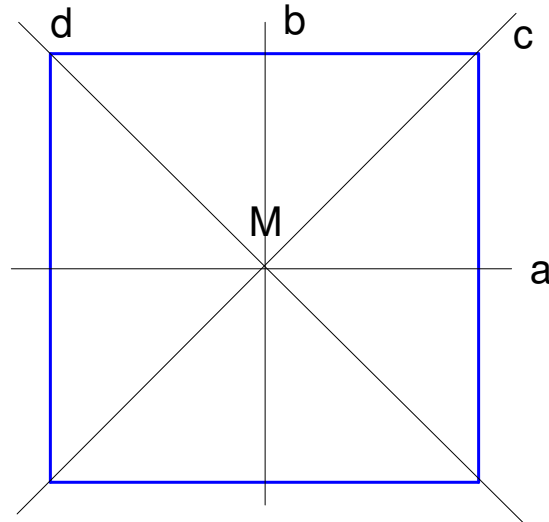
Bemerkung: Damit ist noch keine Minimalität des so gewonnenen Erzeugendensystems garantiert, man kann 'zu viele' Symmetrien zur geometrischen Erzeugung verwandt haben. Aber Grundkenntnisse und Ausprobieren verhilft dann oft auch dazu, ein minimales System zu erhalten. Dazu ein paar praktische Grundregeln:

- 1) Hat man Drehungen um einen Punkt, aber nur mit endlich vielen Winkeln, dann wähle man nur die mit dem kleinsten Winkel.
- 2) Hat man Translationen, so wähle man in der Ebene nur solche mit Vektoren kleinster Länge und in verschiedenen Richtungen, sofern es solche gibt.
- 3) Hat man mehrere Spiegelachsen durch einen Punkt, so verwende man solche in minimalem Winkel und bedenke, dass sich Drehungen daraus ergeben.
- 4) Hat man Spiegelungen an parallelen Achsen, so wähle man sie in minimalem Abstand und denke daran, dass sich entsprechende Translationen daraus bereits ergeben.
- 5) Für Gleitspiegelungen (Schubspiegelungen) wähle man nur echte und mit Translationsvektor minimaler Länge. Auch die zweimalige Anwendung einer Gleitspiegelung ergibt eine Translation.

2. Beispiele für Symmetriegruppen

2.1. Die Symmetrien eines Quadrats. Wir betrachten zunächst ein einfaches Beispiel: $\mathcal{Q}$ ist ein (nicht ausgeartetes) Quadrat (wir denken an die ganze Fläche). Wir wollen die Symmetriegruppe $G_S(\mathcal{Q})$

bestimmen. Zunächst stellen wir fest (Bezeichnungen im Bild sichtbar):



Das Quadrat hat die Symmetrien id , $D_{M,\pi/2}$, $D_{M,\pi}$, $D_{M,3\pi/2}$. (M ist der Mittelpunkt des Quadrats.) Außerdem hat es die Spiegelsymmetrien S_a, S_b, S_c, S_d mit den Achsen a, b durch die Seitenmitten und c, d durch die Diagonalen. Das sind alle. Denn Umkehrungen und Hintereinanderschaltungen führen nicht hinaus. Im Einzelnen erkennt man: $(D_{M,\pi/2})^{-1} = D_{M,3\pi/2}$,

$(D_{M,\pi})^{-1} = D_{M,\pi}$, und die Spiegelungen sind ihre eigenen Inversen, ebenso $id^{-1} = id$. Weiter $D_{M,\pi/2} \circ D_{M,\pi/2} = D_{M,\pi}$, usw., allgemein $D_{M,k \cdot 2\pi/n} \circ D_{M,m \cdot 2\pi/n} = D_{M,(k+m) \cdot 2\pi/n}$. Nun ist $k+m = s \cdot n + r$, mit $s < n$, also $D_{M,(k+m) \cdot 2\pi/n} = D_{M,s \cdot 2\pi/n}$. Daher ist ein Produkt einer der genannten Drehungen mit einer anderen stets unter den genannten Drehungen. Ferner wissen wir schon, dass zwei Spiegelungen der vier genannten hintereinandergeschaltet eine der genannten Drehungen ergibt. Schließlich ergibt eine der Spiegelungen nach einer der Drehungen wieder eine der Spiegelungen, z.B. $S_c \circ D_{M,\pi/2} = S_a$.

Dies war ziemlich mühsam. Die Sache würde sich wesentlich vereinfachen, wenn man nur die Untergruppe der Drehsymmetrien nähme, die viel schneller einzusehen ist, und folgenden Satz heranzieht:

SATZ 25. Wenn eine Figur $\mathcal{M}$ (im dreidimensionalen Anschauungsraum) Drehsymmetrien mit Winkeln $2\pi/n_i$ um die Achsen a_i hat, $i = 1, \dots, N$, dann hat die Figur $1 + \sum_{i=1}^N (n_i - 1)$ Drehsymmetrien (einschließlich id). Wenn sich zusätzlich alle Drehungen durch Spiegelungen der Figur ergeben, dann hat die Symmetriegruppe genau $2 \left(1 + \sum_{i=1}^N (n_i - 1) \right)$ Elemente, also

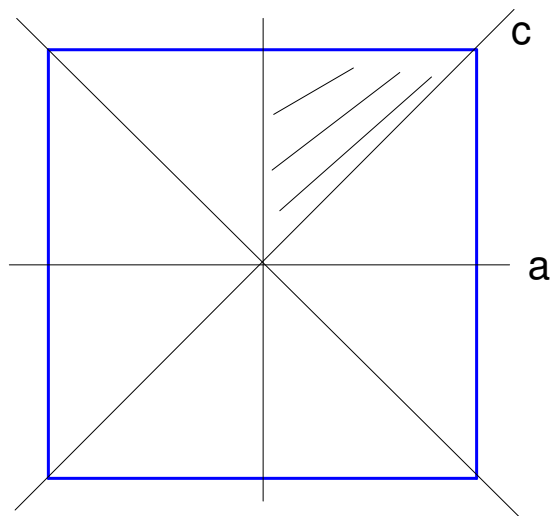
$$|S_G(\mathcal{M})| = 2 \left(1 + \sum_{i=1}^N (n_i - 1) \right).$$

Angewandt auf das Quadrat: Hier haben wir eine Achse durch den Mittelpunkt senkrecht zum Quadrat, und $n = 4$, also hat die Drehgruppe $1 + 3$ Elemente, und die gesamte Symmetriegruppe hat 8 Elemente. Daher muss die oben gegebene Aufzählung vollständig sein, und es bedarf keiner weiteren Prüfung, ob die Hintereinanderschaltungen nicht hinausführen.

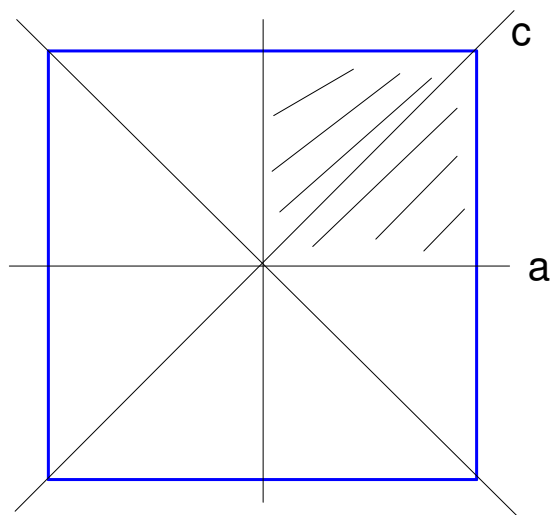
Wir wollen Erzeugendensysteme für die Symmetriegruppe des Quadrats angeben und dafür den letzten Satz des vorigen Abschnitts benutzen:

Ein Erzeugendensystem für die Symmetriegruppe des Quadrats ist $\{S_a, S_c\}$. Ein zweites Erzeugendensystem ist $\{S_a, D_{M,\pi/2}\}$. Wir begründen beide Aussagen mit folgenden Herstellungen des gesamten

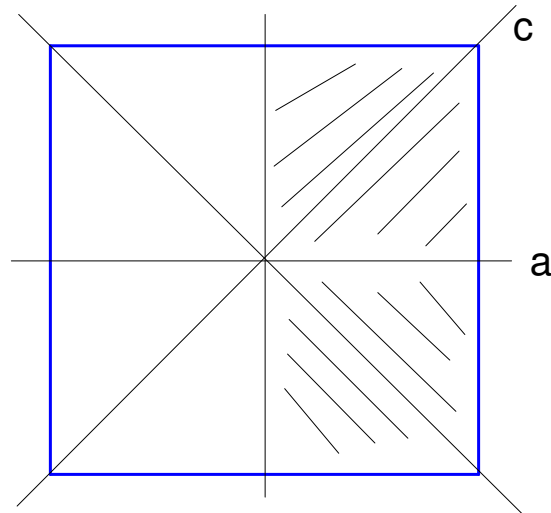
Quadrats aus dem hier (schraffiert) gezeigten Fundamentalgebiet:



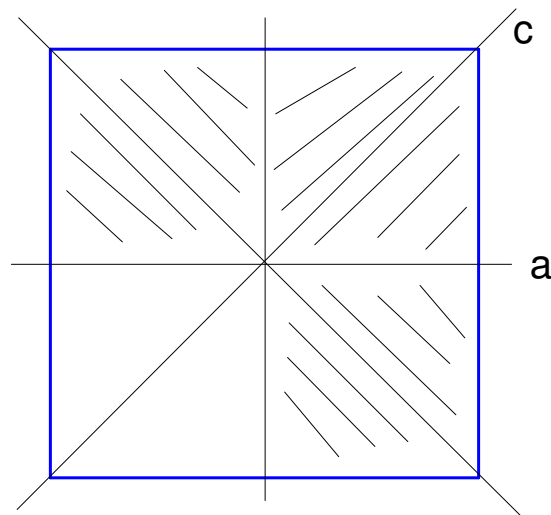
Anwenden von S_c ergibt:



Anwenden von S_a macht daraus:



Wieder S_c angewandt:



Erneute Anwendung von S_a füllt offenbar das ganze Quadrat. Damit ist gezeigt, dass $\{S_a, S_c\}$ ein Erzeugendensystem für $G_S(\mathcal{Q})$ bildet. Dass auch $\{D_{M,\pi/2}, S_a\}$ ein Erzeugendensystem bildet, kann man analog zeigen. Man kann aber auch so argumentieren:

$$\begin{aligned} S_c S_a &= D_{M,\pi/2}, \text{ also} \\ S_c &= D_{M,\pi/2} (S_a)^{-1} \end{aligned}$$

Damit wird S_c von $D_{M,\pi/2}$ und S_a erzeugt. (Es kommt hier nicht darauf an, dass $S_a = S_a^{-1}$, zum Erzeugen dürfen Umkehrabbildungen und Hintereinanderschaltung benutzt werden. Damit folgt die zweite Aussage (dass $\{D_{M,\pi/2}, S_a\}$ ein Erzeugendensystem bildet) aus der ersten (dass $\{S_a, S_c\}$ ein solches ist).

Von vornherein war klar, dass jede Symmetrie des Quadrats einer Permutation der Ecken entsprechen muss. S_4 hat $4! = 24$ Elemente, daher muss $|G_S(\mathcal{Q})|$ ein Teiler von 24 sein. Tatsächlich haben wir also eine echte Untergruppe von der Gruppe S_4 .

2.2. Verallgemeinerung: Die Symmetriegruppe eines regelmäßigen n -Ecks. Die Symmetriegruppe eines regelmäßigen n -Ecks nennt man D_n (n -Diedergruppe). Sie hat offenbar eine Untergruppe von n Drehungen und wegen der Spiegelungen daher $2n$ Elemente. Ein Fundamentalgebiet ist jeweils von zwei Spiegelachsen im kleinstmöglichen Winkel zueinander eingeschlossen (und nur auf einer Seite vom Mittelpunkt aus gesehen). Ein Erzeugendensystem bekommt man jeweils mit zwei Spiegelungen (an Achsen in kleinstmöglichem Winkel zueinander) oder auch mit einer Spiegelung und der Drehung $D_{M, 2\pi/n}$.

2.3. Zur Symmetriegruppe eines Würfels. Bei einem Würfel haben wir:

3 vierzählige Achsen (durch die Seitenmitten)

6 zweizählige Achsen (durch diagonal gegenüberliegende Kantenmitten)

4 dreizählige Achsen (durch diagonal gegenüberliegende Ecken),

das ergibt folgende Anzahl von Drehsymmetrien (einschließlich der Identität):

$$1 + 3 \cdot (4 - 1) + 6 \cdot (2 - 1) + 4 \cdot (3 - 1) = 24,$$

zusammen mit den Spiegelungen also 48 Symmetrien. Das sind viele, aber doch sehr viel weniger als die

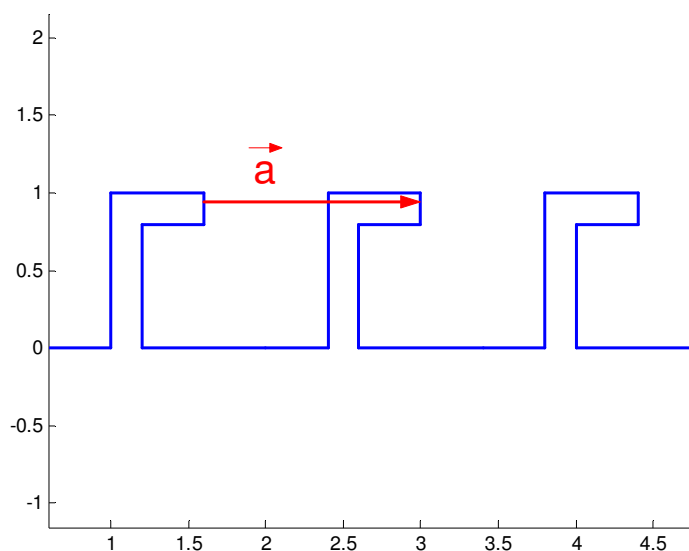
$$|S_8| = 8! = 40320 \text{ Permutationen der Würfecken.}$$

Ebenso kann man die Anzahl der Symmetrien bei den anderen vollkommenen Körpern bestimmen (Tetraeder, Oktaeder, Dodekaeder und Ikosaeder).

3. Symmetrien von Bandmustern

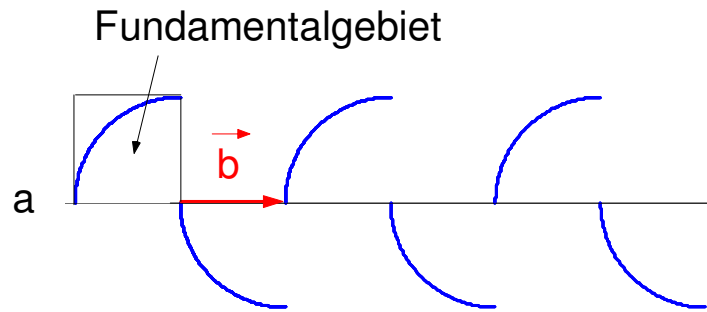
Eine endliche Figur kann offenbar keine nichttriviale (!) Translationssymmetrie ('nichttrivial' heißt: der Translationsvektor ist nicht $\vec{0}$) besitzen, da mit $\vec{a} \neq \vec{0}$ dann auch $T_n \vec{a}$ für alle $n \in \mathbb{N}$ eine Symmetrie sein müsste und damit jeder Punkt auf einen beliebig weit entfernten abgebildet wird. Aber Bandmuster sind gerade definiert durch das Vorhandensein einer Translationssymmetrie in der Richtung des 'Bandes' endlicher konstanter Breite. Man kann herausfinden, wie viele Symmetriegruppen solcher Bandmuster möglich sind. Wir wollen zwei Beispiele vorstellen:

Beispiel 1 (das heißt 'Määnder' in der Kunstgeschichte, in runderer Form auch 'laufender Hund'):



Man denke sich dies nach rechts und links ins Unendliche fortgesetzt, dann hat dies Bandmuster alle Symmetrien $T_{k \cdot \vec{a}}$, $k \in \mathbb{Z}$. Die Symmetriegruppe ist also isomorph zu $(\mathbb{Z}, +)$. Jedes Teilstück der Breite $|\vec{a}|$ bildet ein Fundamentalgebiet, und durch fortgesetztes Anwenden von abwechselnd $T_{\vec{a}}$ und $T_{-\vec{a}} = (T_{\vec{a}})^{-1}$ erzeugt man das ganze Muster aus einem solchen Fundamentalgebiet. Daher wird die Symmetriegruppe des Musters erzeugt von $T_{\vec{a}}$.

Beispiel 2 (auch hier betrachte man das Ganze wieder beidseitig ins Unendliche fortgesetzt):

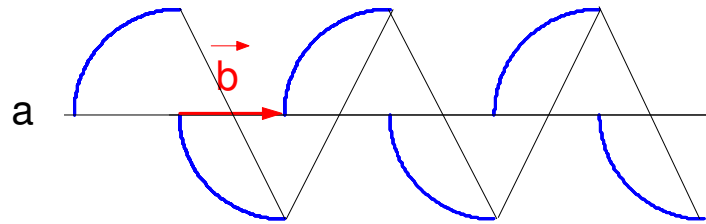


Eingerahmt ist ein Fundamentalgebiet. Die Symmetriegruppe wird erzeugt von $G_{a, \vec{b}}$. Wendet man diese Gleitspiegelung auf das Fundamentalgebiet an, so erhält man den nächsten unteren Bogen. Wendet man dann abwechselnd immer wieder (ins Unendliche fortgesetzt) $G_{a, \vec{b}}$ und $(G_{a, \vec{b}})^{-1}$ an, bekommt man jeweils rechts und links einen Viertelkreisbogen dazu. Also entsteht aus dem Fundamentalgebiet das ganze Muster. Man beachte: Wir haben nicht die Symmetrie S_a und auch nicht die Symmetrie $T_{\vec{b}}$, wohl aber die Hintereinanderschaltung, und das ist $G_{a, \vec{b}}$.

Bemerkung zur Vorsicht: Translationen, Drehungen und Spiegelungen werden im Allgemeinen gut gesehen (kompliziertere Drehungen auch wohl manchmal übersehen), aber Gleitspiegelungen zu entdecken, bedarf einer gewissen Übung. Werden sie übersehen, so bekommt man zu große 'Fundamentalgebiete' (die eben keine mehr sind) und eine zu kleine erzeugte Gruppe, die nicht die ganze Symmetriegruppe umfasst, und man bekommt 'Erzeugendensysteme', die nur eine echte Untergruppe der Symmetriegruppe erzeugen. Im gegebenen Beispiel dürfte man dann an ein doppelt so großes 'Fundamentalsystem' und an die Translation mit $2\vec{b}$, also an die Symmetriegruppe des oben gegebenen Mäanderbeispiels. Aber mit $G_{a, \vec{b}}$ kann man weiter zerteilen, und $T_{2\vec{b}} = G_{a, \vec{b}} \circ G_{a, \vec{b}}$.

Praktischer Hinweis zum Aufspüren einer Gleitspiegelung: Man erkennt eine solche an einer 'Zick-Zack-Linie', mit der man zueinander gehörende markante Punkte auf beiden Seiten der Achse verbinden kann. Allerdings muss man sich bei Mustern, welche die ganze Ebene füllen, noch davon überzeugen, dass die ganze Figur, also auch fernab der Achse, dabei in sich übergeht. Hier ist eine solche

'Zick-Zack-Linie' gezeichnet für das vorliegende Beispiel:



Wenn man diese Linie hat, so geht man mitten hindurch und hat die Achse der Gleitspiegelung. Anschließend findet man leicht die Translation $\vec{b}$ parallel zur Achse.

Beispiel 3 (es werden einfach Rechtecke nebeneinander gelegt, nach rechts und links ins Unendliche fort). Dann wird ein einzelnes Rechteck durch seine beiden Spiegelachsen (durch gegenüberliegende Kantenmittelpunkte) in vier Teile geteilt, einer davon ist ein Fundamentalgebiet, und die Symmetriegruppe wird erzeugt durch eine Spiegelung in Längsrichtung des Musters und zwei weitere Spiegelungen an parallelen Achsen, die beide quer zum Muster liegen, eine dieser Achsen - wir nennen sie a - geht durch die Kantenmitten eines Rechtecks, die andere durch zwei Eckpunkte, welche so nah wie möglich an der Achse a liegen. Übrigens macht es keinen Unterschied, wenn die Rechtecke Quadrate sind; denn die zusätzlichen diagonalen Symmetrieachsen des Quadrats sind keine des Musters.

4. Symmetrien von Pflasterungen der Ebene

Gemeint ist, dass die Ebene durch 'Kacheln' gefüllt wird, indem diese Kacheln immer wieder in zwei verschiedenen Richtungen aneinandergesetzt werden. Natürlich kommt es darauf an, was etwa auf eine Kachel noch eingezeichnet ist. Wir wollen ein paar einfache Beispiele betrachten, dazu auch Fundamentalgebiete und Erzeugendensysteme der Symmetriegruppen angeben.

Das einfachste Beispiel kennen wir schon: Die Punktmenge der gesamten Ebene hat die Gruppe *aller* Symmetrieoperationen (oder Isometrien oder Kongruenzabbildungen) als Symmetriegruppe. Diese Gruppe ist nicht endlich erzeugt, was man allein schon daran sieht, dass die Drehungen mit jedem Winkel um jeden Punkt der Ebene dabei sind. Das sind überabzählbar viele, aber eine endlich erzeugte Gruppe kann höchstens abzählbar unendlich viele Elemente haben.

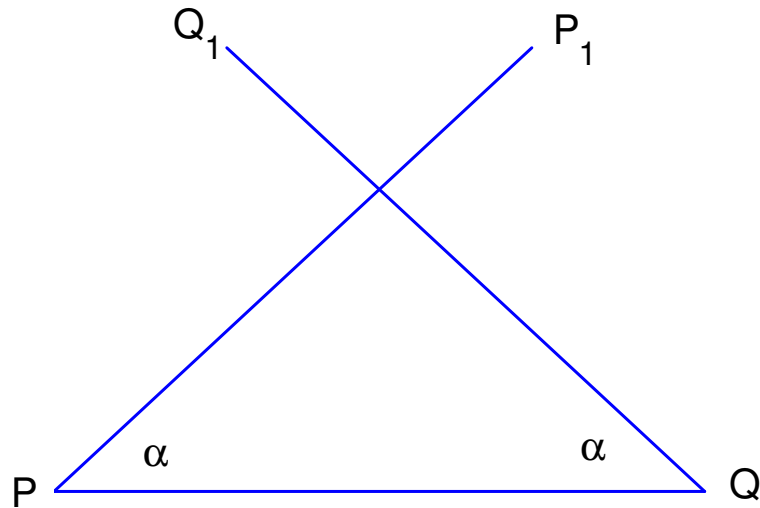
Was uns hier interessiert, sind diskrete Pflasterungen, welche die Eigenschaft besitzen, dass es einen kleinsten Abstand $d_0 > 0$ gibt, so dass ein Punkt P der Muster-Punktmenge in einen anderen Punkt P' durch eine Symmetrie des Musters nur dann abgebildet werden kann, wenn $d(P, P') \geq d_0$. Ein einfaches Beispiel: Quadrate werden aneinandergesetzt, so dass sie die Ebene ausfüllen (das Muster ist dann das Quadratgitter). Dann ist $d_0 =$ Kantenlänge des einzelnen Quadrats. Die Symmetriegruppen sind dann ebenfalls diskret, und sie sind endlich erzeugt.

Zunächst wollen wir einen wichtigen Satz dazu beweisen:

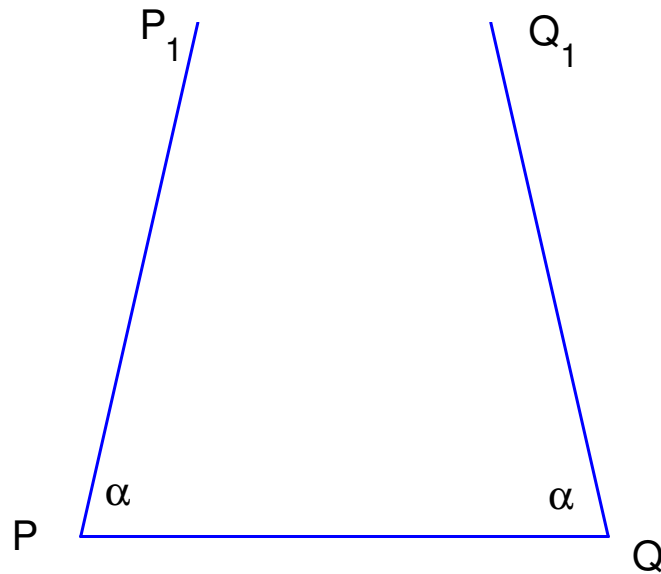
SATZ 26. *Die möglichen Drehwinkel bei diskreten Pflasterungen der Ebene sind Vielfache von $\pi/2$ und von $\pi/3$, keine weiteren.*

Beweis: Es sei P ein Drehpunkt mit Drehwinkel α und d_0 der Minimalabstand. Dann gibt es einen Punkt Q , der ebenfalls Drehpunkt mit Winkel α ist, so dass $d(P, Q) = d_0$. Wäre $0 < \alpha < \pi/3$, dann

wären folgende Punkte P_1, Q_1 ebenfalls Drehpunkte, und man hätte eine Symmetrie der Pflasterung, welche P_1 in Q_1 überführt:



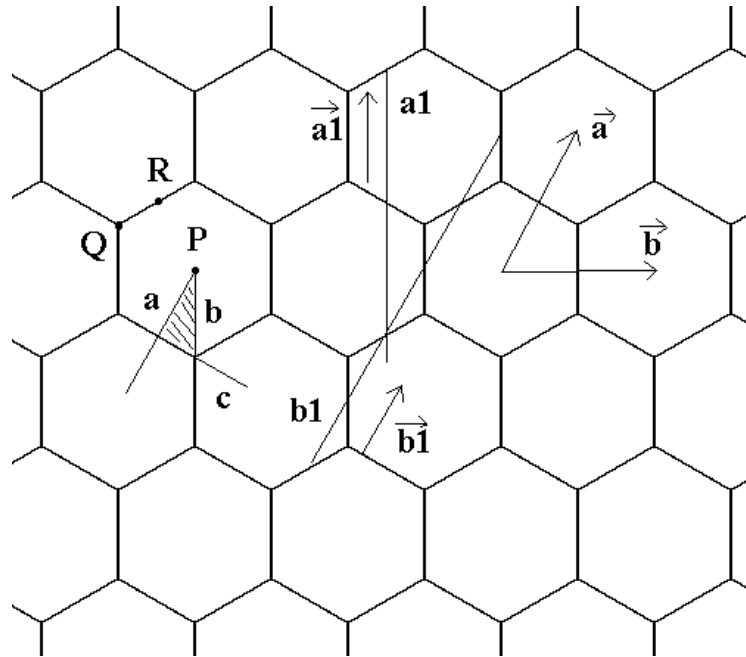
Man denke daran, dass mit der Drehung $D_{Q,\alpha}$ auch $D_{Q,-\alpha} = (D_{Q,\alpha})^{-1}$ eine Symmetrie ist. Aber $d(P_1, Q_1) < d_0$, was unmöglich ist, wenn P_1 in Q_1 durch eine Symmetrie überführbar ist. Daher ist ein Drehwinkel α echt zwischen 0 und $\pi/3$ unmöglich. Völlig analog schließt man mit folgendem Bild, dass ein Drehwinkel α mit $\pi/3 < \alpha < \pi/2$ unmöglich ist:



Aber auch ein Drehwinkel α mit $\pi/2 < \alpha < 2\pi/3$ ist unmöglich; denn dann wäre auch eine Drehung mit 3α möglich, und $3\pi/2 < 3\alpha < 2\pi$. Daher wäre auch Drehung mit -3α möglich, was aber auf einen Drehwinkel echt zwischen 0 und $\pi/2$ hinausläuft. Wir wissen dann schon, dass $-3\alpha = \pi/3$. Damit wäre $-\alpha$ unmöglich und daher auch α unmöglich. Daher ist auch α mit $2\pi/3 < \alpha < \pi$ unmöglich, da Drehung mit -2α wieder eine Drehung mit einem Winkel echt zwischen 0 und $\pi/3$ ergäbe. Für Winkel $> \pi$ nehmen wir den negativen, der ist $< \pi$. ■

Wir stellen aber auch fest, dass alle Vielfachen von $\pi/2$ und $\pi/3$ für diskrete Pflasterungen der Ebene möglich sind; für $\pi/2$ können wir das Quadratmuster nehmen, für $\pi/3$ die Pflasterung der Ebene mit regelmäßigen Sechsecken, die wir nunmehr als Beispiel besprechen:

Beispiel: Die Ebene wird nach Art der Bienenwaben mit regelmäßigen Sechsecken gefüllt:



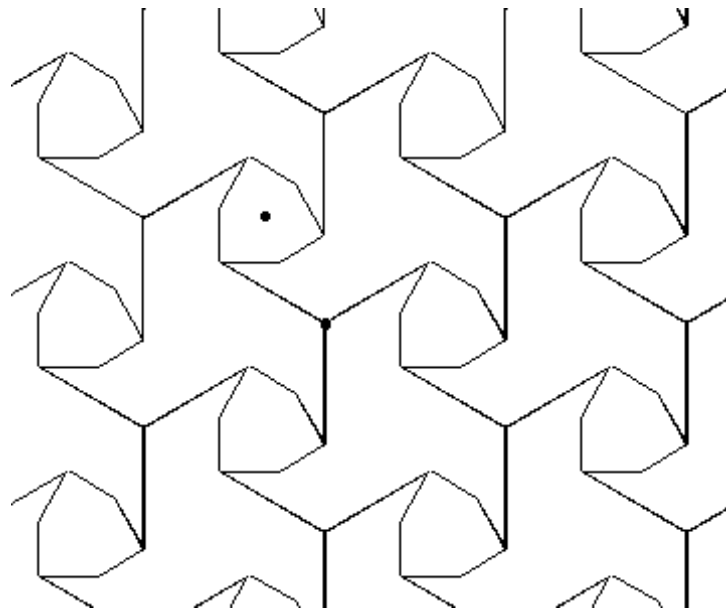
Man erkennt im Einzelnen folgende Repräsentanten der Symmetrietypen, jeweils vollständig für jeden einzelnen Typ genannt:

- 1) Translationen mit $\vec{a}$, $\vec{b}$,
- 2) Drehung um P mit $\pi/3$ (60°), Drehung um Q mit Winkel $2\pi/3$ (120°), Drehung um R mit π (180°),
- 3) Spiegelungen an Achsen a, b, c ,
- 4) Gleitspiegelung an Achse a_1 mit Vektor $\vec{a}_1$ und an Achse a_2 mit Vektor $\vec{a}_2$ (man beachte, dass a_i hier nicht etwa die Länge von $\vec{a}_i$ ist, sondern jeweils eine Gerade bezeichnet).

Ein Fundamentalgebiet (ein solches ist bereits schraffiert eingezeichnet) findet man leicht so: Das ganze Muster entsteht offenbar durch die genannten Translationen aus einem einzigen Sechseck. Nun sind aber sämtliche Symmetrien des einzelnen Sechsecks auch Symmetrien des ganzen Musters! Also zerteilt man wie oben ein einzelnes Sechseck, so dass z.B. das schraffierte Fundamentalgebiet entsteht. Wir wissen daher bereits, dass $\{S_a, S_b, T_{\vec{a}}, T_{\vec{b}}\}$ ein Erzeugendensystem für die Symmetriegruppe ist. Aber das ist nicht minimal. Es genügt bereits $\{S_a, S_b, T_{\vec{a}}\}$. Denn der Vektor $\vec{b}$ entsteht aus $\vec{a}$ durch Drehung mit $\pi/3$, und diese Drehung wird mit $S_b \circ S_a$ produziert. Aber auch Folgendes ist ein Erzeugendensystem: $\{S_a, S_b, S_c\}$. Das wollen wir noch einmal konstruktiv einsehen, indem wir aus dem Fundamentalgebiet das ganze Muster machen durch sukzessives Anwenden dieser drei Spiegelungen: Zunächst erzeugen wir das erste Sechseck durch Anwenden von S_a, S_b abwechselnd. Spiegelung S_c ergibt dann ein zweites Sechseck. Das wird aber dann durch abwechselnde Anwendung von S_a, S_b rings um das erste Sechseck gedreht. Dann haben wir einen Kranz von Sechsecken um das erste. Spiegelung an c liefert mit anschließendem Drehen (wieder nur um P , wieder nur erzeugt mit S_a, S_b) einen größeren Kranz von Sechsecken, usw., usw.

Schließlich erkennt man auch recht einfach, dass $\{S_a, S_b, G_{a_1, \vec{a}_1}\}$ auch ein minimales Erzeugendensystem liefert: Mit S_a, S_b entsteht wieder das ganze Sechseck, mit der Gleitspiegelung und ihrer Inversen bekommt man eine unendliche Doppelreihe (vertikal) von Sechsecken, Spiegelung an b gibt dann eine weitere Reihe links, die Gleitspiegelung dann wieder eine weiter rechts, usw.

Zweites Beispiel:



Dies Beispiel soll vor allem illustrieren, dass man durch Zerstören der Spiegelsymmetrien Muster erzeugen kann, welche eher dynamisch wirken, durch Drehungen, die nicht von Spiegelungen erzeugt sind. Die Symmetriegruppe des gezeigten Musters wird von den Drehungen mit Winkel $2\pi/3$ (120°) um die hervorgehobenen Punkte erzeugt.

Abschließend wollen wir einen Satz zu solchen Mustern noch wenigstens nennen:

SATZ 27. *Es gibt genau siebzehn verschiedene Symmetriegruppen (bis auf Isomorphie) für diskrete ebene Muster, welche die Ebene füllen und zwei unabhängige Translationen besitzen.*

Bemerkung: Es ist natürlich in solchen Fällen leichter, Realisierungen für die siebzehn Symmetriegruppen zu finden, als zu beweisen, dass es keine weiteren gibt.

Inzidenzgeometrie und Eulercharakteristik

Es handelt sich in diesem Kapitel um einen mehr kombinatorischen Aspekt einer abstrakteren Form von Geometrie. Jedoch ist es für einige durchaus sehr praktische Aufgaben nützlich, sich diesem Aspekt ein wenig zu widmen. Grob sieht die Sache so aus: Man hat eine Relation wie 'liegt auf', z.B. findet man das wieder bei 'Der Punkt P liegt auf der Geraden g ', oder auch umgekehrt 'die Gerade g schneidet in den Punkt P ' ('incidere' heißt (lat.) 'Einschneiden', Inzidenz' ist also ein 'Einschneiden'. Aber auch 'Die k -elementige Teilmenge B von A liegt in der m -elementigen Menge C von A ' kann man so auffassen. Kurzum:

DEFINITION 21. *Eine Inzidenzgeometrie ist ein Paar von nicht leeren Mengen $\mathbf{E}$ und $\mathbf{K}$ mit einer ausgezeichneten Teilmenge $I \subset \mathbf{E} \times \mathbf{K}$. Für $(E, K) \in I$ sagt man auch 'E inzidiert mit K' bzw. benennt konkretere Formen von 'Inzidenz' entsprechend. Entweder $(E, K) \in I$ oder $(E, K) \notin I$. Man wird auch die Buchstaben E, K für die Elemente der vorgegebenen Mengen sowie diese Mengen jeweils mit im Kontext suggestiven Buchstaben bezeichnen.*

Bemerkung: Die Buchstaben $\mathbf{E}$ und $\mathbf{K}$, für die Elemente E und K , wurden gewählt, weil hier meist von Ecken und Kanten die Rede sein wird. Außerdem benutzt das Buch von Schwarz/Scheid diese Buchstaben. Handschriftlich sollte man sich helfen mit ' $\mathbf{E}$ ' und ' $\mathbf{K}$ '.

'Irgendeine Relation zwischen Elementen von $\mathbf{E}$ und Elementen von $\mathbf{K}$ ' ist nun so abstrakt, dass man keine interessanten Aussagen erwarten kann. Aber man kann zusätzliche kombinatorischen Forderungen stellen und damit durchaus interessante Sachverhalte modellieren und darüber interessante Aussagen machen. Darunter sind durchaus subsantielle abstraktere geometrische Sachverhalte, die naturgemäß erst in neuerer Mathematik voll entfaltet wurden, aber für konkrete Probleme von Euler schon betrachtet.

Betrachten wir Folgendes: Punkte sind durch 'Kanten' miteinander verbunden, eventuell durch mehrere Kanten - oder auch nicht. Konkretes Beispiel dafür: Zwischen Städten gibt es eine Eisenbahnverbindung - oder auch nicht. Hier sind die Städte die Punkte. Die Relation I ist dann: $(E, K) \in I : \iff E$ ist ein Endpunkt der Kante K . Man nennt die Punkte dann auch gern 'Ecken'. Wenn an jeder Kante genau zwei Punkte (als Endpunkte) liegen, so nennt man diese Inzidenzstruktur einen Graphen, genauer einen ungerichteten Graphen, wenn die Kanten nicht mit Richtungen versehen sind. Wir werden in der Hauptsache einige Aussagen über sogenannte Graphen (als Spezialfälle von Inzidenzgeometrien zeigen und anwenden. Zuvor bemerken wir jedoch eine sehr einfache, doch kombinatorisch fundamentale Tatsache:

SATZ 28. *Sei $(\mathbf{E}, \mathbf{K}, I)$ eine Inzidenzstruktur. Seien $\mathbf{E}$ und $\mathbf{K}$ endliche Mengen. Es gelte nun (Erinnerung: $|A|$ ist die Anzahl von A für eine Menge A)*

$$(E) := |\{K \in \mathbf{K} \mid (E, K) \in I\}| = n \text{ für alle } E \in \mathbf{E}.$$

Dann hat man

$$(i) \quad |I| = n |\mathbf{E}|.$$

Wenn

$$(K) := |\{E \in \mathbf{E} \mid (E, K) \in I\}| = k \text{ für alle } K \in \mathbf{K},$$

dann gilt:

$$(ii) \quad |I| = k |\mathbf{K}|.$$

Wenn beide letzten Bedingungen zutreffen, dann also

$$(iii) \quad |I| = n |\mathbf{E}| = k |\mathbf{K}|$$

Beweis: (i) Unter $(E) = n$ für alle $E \in \mathbf{E}$ hat man $|I| = n |\mathbf{E}|$; denn seien $E_1, \dots, E_r$ die Punkte in $\mathbf{E}$. Dann kann man nach den Voraussetzungen I so aufzählen (ohne ein Paar zu wiederholen!):

$$\begin{aligned} & (E_1, K_{1,1}), (E_1, K_{1,2}), \dots, (E_1, K_{1,n}), \\ & (E_2, K_{2,1}), (E_2, K_{2,2}), \dots, (E_2, K_{2,n}), \\ & \vdots \\ & (E_r, K_{r,1}), (E_r, K_{r,2}), \dots, (E_r, K_{r,n}). \end{aligned}$$

Völlig analog sieht man (ii) - man lasse dann $K_1, \dots, K_s$ durch $\mathbf{K}$ laufen: $|I| = k |\mathbf{K}|$. (iii) folgt aus (i) und (ii). ■

Wir werden diese elementare Tatsache mehrmals anwenden. Hier eine erste

Anwendung (nur von (ii)): Es seien n Gäste bei einem Fest, und r sei die Zahl der Gäste, die einer ungeraden Zahl (anderer Gäste) die Hand schütteln. Behauptung: r ist gerade. **Beweis:** 'Hand schütteln' ist symmetrisch. Also gehören zu jedem Hände schüttelnden Paar genau zwei Personen. Daher mit $\mathbf{E} :=$ Menge der Gäste und $\mathbf{K} :=$ Menge der handschüttelnden Teilmengen der Menge von Gästen, welche aus zwei handschüttelnden Gästen bestehen, und wir definieren

$$(E, \{F, R\}) \in I : \iff E \in (F, R).$$

Dann hat man: $(K) = 2$ für alle $K \in \mathbf{K}$. Also $|I| = 2 |\mathbf{K}|$. Aber nunmehr betrachten wir zu jeder Person die Paare, die sie handschüttelnd mit anderen bildet. Das sind gerade Anzahlen für Personen, die einer geraden Zahl von Leuten die Hand schütteln, und ungerade Anzahlen für die anderen. Nun ist I eine gerade Zahl und eine Summe von geraden Zahlen plus eine Summe von r ungeraden Zahlen. Also muss die zweite Summe gerade sein und damit r gerade. ■

Ebenso folgert man, dass auch die Zahl der Paare, welche sich nicht die Hand schüttelt, gerade sein muss.

DEFINITION 22 (ungerichteter Graph). *Ein ungerichteter Graph ist eine Inzidenzstruktur $(\mathbf{E}, \mathbf{K}, I)$ mit $I \subset \mathbf{E} \times \mathbf{K}$, so dass es zu jeder Kante $K \in \mathbf{K}$ genau zwei Punkte (oder 'Ecken') gibt, die mit K inzidieren, formal also: Für jede Kante $K \in \mathbf{K}$ gibt es $E_1, E_2 \in \mathbf{E}$ mit $E_1 \neq E_2$, so dass $(E_1, K) \in I$ und $(E_2, K) \in I$ und für alle $F \in \mathbf{E}$: Wenn $(F, K) \in I$, dann $F = E_1$ oder $F = E_2$.*

Bemerkungen: Manchmal definiert man einen ungerichteten Graphen auch als Inzidenzstruktur der Form $(\mathbf{E}, \mathbf{E}, I)$, das läuft dann darauf hinaus, dass es zwischen zwei Punkten höchstens eine Verbindungskante gibt. Aber das wäre zu eng für unsere Zwecke. Andererseits ist es für manche Zwecke nötig, die Kanten mit Richtungen zu versehen und so 'gerichtete' Graphen zu behandeln. Aber das brauchen wir hier nicht. Beachten Sie auch, dass es nach Definition hier keine 'Schleifen' gibt, d.h. Kanten, welche von einem Punkt E wieder zu ihm selbst führen.

Ein Graph ist also ein Beispiel einer Inzidenzgeometrie, immer noch sehr allgemein. Aber in diesem Zusammenhang hat man bereits eine Fülle sehr interessanter Probleme. Fangen wir gleich an mit einem 'Jahrhundertproblem': Ein planarer Graph ist ein solcher, den man in die Ebene zeichnen kann (mit 'Kanten' zwischen den 'Ecken', so dass keine Überschneidung von Kanten entsteht. Beispiele dafür sind beliebige 'Landkarten', bei denen die 'Ecken' die Länder sind (oder Staaten) und die Kanten die gemeinsamen Grenzen. Zwei Länder stehen also in der Relation I genau dann, wenn sie eine gemeinsame Grenze haben (die aber mehr als nur ein gemeinsamer Punkt ist (!)). Das 'Jahrhundertproblem': Kann man mit nur vier Farben jede Landkarte so färben, dass keine zwei Länder mit gemeinsamer Grenze dieselbe Farbe bekommen? Das ist das berühmte und äußerst schwierige 'Vierfarbenproblem', das nach unglaublich viel mathematischer Arbeit (mehr als 120 Jahre lang) 1976 gelöst wurde und an dessen Lösung viele berühmte Mathematiker mit wichtigen Beiträgen mitgearbeitet haben - eine besser verständliche Lösung zu produzieren, ist immer noch aktueller Forschungsgegenstand (!) - der Beweis musste eine Berg von menschlich unbewältigbarer Arbeit dem Computer überlassen. Abstrakter kann man es so beschreiben: Kann man die Ecken eines planaren Graphen mit nur vier Farben so einfärben, dass keine zwei mit einer Kante verbundenen Ecken dieselbe Farbe bekommen?

Aber es gibt viel einfachere Probleme, die auch naheliegen und von Interesse sind, dazu immer noch knifflig genug. Z. B. die Frage, ob man wenigstens mit fünf Farben auskommt ist eine solche. Aber es gibt eine Fülle weiterer solcher Probleme, wir nennen zunächst:

1) Unter welcher Bedingung genau kann ein Graph so durchlaufen werden, dass alle Ecken dabei besucht werden (eventuell mehrfach) und *jede Kante genau einmal* durchlaufen wird? (Vgl. dazu das konkrete 'Königsberger Brückenproblem' im nächsten Abschnitt und weitere Beispiele.)

1a) Kann der Graph wie in 1) so durchlaufen werden, dass Start- und Zielpunkt gleich sind?

2) Unter welchen Umständen ist ein Graph 'planar', kann er also in der Ebene kreuzungsfrei dargestellt werden?

Zu 1), 1a) folgende Definition und eine einfache Beobachtung.

DEFINITION 23. *Ein Weg in einem (ungerichteten) Graphen $(\mathbf{E}, \mathbf{K}, I)$ ist eine Folge von Ecken und Kanten der Form $(E_1 K_{1,2} E_2 K_{2,3} E_3 \dots E_{n-1} K_{n-1,n} E_n)$ mit $E_0, \dots, E_n \in \mathbf{E}$ und $K_{i,i+1} \in \mathbf{K}$, für $i = 0, \dots, n-1$, so dass jeweils $(E_i, K_{i,i+1}) \in I$ und $(E_{i+1}, K_{i,i+1}) \in I$. (E_0 heißt dann Anfangspunkt und E_n Endpunkt des Weges.) Der Graph heißt *zusammenhängend*, wenn es für zwei Punkte $E \neq F$ einen Weg von E nach F gibt. Der Graph heißt *unikursal*, wenn es einen Weg gibt, in dem alle Ecken und alle Kanten vorkommen, dabei jede Kante aber nur einmal. Ein solcher Weg heißt *Eulerweg*. Der Graph heißt *geschlossen unikursal*, wenn es einen geschlossenen Eulerweg gibt, also einen Eulerweg, der in einem Punkt E_0 anfängt und endet.*

Bemerkungen: Man beachte, dass für einen Eulerweg nicht verlangt ist, dass die Ecken nicht mehrfach besucht werden. Ein unikursaler Graph muss offenbar zusammenhängend sein. Wir werden uns also nur für zusammenhängende Graphen interessieren. Diese Eigenschaft setzen wir nunmehr voraus. In einem zusammenhängenden Graphen geht also von jedem Punkt mindestens eine Kante aus. Unter dieser Voraussetzung können wir aber einen Eulerweg so beschreiben:

SATZ 29. *Ein Eulerweg in einem zusammenhängenden Graphen ist eine Folge von Kanten $(K_{1,2}, K_{2,3}, \dots, K_{n-1,n})$, so dass jede Kante $K \in \mathbf{K}$ darin genau einmal vorkommt und stets $K_{i,i+1}$ und $K_{i+1,i+2}$ einen gemeinsamen Punkt haben, also ein $F \in \mathbf{E}$ existiert mit $(F, K_{i,i+1}), (F, K_{i+1,i+2}) \in I$, für $1 \leq i \leq n-2$. Geschlossen ist ein solcher Eulerweg offenbar, wenn zusätzlich ein $R \in \mathbf{E}$ existiert mit $(R, K_{1,2}), (R, K_{n-1,n}) \in I$.*

Begründung: Wenn dabei alle Kanten vorkommen, dann auch alle Punkte mindestens einmal als Anfangs- oder Endpunkt einer solchen Kante. Die zweite Bedingung gibt offenbar wieder, dass Anfangs- und Endpunkt des Weges gleich sind. ■

Wir sind nunmehr in der Lage, die Probleme 1) und 1a) zu lösen. Dazu definieren wir noch:

DEFINITION 24. *Die Ordnung eines Punktes E in einem Graphen, dessen Punkte alle auf mindestens einer Kante liegen, ist die Anzahl der Kanten, die von E ausgehen. (Sie ist also stets ≥ 1 .)*

SATZ 30. 1) *Ein zusammenhängender Graph ist unikursal genau dann, wenn die Anzahl der Ecken ungerader Ordnung genau 0 oder genau 2 ist.*

2) *Ein zusammenhängender Graph ist geschlossen unikursal genau dann, wenn es keinen Punkt mit ungerader Ordnung gibt.*

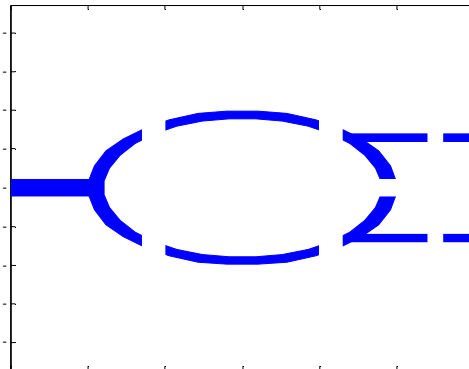
FOLGERUNG 5. *Ein zusammenhängender Graph ist unikursal, aber nicht geschlossen unikursal genau dann, wenn die Zahl der Ecken ungerader Ordnung genau 2 ist. Jeder Eulerweg muss dann in einem dieser Punkte starten und im anderen enden.*

Beweis: Nach dem vorigen Satz kommen die Indizes $2 \dots n-1$ in einem Eulerweg alle genau zwei mal vor, also alle Zwischenpunkte des Weges in gerader Zahl. Sie dürfen sich wiederholen, und das bedeutet, dass alle Zwischenpunkte gerade Ordnung haben müssen, weil im Eulerweg alle Kanten vorkommen. Lediglich Anfangs- und Endpunkt können ungerade Ordnung haben. Dann müssen sie verschieden sein, es gibt dann also genau zwei Ecken ungerader Ordnung, und ein Eulerweg muss vom einen zum anderen führen. Wenn der Eulerweg geschlossen ist, so fallen Anfangs- und Endpunkt zusammen, und dann muss er auch wieder gerade Ordnung haben, da die erste Kante und die letzte von ihm ausgehen, d.h. zwei Kanten plus die gerade Anzahl der weiteren Kanten, an denen der Punkt (als Zwischenpunkt des Eulerweges) liegt. ■

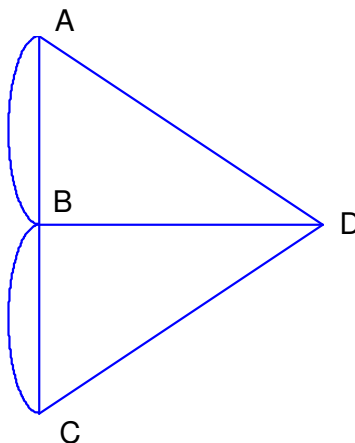
Bemerkung: Wir wissen bereits (die Sache liegt wie beim Händeschütteln), dass die Zahl der Ecken mit ungerader Ordnung gerade sein muss. Aber das kann im allgemeinen eine beliebige gerade Zahl sein, nicht aber bei unikursalen Graphen.

1. Das Königsberger Brückenproblem und Verwandtes

Das kombinatorisch-geometrische Problem: Im alten Königsberg führten sieben Brücken über den Pregel-Fluss, und so stellte sich die Frage, ob man einen Weg über alle Brücken gehen könne, jede Brücke aber nur einmal überquerend. Hier ist ein schematisches Bild davon, immer noch nahe bei einem Stadtplan - die sieben Brücken sind als weiße Unterbrechungen des blauen verzweigten Flusses gezeichnet, die Gebiete benannt, die mit den Brücken verbunden sind:



Um das Problem mathematisch lösen zu können, fasst man es zweckmäßig graphentheoretisch auf, wie Euler das getan hat: Die Stadtgebiete werden als Ecken aufgefasst, die Brücken als Kanten, welche diese Ecken verbinden:

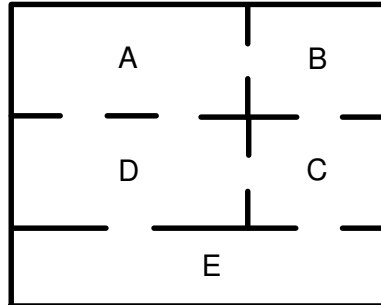


Das Problem stellt sich daher mathematisch so: Ist der skizzierte Graph unikursal? Inspektion der Eckenordnungen zeigt sofort, dass er es nicht ist, weil wir mehr als zwei Ecken ungerader Ordnung haben. Die Eckenordnungen sind passend zu (A, B, C, D) : $(3, 5, 3, 3)$. Das Problem ist also negativ gelöst, ohne dass man langwierig alle erdenklichen Möglichkeiten von Wegen probieren müsste.

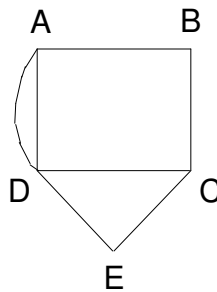
1.1. Das Problem des Museumswärters. Kann man einen Weg durch eine Etage eines Museums gehen, bei dem man jede Tür genau einmal durchläuft (zusätzlich etwa noch beim Ausgangspunkt wieder ankommt)?

Es liegt nahe, die Situation analog zum Brückenproblem zu beschreiben, die Brücken, also die Kanten sind hier die Türen, und die 'Ecken' sind die Museumsräume. Man wird vermuten, dass die Sache vom genauen Grundriss der Etage abhängt. Hier ein Beispiel (Türen sind als Lücken in den Wänden

gezeichnet):

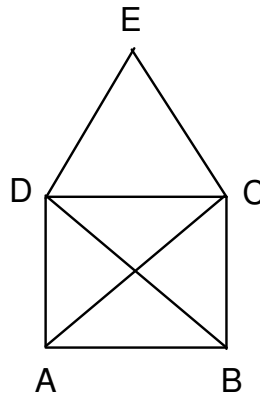


Der zugehörige Graph sieht so aus:



Man sieht die Eckenordnungen $(3, 2, 3, 4, 2)$, zur Punktefolge (A, B, C, D, E) . Also gibt es einen Eulerweg, aber keinen geschlossenen. Man muss bei A starten und bei C aufhören, oder umgekehrt. Einen solchen Weg findet man dann auch leicht, z.B. (A, D, A, B, C, D, E, C) .

1.2. Das 'Haus vom Nikolaus'. Das ist das allseits bekannte Spielchen, folgendes 'Haus vom Nikolaus' in einem Zug zu zeichnen, ohne über eine Kante mehrfach zu laufen:



Wir beobachten genau die beiden Punkte A, B mit ungerader Ordnung. Also müssen wir in A beginnen und in B enden (oder umgekehrt), und nach dem Satz über unikursale Graphen ist auch versprochen, dass das funktioniert. Beispiel eines Eulerweges: $(A, B, C, E, D, A, C, D, B)$, es geht aber auf noch viel mehr Weisen.

Auf dieselbe Weise kann man alle erdenklichen Modifikationen untersuchen, indem man allein die Eckenordnungen heranzieht. Allerdings darf man nicht immer blindwütig zeichnen, auch wenn ein Eulerweg bzw. sogar ein geschlossener existiert. Wir nennen drei Beispiele:

- 1) Wenn man symmetrisch noch unten am 'Haus des Nikolaus' ein weiteres Dach anbringt, so entsteht ein geschlossen-unikursaler Graph. Man kann dann also in einem Zug zeichnen und dabei am Startpunkt enden.
- 2) Zwei 'Häuser vom Nikolaus' nebeneinander, mit gemeinsamer Wand, oder
- 3) zwei mit den Böden aneinanderstoßende 'Häuser vom Nikolaus' - die typische Leserin wird solche und kompliziertere Beispiele nun ohne weiteres selber untersuchen können.

2. Vollständige Graphen und planare Graphen, Eulercharakteristik

Planare Graphen kann man 'entwirrt' in der Ebene zeichnen, also ohne Überschneidung der Kanten. So ist 'das Haus vom Nikolaus' ein planarer Graph (obwohl ursprünglich nicht entwirrt). Damit ist klar, dass 'zu viele Kanten' ein Hindernis für die Eigenschaft der Planarität sind. Man kann sich für die Frage der Planarität auf einfache Graphen beschränken, d.h. solche ohne mehrfache Kantenverbindungen zwischen zwei Punkten, weil man mehrfache Kanten zwischen zwei Punkten beliebig dicht nebeneinander gezeichnet werden könnten, ohne die Planarität zu stören. Wir setzen auch stets wieder zusammenhängende Graphen voraus, was ebenfalls keine Einschränkung bedeutet. Wir definieren nunmehr zwei Klassen wichtiger Beispiele:

DEFINITION 25. *Ein vollständiger Graph ist ein solcher, bei dem jedes Paar von zwei verschiedenen Punkten durch genau eine Kante verbunden ist. K_n ($n \geq 2$) ist der vollständige Graph mit n Punkten. Seien $k, m \geq 1$, $k, m \in \mathbb{N}$. Ein (k, m) -vollständiger Graph hat zwei disjunkte Teilmengen A, B seiner Punktmenge mit je k bzw. m Elementen, so dass jeder Punkt von A mit jedem Punkt von B durch genau eine Kante verbunden ist, aber keine Verbindungskanten zwischen den Punkten von A oder zwischen denen von B bestehen. Der (k, m) -vollständige Graph wird mit $K_{k,m}$ bezeichnet.*

Bemerkung: 'Der vollständige Graph mit n Punkten K_n , 'der (k, m) -vollständige Graph $K_{k,m}$ ' ist sind berechnete Formulierungen, da zwei Graphen mit einer dieser Eigenschaften stets isomorph sind (auch wenn Bilder davon sehr verschieden gezeichnet werden können), d.h. man kann die Punkte und die Kanten bijektiv so aufeinander abbilden, mit $f : \mathbf{E}_1 \rightarrow \mathbf{E}_2$ und $g : \mathbf{K}_1 \rightarrow \mathbf{K}_2$, dass $(E, K) \in I_1 \iff$

$(f(E), g(K)) \in I_2$. Isomorphe Graphen stimmen insbesondere in allen Aussagen über ihre Relation I überein.

Wir haben dazu folgenden (hier nicht zu beweisenden)

SATZ 31. *Ein Graph ist genau dann planar, wenn er keinen Teilgraphen enthält, der K_5 ist oder $K_{3,3}$.*

Bemerkung: Lokal hat man dann einfach zu viele Kanten, als dass man sie überschneidungsfrei zeichnen könnte. Der Satz bestätigt die grobe Anschauung und präzisiert sie. Man überzeuge sich davon, dass man K_4 und $K_{3,2}$ noch zeichnen kann. Es ist klar, was ein Teilgraph sein sollte: eine Teilmenge $E' \subset E$ der Ecken mit allen Kanten zwischen den Punkten aus E' .

Beispiel: Wenn man alle drei Versorger mit Strom, Gas und Wasser mit auch nur drei Haushalten verbinden will, so kann man das nicht mehr planar machen. Der Graph ist $K_{3,3}$ selbst. Wenn nur einer der Haushalte sich mit Strom und Wasser begnügt, so entsteht ein planarer Graph. $K_{3,2}$ und K_4 sind noch planar, was man mit einer kleinen Skizze leicht realisiert. Zwar kann man sich empirisch davon überzeugen, dass K_5 und $K_{3,3}$ (und klar dann auch jeder 'Obergraph' davon) nicht planar sind, aber das ist doch wenig elegant und zudem schon unsicher. Aber man kann folgenden Satz beweisen, der das theoretisch einzusehen gestattet, dazu noch eine fundamentale

DEFINITION 26. *Sei ein beliebiger planarer Graph gegeben. Dann bezeichnen wir mit e die Anzahl der Ecken, mit k die Anzahl der Kanten und mit f die Anzahl der Flächen.*

Bemerkung: Diese Anzahlen werden sich kombinatorisch als höchst bedeutsam erweisen, wie die beiden nächsten Sätze zeigen. Natürlich ist dann $e = |E|$ und $k = |K|$, aber so ist es praktischer.

SATZ 32. *Sei G ein planarer Graph, in dem jede Fläche von genau n Kanten begrenzt wird. Dann gilt:*

$$n(e - 2) = k(n - 2).$$

Bemerkung: Man sieht nicht vor dem nächsten Satz zur Eulercharakteristik, wie man das zeigen sollte, aber dann wird es ganz einfach.

Folgerung: K_5 und $K_{3,3}$ sind nicht planar. Denn für K_5 gilt: $e = 5$, $k = 10$, $n = 3$. Wäre er planar, so müsste $n(e - 2) = 9 = k(n - 2) = 10$ sein, was offensichtlich falsch ist. Für $K_{3,3}$ hat man: $e = 6$, $k = 9$, $n = 4$, bei Planarität müsste also gelten: $n(e - 2) = 16 = k(n - 2) = 18$, was wieder falsch ist.

Nun kommt der Satz von der Eulercharakteristik:

SATZ 33. *1) Sei F eine kompakte Fläche in der Ebene, d.h. beschränkt und abgeschlossen, mit einem Rand, der stückweise glatt ist und endlich viele 'Knicke' haben darf. F habe keine Löcher. Teilt man F auf beliebige Weise in 'Länder' ein, so dass ein planarer Graph entsteht, dann gilt:*

$$\chi(F) = e + f - k = 1^*.$$

Diese Zahl ist eine bedeutsame topologische Invariante, sie bleibt bei beliebigen stetigen Deformationen von F erhalten. Dasselbe gilt für die Zahl der Löcher. (Diese Aussage heißt auch 'Eulersche Polyederformel'.)

2) Sei F eine kompakte Fläche, die man stetig in eine (ganze, wieder ohne Löcher) Kugeloberfläche deformieren kann. Dann hat F die Eulercharakteristik 2, also

$$\chi(F) = e + f - k = 2 \text{ (auch: 'Eulersche Polyederformel')}$$

für jede Einteilung von F in Länder nach Art eines überschneidungsfreien Graph auf der Fläche F . (Hier ist der Begriff 'planar' sinngemäß auf einen anderen Raum als die Ebene verallgemeinert.)

Bemerkung zu *) : Der Satz ist so gemeint, dass die Fläche außerhalb der kompakten Fläche F nicht mitgezählt wird. Zählt man diese mit, dann läuft die Aussage auf 2) hinaus, weil man dann ein planares Netzmodell für eine Landkarte auf der Kugel vor sich hat. Außerdem liegen dann an jeder Kante genau zwei Flächen. Bei dieser Zählung hat man natürlich $e + f - k = 2$, wie in Aussage 2. (So in Schwarz/Scheid.) Wir wollen hier jedoch auf die Tatsache hinweisen, dass die Eulercharakteristik eine topologische Angelegenheit ist, also nur mit der elementaren Geometrie der Topologie zu tun hat, bei der Winkel, Längen, Ecken keine Rolle spielen. Für die Topologie ist gleich, was man durch stetige Deformation ineinander überführen kann. Und hier gilt: Eine kompakte Fläche in der Ebene ohne Loch hat Eulercharakteristik 1, und jedes Loch vermindert die Eulercharakteristik um 1. Die Zahl der Löcher ist

also in der Topologie nicht gleichgültig. Ebenso ist es die Dimension nicht: Die Ebene ist nicht topologisch gleichwertig zur Kugel (wohl ist eine Kugel topologisch dasselbe wie ein Würfel oder auch ein Ikosaeder.) Sondern eine kompakte Menge in der Ebene ist topologisch gleichwertig zu einer Kugel mit einem Loch - und eine Kugel mit Loch hat mit Aussage 2) die Eulercharakteristik 1.

Bevor wir die Aussage 1) beweisen, wollen wir den vorigen Satz daraus herleiten:

Folgerung von Satz 33 aus Satz 34: Für diesen Zweck zählen wir die äußere Fläche mit und haben dann nach der Bemerkung

$$e + f - k = 2,$$

außerdem nach Satz 23, da nunmehr an jede Kante genau zwei Flächen stoßen:

$$nf = 2k.$$

Daraus folgt:

$$\begin{aligned} n(e + f - k) &= 2n, \\ n(e + f - k) &= ne + (2 - n)k, \text{ also} \\ ne + (2 - n)k &= 2n, \text{ daher} \\ n(e - 2) &= k(n - 2). \blacksquare \end{aligned}$$

Beweis von Satz 34 (er ist denkbar einfach), für Teile 1) und 2) zugleich: Wir führen einen Induktionsbeweis über die Zahl der Kanten: 1) und 2) unterscheiden sich nur durch den Induktionsanfang: Für 1) beginnen wir mit einem Dreieck und zählen ab: $e + f - k = 3 + 1 - 3 = 1$ (die Fläche außerhalb des Dreiecks nicht mitgezählt). Für 2) beginnen wir mit einer Kante und seinen Endpunkten und zählen: $e + f - k = 2 + 1 - 1 = 2$. Wenn wir nunmehr eine Kante hinzunehmen, so behaupten wir, dass sich $e + f - k$ nicht verändert. Für die neue Kante gibt es zwei Möglichkeiten:

a) Die neue Kante wird zwischen zwei schon existierende Ecken gezeichnet, die bisher unverbunden waren. Dann hat man, wenn die neuen Ecken- Kanten- und Flächenzahlen mit Strichen bezeichnet werden, die alten ohne Striche:

$$\begin{aligned} e' &= e, \quad f' = f + 1, \quad k' = k + 1 ; \\ \text{also} \quad e' + f' - k' &= e + f - k . \end{aligned}$$

Also hat sich diese Summe nicht verändert.

b) Die neue Kante geht von einer schon vorhandenen Ecke ins Leere, zu einer neuen Ecke, ohne eine feinere Flächeneinteilung hervorzubringen, dann hat man

$$e' = e + 1, \quad f' = f, \quad k' = k + 1.$$

Wieder ist diese Summe unverändert geblieben $\blacksquare$

2.1. Anwendungen der Eulercharakteristik in der Ebene und auf der Kugel. Man kann nunmehr die Invarianz der Eulercharakteristik in vielen Situationen anwenden, weil sie eine grundlegende kombinatorisch-topologische Eigenschaft darstellt. Im Buch von Schwarz/Scheid finden Sie ein Argument, das die Eulercharakteristik benutzt, für folgenden

Satz 34. Wenn man in der Ebene ein gleichmäßiges Parallelogramm-Gitter von Punkten hat und eine Fläche F begrenzt wird durch einen geschlossenen Polygonzug (also Vieleckzug), dessen Ecken sämtlich Gitterpunkte sind, dann hat man Für den Flächeninhalt $I(F)$ mit der Zahl e_i der inneren Gitterpunkte und mit der Zahl e_a der auf dem Rand liegenden Gitterpunkte:

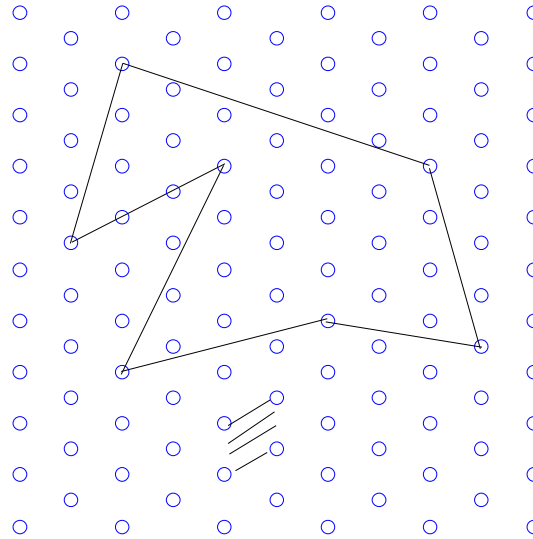
$$I(F) = e_i + \frac{1}{2}e_a - 1 \quad [FE],$$

wobei die Flächeneinheit ein einzelnes der Gitter-Parallelogramme ist. Oder anders gesagt: Wenn F_1 eines der Grundparallelogramme ist, dann gilt:

$$I(F) = \left(e_i + \frac{1}{2}e_a - 1 \right) I(F_1).$$

Es sei bemerkt, dass man diese Tatsache auch anders einsehen kann, vor allem aber hervorgehoben, dass die Sache nicht nur für Quadrate, sondern für beliebige Parallelelogramme funktioniert.

Ein Beispiel zur Anwendung dieses Satzes:



Man zählt:

$$e_i = 23, e_a = 9, \text{ also hat man:}$$

$$F = (23 + 4.5 - 1) F_1 = 17.5 \cdot F_1,$$

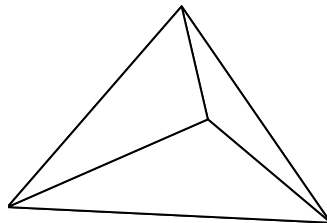
wobei F_1 der Flächeninhalt eines der Grundparallelogramme ist wie z.B. schraffiert. - Man kann auch andere Grundparallelogramme sehen, aber deren Flächeninhalt ist derselbe wie F_1 .

Von zentraler Bedeutung ist die Eulercharakteristik (nebenbei auch unser grundlegender kombinatorischer Satz 23) im Beweis des folgenden Satzes. Bereits in der Antike waren fünf 'vollkommene' Körper bekannt, wobei die 'Vollkommenheit' durch folgende Eigenschaften charakterisiert sind:

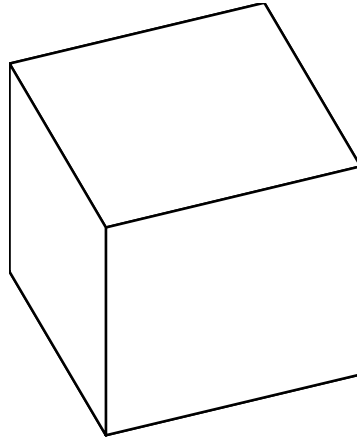
- 1) Jede Seitenfläche ist ein regelmäßiges n - Eck ($n \geq 3$),
- 2) Von jeder Ecke gehen m Kanten aus (offenbar auch $m \geq 3$).

Hier sind die bekannten fünf 'vollkommenen' oder 'platonischen' Körper:

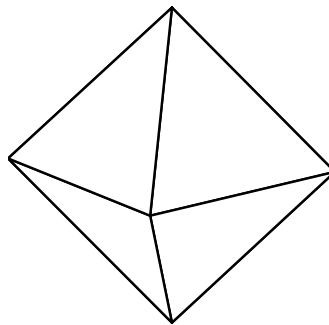
- 1) Tetraeder (von: 'vier Begrenzungspolygone', hier gleichseitige Dreiecke)



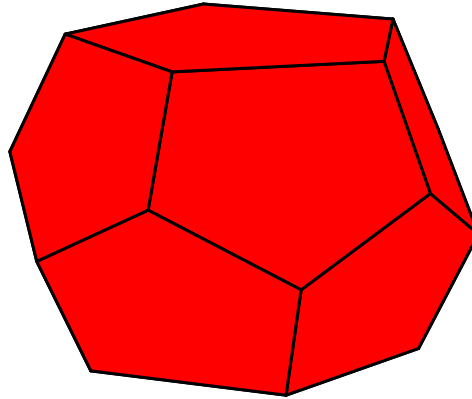
2) Hexaeder oder Würfel (sechs Quadrate sind die Begrenzungsflächen):



3) Oktaeder (acht Begrenzungsflächen):

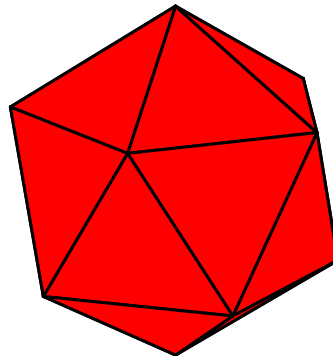


4) Dodekaeder ($\delta\omega\delta\varepsilon\kappa\alpha$ - dodeka - heißt zwölf):



und schließlich

3) Ikosaeder ($\varepsilon\iota\kappa\omicron\sigma\iota$ (eikosi) heißt 'zwanzig'):



Man beachte: Hexaeder und Oktaeder sind dual in dem Sinne, dass der eine durch Vertauschen von Ecken und Flächen aus dem anderen hervorgeht. Das kann man konkret so realisieren, dass man die Seitenflächen-Mittelpunkte des einen mit neuen Kanten verbindet und den anderen bekommt. In derselben Weise sind Dodekaeder und Ikosaeder dual zueinander. (Dualität ist eine wichtige begriffliche Figur der Mathematik, so sind auch Punkte und Geraden dual zueinander in der projektiven Geometrie der Ebene. Auch bei Vektorräumen findet man sie wieder.)

Nun stellt sich die Frage, ob man damit alle vollkommenen Körper besitzt oder ob es noch mehr davon gibt, zwangsläufig viel kompliziertere. Der folgende Satz gibt die Antwort, dass die Liste vollständig ist, es gibt keine weiteren. Man sollte dabei jedoch beachten, dass schon die topologische Struktur der Kugel und davon allein die Eulercharakteristik weitere vollkommenen Körper verhindert, nur die Konstanz der

Zahlen n für alle Flächen und k für alle Ecken verhindert weitere, auch dann, wenn die n -Ecke beliebig schief sein dürften. (Man beachte, dass nur diese schwächeren Voraussetzungen gemacht werden)

SATZ 35. *Wenn von jeder Ecke eines Polyeders m Kanten ausgehen und jede Seitenfläche begrenzt wird durch ein n -Eck, dann gibt die folgende Tabelle alle Möglichkeiten dafür - es sind genau die, welche man auch als vollkommen regelmäßige Polyeder mit den fünf 'vollkommenen Körpern' realisieren kann:*

	n	m	e	k	f
<i>Tetraeder</i>	3	3	4	6	4
<i>Würfel (Hexaeder)</i>	4	3	8	12	6
<i>Oktaeder</i>	3	4	6	12	8
<i>Dodekaeder</i>	5	3	20	30	12
<i>Ikosaeder</i>	3	5	12	30	20

Beweis: Man hat mit Satz 23:

$$nf = 2k, \text{ also } f = \frac{2}{n}k, \quad (1)$$

$$me = 2k, \text{ also } e = \frac{2}{m}k \quad (2)$$

Denn an jeder Kante liegen zwei Flächen und auch zwei Ecken. Wir wenden Satz 23 also an auf die beiden Inzidenzstrukturen Ecken/Kanten und Flächen/Kanten. Mit der Eulerschen Polyederformel gilt

$$e + f - k = 2,$$

also

$$e + f - k = \left(\frac{2}{m} + \frac{2}{n} - 1 \right) k = 2 \quad (3)$$

Es folgt, dass die Zahl

$$\frac{1}{m} + \frac{1}{n} - \frac{1}{2} > 0$$

sein muss, weil k und 2 positive Zahlen sind. Nun sind nach Definition $m, n \geq 3$. Außerdem $\frac{1}{6} + \frac{1}{3} - \frac{1}{2} = 0$ und $\frac{1}{4} + \frac{1}{4} - \frac{1}{2} = 0$, also folgt, dass es für m, n nur folgende Möglichkeiten gibt:

- 1) $m = n = 3$
- 2) $n = 4, m = 3$
- 3) $n = 3, m = 4$
- 4) $n = 5, m = 3$
- 5) $n = 3, m = 5$.

Das sind genau die behauptet einzigen. Aber wir wollen noch sehen, dass sich auch die restliche Information der Tabelle daraus eindeutig ergibt, also die zugehörigen Zahlen e, k, f ausrechnen:

Zu 1): Es liegt stets nahe, (3) für die Bestimmung von k zu nutzen und dann (1), (2) für f und e . Mit $m = n = 3$ also

$$\begin{aligned} k &= \frac{1}{\left(\frac{1}{n} + \frac{1}{m} - \frac{1}{2}\right)} = \frac{1}{1/6} = 6, \\ f &= \frac{2k}{n} = \frac{12}{3} = 4, \\ e &= \frac{2k}{m} = \frac{12}{3} = 4. \end{aligned}$$

Alternativ kann man e natürlich auch über die Eulercharakteristik bekommen, $e = 2 + k - f$.

Analog schließen wir für 2):

$$\begin{aligned} k &= \frac{1}{\left(\frac{1}{4} + \frac{1}{3} - \frac{1}{2}\right)} = \frac{1}{1/12} = 12, \\ f &= \frac{2k}{n} = 6, \quad e = 8, \end{aligned}$$

daraus ergibt sich 3) mit Dualität.

Zu 4):

$$k = \frac{1}{\left(\frac{1}{5} + \frac{1}{3} - \frac{1}{2}\right)} = \frac{1}{1/30} = 30,$$

$$f = \frac{60}{5} = 12, \quad e = 20,$$

5) ergibt sich dann wieder mit Dualität. ■

3. Elementares zu Flächeninhalten und Volumina

1.) Flächeninhalte von Dreiecken und Parallelogrammen

Sie wurden bereits besprochen.

2.) Flächeninhalte von Trapezen:

Ein Trapez ist ein Parallelogramm mit zwei parallelen Seiten. Sind deren Längen a und b und ist der Abstand zwischen den Geraden durch diese Seiten h , so sieht man mit einer Skizze leicht ein, dass für den Flächeninhalt $I(F)$ der Trapezfläche gilt: $I(F) = \frac{a+b}{2}h$.

3.) Flächeninhalte von allgemeineren Polygonen setzt man aus denen von Dreiecken und etwa noch größeren Figuren zusammen. Speziell für den Flächeninhalt eines Polygons, dessen Ecken auf einem Parallelogrammgitter liegen, haben wir eine einfache Formel kennengelernt.

4.) Der Flächeninhalt eines Kreises vom Radius R ist $R^2\pi$, sein Umfang ist $2R\pi$.

5.) Der Oberflächeninhalt einer Kugel vom Radius R ist $4\pi R^2$.

6.) Das Volumen einer Kugel vom Radius R ist $\frac{4}{3}\pi R^3$.

7.) Das Volumen eines Zylinders mit Querschnittsflächeninhalt A ist Ah , wobei h die Höhe des Zylinders ist (senkrecht zu den Querschnittsflächen zu messen!)

8.) Das Volumen eines Kegels mit beliebiger Grundfläche und Grundflächeninhalt G ist $\frac{1}{3}Gh$, wieder mit der Höhe h .

3.1. Scherung und allgemeineres Cavalieri-Prinzip: Wenn man in einem Dreieck (ähnlich auch für Parallelogramme) eine Grundseite der Länge a hat und die Höhe h , so ist der Flächeninhalt $\frac{1}{2}ah$. Wenn man diese Aussage nur für gleichschenklige Dreiecke hätte, so folgt sie über Scherung allgemein. Denn dann ersetzt man die Strecken parallel zu a , aus denen das Dreieck zusammengesetzt ist, durch gleich lange. Völlig analog geht das mit Volumina allgemeiner Zylinder oder Kegel: Jede Querschnittsfläche wird bei Scherung zu einer mit gleichem Flächeninhalt. Man kann dabei auch Allgemeineres tun als scheren und etwa das Ganze zu einer S -Form deformieren (im Aufriss gesehen).

3.2. Die Idee des Integrierens. Die Flächeninhalte oder Volumina von krummlinig Begrenztem erfordern prinzipiell etwas Analysis. Beispiel: Die Länge einer stetig differenzierbaren Kurve $\vec{x}(t)$, $a \leq t \leq b$, ist

$$L = \int_a^b |\vec{x}'(t)| dt.$$

Das kann man auch genauer begründen, wir begnügen uns mit folgendem Hinweis: Der Vektor $\vec{x}'(t)$ ist am Kurvenpunkt $\vec{x}(t)$ tangential zur Kurvenbahn. Ersetzt man nun die Kurve durch entsprechende gerade Stückchen, so bekommt man die Näherung - dabei ist das Intervall $[a, b]$ eingeteilt in $a = t_0 < t_1 < \dots < t_n = b$, und $\Delta t_i = t_{i+1} - t_i$, für $0 \leq i < n$:

$$L \approx \sum_i |\vec{x}'(t_i)| \Delta t_i,$$

Grenzwertbildung für $\max(\Delta t_i) \rightarrow 0$ liefert das oben stehende Integral, die Näherungssumme ist eine zugehörige Riemann-Summe. Anwendung auf das Beispiel der Kreiskurve vom Radius R ergibt das:

$$\vec{x}(t) = \begin{pmatrix} R \cos(t) \\ R \sin(t) \end{pmatrix}, \quad 0 \leq t < 2\pi,$$

also mit

$$\vec{x}'(t) = R \begin{pmatrix} -\sin(t) \\ \cos(t) \end{pmatrix} \quad \text{und} \quad |\vec{x}'(t)| = R \cdot 1 = R$$

hat man

$$L = \int_0^{2\pi} |\vec{x}'(t)| dt = \int_0^{2\pi} R dt = 2\pi R,$$

und damit hat man den Umfang des Kreises berechnet.

Daraus können wir aber mit einer weiteren Integration den Inhalt $I(F)$ der Kreisfläche F vom Radius R so ausrechnen: Geometrische Idee ist es, die Umfänge aller konzentrischen Kreise innerhalb des angegebenen mit dem Wert $L(F_r) = 2\pi r$ per Integration 'aufzusummieren' (man nennt das auch: 'Aufintegrieren'):

$$\begin{aligned} I(F) &= \int_0^R L(F_r) dr = \int_0^R 2\pi r dr \\ &= 2\pi \left[\frac{1}{2} r^2 \right]_0^R = \pi R^2. \end{aligned}$$

Damit ist die Formel für den Kreisflächeninhalt hergeleitet. Außerdem sieht man, dass nicht von ungefähr die Ableitung von πr^2 nach r gerade der Umfang $2\pi r$ ist.

Völlig analog verhalten sich Kugelvolumen und Oberflächeninhalt zueinander:

$$\frac{d}{dr} \left(\frac{4}{3} \pi r^3 \right) = 4\pi r^2.$$

Wir leiten das Kugelvolumen aus dem Oberflächeninhalt der Kugeln her, denken uns den Kugelkörper zusammengesetzt aus lauter Kugeloberflächen (vom Radius 0 bis zum Radius R) und integrieren auf:

$$V(K_R) = \int_0^R O(K_r) dr = \int_0^R 4\pi r^2 dr = 4\pi \left[\frac{1}{3} r^3 \right]_0^R = \frac{4}{3} \pi R^3.$$

Dabei haben wir mit K_r den Kugelkörper vom Radius r bezeichnet und mit $O(K_r)$ dessen Oberflächeninhalt, $V(K_R)$ meint natürlich das Volumen des Kugelkörpers vom Radius R .

Schließlich wollen wir das allgemeine Kegelvolumen (mit beliebigem Grundflächeninhalt G auf diese Weise herleiten: Dazu lassen wir die Spitze des Kegels den z -Wert Null haben, die Grundfläche den z -Wert h und legen die Grundfläche entsprechend parallel zur xy -Ebene. Mit $G(x)$ bezeichnen wir den Inhalt der zur Grundfläche parallelen Querschnittsfläche auf der Höhe x ($0 \leq x \leq h$) und integrieren auf. Dazu müssen wir nur $G(x)$ bestimmen:

Mit $G(x) = kx^2$ (proportional zu x), Proportionalitätsfaktor k , haben wir mit

$$\begin{aligned} G(0) &= 0, G(h) = G: \\ k &= \frac{1}{h^2} G. \end{aligned}$$

Das ergibt für das Kegelvolumen:

$$V_{Kegel} = \int_0^h \frac{1}{h^2} G x^2 dx = \frac{G}{h^2} \left[\frac{1}{3} x^3 \right]_0^h = \frac{G}{h^2} \cdot \frac{1}{3} h^3 = \frac{1}{3} G h.$$

Sammlung wichtiger Sätze und Formeln

1. Sätze und Formeln zu Dreiecken:

- 1) Satz: Die Summe der Innenwinkel im Dreieck ist 180°
- 2) Satz (Pythagoras): Im rechtwinkligen Dreieck (rechter Winkel zwischen den Seiten mit den Längen a, b gilt: $a^2 + b^2 = c^2$
- 3) Höhensatz: Mit denselben Voraussetzungen wie in 2) und den Hypotenusenabschnitten p, q gilt $h^2 = pq$
- 4) Kathetensatz: Wieder mit denselben Voraussetzungen wie bei 2) und 3) gilt $pc = b^2$ und $qc = a^2$, wobei p der Hypotenusenabschnitt zur Seite der Länge b ist, q entsprechend zu a .
- 5) Kosinussatz: $c^2 = a^2 + b^2 - 2ab \cos(\gamma)$ (in jedem nicht ausgearteten Dreieck, mit den üblichen Bezeichnungen)
- 6) Sinussatz: $\frac{\sin(\alpha)}{a} = \frac{\sin(\beta)}{b} = \frac{\sin(\gamma)}{c}$ (in jedem nicht ausgearteten Dreieck, mit üblichen Bezeichnungen)
- 7) Thales-Satz: Im rechtwinkligen Dreieck mit Hypotenuse c liegt der Punkt C auf dem Thaleskreis um den Mittelpunkt der Hypotenuse.
- 8) Satz von den Mittelpunktswinkel- und Peripheriewinkeln im Kreis (Verallgemeinerung des Thales-Satzes):
Zu einer Sehne im Kreis gehöre der Mittelpunktswinkel α (Schenkel: Strecken von den Enden der Sehne zum Mittelpunkt).
Von den Endpunkten der Sehne werden Schenkel zu einem Kreisbogen außerhalb des Sehnensbogens gezogen. Dann ist der Peripheriewinkel β zwischen diesen Schenkeln: $\beta = \alpha/2$.
- 9) Formeln für den Dreiecks-Flächeninhalt: $I(F) = \frac{1}{2}ab \sin(\gamma) = \frac{1}{2}ch_c$ (übliche Bezeichnungen)

2) Formeln zur Berechnung von Flächeninhalten und Volumina:

- 1) Kreisumfang bei Radius R (oder Länge der Kreiskurve): $U = 2\pi R$
- 2) Kreisflächeninhalt bei Radius R : $I(F) = \pi R^2$
- 3) Oberflächeninhalt der Kugel vom Radius R : $O_{Kugel} = 4\pi R^2$
- 4) Kugelvolumen bei Radius R : $V_{Kugel} = \frac{4}{3}R^3$
- 5) Zylindervolumen bei Querschnittflächeninhalt A und Höhe h : $V_{Zylinder} = A \cdot h$
- 6) Kegelvolumen bei (allgemeiner) Grundfläche mit Inhalt G und Höhe h :

3) Definitionen, Sätze und Formeln zur Vektorrechnung:

Definition der Vektorraumoperationen:

1) $\vec{x} + \vec{y}$ ergibt sich geometrisch durch Hintereinanderschalten, rechnerisch im $\mathbb{R}^2$ bzw. $\mathbb{R}^3$ komponentenweise

2) $\lambda \vec{x}$ ergibt sich geometrisch durch Strecken mit $|\lambda|$ (dazu Umdrehen des Pfeils bei $\lambda < 0$), rechnerisch wieder komponentenweise.

3) Definition des Betrages $|\vec{x}|$ eines Vektors $\vec{x}$ im $\mathbb{R}^n$ mit Komponenten x_i , $1 \leq i \leq n$: $|\vec{x}| = \sqrt{\sum_{i=1}^n x_i^2}$.

4) Definition des Skalarproduktes $\vec{x} \cdot \vec{y}$ von Vektoren $\vec{x}, \vec{y} \in \mathbb{R}^n$ mit Komponenten x_i von $\vec{x}$ und y_i von $\vec{y}$: $\vec{x} \cdot \vec{y} := \sum_{i=1}^n x_i y_i$. Bem. $\vec{x} \cdot \vec{x} = |\vec{x}|^2$ gilt stets.

5) Satz (Skalarprodukt und Winkel):

$\vec{x} \cdot \vec{y} = 0 \iff \vec{x} \perp \vec{y}$, allgemeiner $\cos(\angle(\vec{x}, \vec{y})) = \frac{\vec{x} \cdot \vec{y}}{|\vec{x}| |\vec{y}|}$ für $\vec{x}, \vec{y} \neq \vec{0}$

6) Die senkrechte Projektion von $\vec{b}$ auf $\vec{a} \neq \vec{0}$ ist der Vektor $\frac{\vec{b} \cdot \vec{a}}{\vec{a} \cdot \vec{a}} \vec{a}$. Damit ist $\vec{b} - \frac{\vec{b} \cdot \vec{a}}{\vec{a} \cdot \vec{a}} \vec{a} \perp \vec{a}$.

7) Parametrisierung von Kurvenbahnen: Man braucht genau einen freien Parameter, Beispiele:

7.1 Die Gerade g durch P und Q ist parametrisiert mit $\vec{x}_g(\lambda) = \vec{x}_P + \lambda(\vec{x}_Q - \vec{x}_P)$, $\lambda \in \mathbb{R}$.

Bem. Die Strecke von P nach Q ist ebenso parametrisiert, nur ist λ auf $0 \leq \lambda \leq 1$ einzuschränken.

7.2 Die Kreisbahn (nicht: Fläche) im $\mathbb{R}^2$ um den Punkt M mit Radius R ist parametrisiert mit:

$$\vec{x}(t) = \vec{x}_M + R \begin{pmatrix} \cos(t) \\ \sin(t) \end{pmatrix}, \quad 0 \leq t < 2\pi.$$

8) Parametrisierung von Flächen: Man braucht genau zwei freie Parameter, Beispiele:

8.1 Die Ebene E durch P, Q, R (nicht kollinear) ist parametrisiert mit

$$\vec{x}_E(\lambda, \mu) = \vec{x}_P + \lambda(\vec{x}_Q - \vec{x}_P) + \mu(\vec{x}_R - \vec{x}_P), \quad \lambda, \mu \in \mathbb{R}.$$

Bem: Wenn man nur die Fläche des Parallelogramms mit den Eckpunkten P, Q, R, S und

$\vec{x}_S - \vec{x}_R = \vec{x}_Q - \vec{x}_P$ haben will, so setze man nur $0 \leq \lambda, \mu \leq 1$.

8.2 Die Kreisfläche F in der Ebene mit Mittelpunkt M und Radius R ist parametrisiert mit

$$\vec{x}_F(r, \alpha) = \vec{x}_M + r \begin{pmatrix} \cos(\alpha) \\ \sin(\alpha) \end{pmatrix}, \quad 0 \leq r \leq R, \quad 0 \leq \alpha < 2\pi.$$

9) Parametrisierung von Körpern: Man braucht drei freie Parameter, Bsp.:

9.1 Der Zylinderkörper, dessen Bodenfläche die Kreisfläche um den Ursprung mit Radius R ist und

dessen Höhe h beträgt, d.h. dessen Deckelfläche die Kreisfläche um den Punkt $\begin{pmatrix} 0 \\ 0 \\ h \end{pmatrix}$ ist mit $h > 0$,

wird parametrisiert durch $\vec{x}(r, \alpha, z) = \begin{pmatrix} r \cos(\alpha) \\ r \sin(\alpha) \\ z \end{pmatrix}$, $0 \leq r \leq R$, $0 \leq \alpha < 2\pi$, $0 \leq z \leq h$.

4) Definitionen und Sätze über Matrizen (nur 2×2) und lineare Abbildungen (auf $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ beschränkt):

1) Definition (lineare Abbildung): Eine Abbildung $\vec{f}: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ heißt linear genau dann, wenn
(i) $\vec{f}(\vec{x} + \vec{y}) = \vec{f}(\vec{x}) + \vec{f}(\vec{y})$ und (ii) $\vec{f}(\lambda \vec{x}) = \lambda \vec{f}(\vec{x})$ für alle $\vec{x}, \vec{y} \in \mathbb{R}^2$ und $\lambda \in \mathbb{R}$.

2) Definition 'Matrix mal Vektor': $\begin{pmatrix} a & b \\ c & d \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} := \begin{pmatrix} ax + by \\ cx + dy \end{pmatrix}$

3) Satz: Jede Matrix definiert eine lineare Abbildung, und zu jeder linearen Abbildung $\vec{f}: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ existiert genau eine (2×2) -Matrix A , so dass $\vec{f}(\vec{x}) = A\vec{x}$.

Bem.: Dabei ist A eindeutig bestimmt: $\vec{f}(\vec{e}_1) =$ erste Spalte von A , und $\vec{f}(\vec{e}_2) =$ zweite Spalte von A .
Dabei ist $\vec{e}_1 = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ und $\vec{e}_2 = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$.

4) Satz: Eine Matrix stellt eine bijektive Abbildung genau dann dar, wenn ihre Spaltenvektoren linear unabhängig sind.

5) Def.: Eine Matrix heißt orthogonal genau dann, wenn ihre Spaltenvektoren senkrecht aufeinander stehen und beide Betrag 1 haben. Eine orthogonale Abbildung ist eine solche, die sich durch eine orthogonale Matrix darstellen lässt.

6) Satz: Orthogonale Abbildungen erhalten Längen und Winkel, da sie das Skalarprodukt erhalten.

7) Definition des Matrizenproduktes (hier nur für (2×2) -Matrizen:

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix} \begin{pmatrix} u & v \\ w & x \end{pmatrix} := \begin{pmatrix} au + bw & av + bx \\ cu + dw & cv + dx \end{pmatrix}.$$

Bem.: Die Produktmatrix stellt die Hintereinanderschaltung der Matrixabbildungen dar, d.h. $(AB)\vec{x} = A(B\vec{x})$ für alle $\vec{x} \in \mathbb{R}$. Man beachte, dass im Allgemeinen $AB \neq BA$.

8) Beispiel Drehmatrizen: Die Matrix $\begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix}$ stellt die Drehung um den Ursprung (!) mit dem Winkel α entgegen dem Uhrzeigersinn dar.

9) Beispiel Spiegelungsmatrizen: Die Matrix $\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$ stellt die Spiegelung an der Achse $y = x$ dar. Allgemein ist eine Spiegelungsmatrix (nur für die Spiegelungen an Ursprungsachsen (!)) von der Form $\begin{pmatrix} \cos(\alpha) & \sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix}$ (man beachte die negative Determinante!)

10) Beispiel Streckungsmatrizen: Die Matrix für die Streckung mit dem Ursprung als Zentrum und dem Streckungsfaktor k ist $\begin{pmatrix} k & 0 \\ 0 & k \end{pmatrix}$. (Streckungen mit anderen Zentren sind nicht linear!)

5) Abbildungsgeometrie: Translationen, Drehungen, Spiegelungen und Gleitspiegelungen, Isometrien (Kongruenzabbildungen) dazu auch zentrische Streckungen und Ähnlichkeitsabbildungen in der Ebene (zuweilen allgemeiner für $\mathbb{R}^n$ gefasst):

- 1) Definition von $\circ$ (**Hintereinanderschaltung**): Wenn $\sigma, \tau : \mathbb{R}^n \rightarrow \mathbb{R}^n$, dann $(\sigma \circ \tau)(\vec{x}) := \sigma(\tau(\vec{x}))$ für alle $\vec{x} \in \mathbb{R}^n$. Definition von **Umkehrabbildung** σ^{-1} zu σ , wenn σ bijektiv ist: $\sigma^{-1} \circ \sigma = id$, mit $id(\vec{x}) = \vec{x}$. Bem. Beachte $(\sigma \circ \tau)^{-1} = \tau^{-1} \circ \sigma^{-1}$.
- 2) Die Abbildung $T_{\vec{a}} : \mathbb{R}^n \rightarrow \mathbb{R}^n$ mit $\vec{a} \in \mathbb{R}^n$ ist definiert durch $T_{\vec{a}}(\vec{x}) := \vec{a} + \vec{x}$, für $\vec{x} \in \mathbb{R}^n$. Sie heißt Translation mit $\vec{a}$ oder auch Parallelverschiebung mit $\vec{a}$.
- 3) Die Abbildung $D_{P,\alpha}$ mit einem Punkt P in der Ebene und dem Winkel α ist definiert durch: $D_{P,\alpha}(Q) :=$ Resultat der Drehung von Q um P mit dem Winkel α *entgegen dem Uhrzeigersinn*.
- 4) Die Abbildung S_a für eine Gerade a in der Ebene ist die Spiegelung an a . $S_a(P)$ wird so produziert: Gehe von P senkrecht auf a und dann noch einmal so weit.
- 5) Die Abbildung $G_{a,\vec{b}}$ für eine Gerade a in der Ebene und einen Vektor $\vec{b}$ ist definiert durch: $G_{a,\vec{b}} := T_{\vec{b}} \circ S_a$, also $G_{a,\vec{b}}(\vec{x}) = T_{\vec{b}} \circ S_a(\vec{x}) = \vec{b} + S_a(\vec{x})$ für alle $\vec{x} \in \mathbb{R}^2$.
- 6) Satz.: Wenn a, b parallele Achsen sind und $\vec{c}$ der Vektor, welcher von a in senkrechter Richtung auf b weist und $|\vec{c}|$ der doppelte Abstand zwischen a und b ist, dann gilt: $S_b \circ S_a = T_{\vec{c}}$. Wenn die Achsen a, b sich in P schneiden und a in b gedreht wird mit Drehung um P mit Winkel α *entgegen dem Uhrzeigersinn*, dann ist $S_b \circ S_a = D_{P,2\alpha}$.
- 7) Def.: Eine Isometrie σ des $\mathbb{R}^2$ ist eine Abbildung $\mathbb{R}^2 \rightarrow \mathbb{R}^2$, welche Längen (und damit auch Winkel) erhält, also $|\sigma(\vec{x}) - \sigma(\vec{y})| = |\vec{x} - \vec{y}|$ (analog für $\mathbb{R}^n$). Eine Isometrie heißt auch Kongruenzabbildung. Figuren, welche durch eine Isometrie ineinander überführt werden können, heißen kongruent.
- 8) Satz: Alle Abbildungen der Typen 1) bis 5) sind Isometrien, aber unter den nichttrivialen davon sind nur die Drehungen um den Ursprung und die Spiegelungen an Ursprungsgeraden auch orthogonale Abbildungen (die anderen sind nicht linear).
- 9) Satz: Jede Isometrie des $\mathbb{R}^2$ lässt sich schreiben als $\tau \circ \sigma$, mit einer Translation τ und einer orthogonalen Abbildung σ . Jede solche Isometrie lässt sich schreiben als Produkt von höchstens drei Spiegelungen.
- 10) Def.: Die zentrische Streckung um P mit dem Streckungsfaktor $k > 0$ ist die Abbildung $Z_{P,k}(\vec{x}) = \vec{x}_P + k(\vec{x} - \vec{x}_P)$. Bem. Jede zentrische Streckung ist offenbar zu schreiben als Produkt einer Translation nach einer zentrischen Streckung um den Ursprung.
- 11) Def.: Jedes Produkt von Isometrien und zentrischen Streckungen heißt Ähnlichkeitsabbildung.
- 12) Def.: Zwei Figuren heißen ähnlich, wenn sie durch eine Ähnlichkeitsabbildung ineinander überführt werden können.
- 13) Satz: Jede Ähnlichkeitsabbildung lässt sich schreiben in der Form $\tau \circ Z_{O,k} \circ \sigma$, mit einer Translation τ und einer orthogonalen Abbildung σ . Dabei lässt sich noch $\tau \circ Z_{O,k}$ schreiben als $Z_{P,k}$ mit geeignetem P . Es folgt: Zwei Figuren sind ähnlich genau dann, wenn die eine kongruent zu einer zentrischen Streckung der anderen ist.

6) Gruppen von Abbildungen, Untergruppen und Erzeugendensysteme, Symmetriegruppen von Punktmengen in der Ebene.:

- 1) Definition: Eine **Gruppe** ist eine Struktur $(G, \cdot, ()^{-1}, e)$ mit: $e \in G$ und $\cdot : G \times G \rightarrow G$ mit folgenden Eigenschaften: Für alle $a, b, c \in G$ gilt:
 $a \cdot (b \cdot c) = (a \cdot b) \cdot c$, (man schreibt auch ab statt $a \cdot b$)
 $ea = a$
 $a^{-1}a = e$. Bem: Es folgt dann auch $ae = a$ und $aa^{-1} = e$. Wichtig: $ab = c \iff b = a^{-1}c$ usw.
- 2) Eine Untergruppe von $(G, \cdot, ()^{-1}, e)$ ist eine Teilmenge $U \subset G$ mit $e \in U$ und $a, b \in U \implies a^{-1}b \in U$.
- 2) Beispiele:
- 2.1 Alle bijektiven Abbildungen $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ bilden mit $\circ$ für die Multiplikation eine Gruppe.
- 2.2 Alle bijektiven linearen Abbildungen $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ bilden eine Untergruppe von der Gruppe aus 2.1.
- 2.3 Alle orthogonalen Abbildungen $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ bilden eine Untergruppe von der Gruppe aus 2.2.
- 2.4 Alle Isometrien $\mathbb{R}^2 \rightarrow \mathbb{R}^2$ bilden eine Untergruppe von der Gruppe aus 2.1, aber nicht von den Gruppen in 2.2 und 2.3, doch eine Obergruppe von der aus 2.3.
- 3) Erzeugendensysteme: Eine Menge $W \subset G$ (G eine Gruppe) heißt Erzeugendensystem für G , wenn für alle $g \in G$ gilt: Es gibt $h_1, \dots, h_n$ mit $h_i \in W$ oder $(h_i)^{-1} \in W$, so dass $g = h_n \circ \dots \circ h_1$. Bem. Wenn man e als leeres Produkt auffasst, dann hat die Gruppe $\{e\}$ die leere Menge als Erzeugendensystem.
- 4) Def.: Die Symmetriegruppe $G_S(\mathcal{M})$ einer geometrischen Punktmenge $\mathcal{M} \subset \mathbb{R}^2$ ist die Menge aller Isometrien $\mathbb{R}^2 \rightarrow \mathbb{R}^2$, welche $\mathcal{M}$ in sich überführen (analog für $\mathbb{R}^3$).
 Bem.: Die Gruppenoperation ist wieder die Hintereinanderschaltung, und aus den Eigenschaften der Menge $G_S(\mathcal{M})$ folgt, dass sie mit dieser Gruppenoperation und der Inversenbildung sowie $e := id$ eine Untergruppe von der Gruppe aus 2.4 bildet.
- 5) Beispiele:
- 5.1 Die Symmetriegruppe eines regelmäßigen n -Ecks wird durch zwei Spiegelungen erzeugt.
- 5.2 Die Symmetriegruppe der gewöhnlichen Pflasterung der ganzen Ebene mit Quadraten wird von drei Spiegelungen erzeugt.

7) Die Symmetrien von Kachelungen in der Ebene:

Satz: Eine beschränkte Punktmenge kann keine Translation mit einem Vektor $\vec{a} \neq \vec{0}$ als Symmetrie besitzen.

Satz: Die Symmetriegruppe einer Kachelung der Ebene enthält mindestens zwei unabhängige Translationen, mögliche Drehungen sind aber nur solche mit Drehwinkeln, welche Vielfache von $\pi/3$ oder von $\pi/4$ sind. Die Drehungen können durch Spiegelungen erklärt sein, welche Symmetrien der Kachelung sind, müssen es aber nicht.

Definition: Ein Fundamentalgebiet einer geometrischen Punktmenge ist eine Teilmenge, welche sich durch keine Symmetrie des Ganzen mehr zerlegen lässt.

Satz: Wenn aus einem Fundamentalgebiet $\mathcal{G}$ von $\mathcal{M}$ durch fortwährende Anwendung von den Symmetrien $\sigma_1, \dots, \sigma_n$ sowie $\sigma_1^{-1}, \dots, \sigma_n^{-1}$ die ganze Menge $\mathcal{M}$ (geometrisch!) erzeugt wird, dann ist $\sigma_1, \dots, \sigma_n$ ein Erzeugendensystem der Symmetriegruppe von $\mathcal{M}$.

8) Inzidenzstrukturen und Graphen:

- 1) Definition: Eine Inzidenzstruktur ist ein Tripel $(\mathbb{E}, \mathbb{K}, I)$ mit $I \subset \mathbb{E} \times \mathbb{K}$.
- 2) Satz: (i) Wenn jedes $E \in \mathbb{E}$ mit genau k ($k \geq 1$) Kanten aus $\mathbb{K}$ inzidiert und $\mathbb{E}$ endlich ist, dann $|I| = k \cdot |\mathbb{E}|$
 (ii) Wenn jedes $K \in \mathbb{K}$ mit genau n Punkten ($n \geq 1$) inzidiert und $\mathbb{K}$ endlich ist, dann $|I| = n \cdot |\mathbb{K}|$
- 3) Def.: Ein (ungerichteter) Graph ist eine Inzidenzstruktur, bei welcher zu jedem $K \in \mathbb{K}$ genau zwei Punkte P_1, P_2 mit $P_1 \neq P_2$ existieren, so dass (P_1, K) und (P_2, K) aus I sind.
- 4) Def.: Ein Weg in einem ungerichteten Graphen ist eine Folge von Kanten, wobei jede Kante mit einer Richtung versehen wird und der Anfangspunkt der folgenden Kante jeweils der Endpunkt der früheren ist. Der Anfangspunkt des Weges ist der Anfangspunkt der ersten Kante, der Endpunkt der letzten Kante ist der Endpunkt des Weges. Der Weg heißt geschlossen, wenn sein Anfangspunkt auch der Endpunkt ist.
- 5) Def.: Ein (ungerichteter) Graph heißt zusammenhängend, wenn von jeder Ecke E_1 aus $\mathbb{E}$ ein Weg zu jeder anderen Ecke $E_2 \in \mathbb{E}$ existiert.
- 6) Def.: Ein zusammenhängender Graph heißt unikursal, wenn ein sog. Eulerweg existiert, der jede Kante genau einmal enthält. Der Graph heißt geschlossen-unikursal, wenn ein geschlossener Eulerweg existiert.
- 7) Def.: Die Ordnung einer Ecke ist die Zahl der damit inzidierenden Kanten.
- 8) Satz: Ein zusammenhängender Graph ist genau dann unikursal, wenn er keine oder genau zwei Ecken mit ungerader Ordnung enthält. Enthält er keine Ecke ungerader Ordnung, so ist er geschlossen unikursal. Enthält er genau zwei Ecken ungerader Ordnung, so muss jeder Eulerweg in einer dieser Ecken beginnen und in der anderen enden.
- 9) Def.: Ein Graph ist planar genau dann, wenn er in der Ebene ohne Überschneidungen von Kanten gezeichnet werden kann. (Die Kanten dürfen krumm gezeichnet werden!)
- 10) Satz: Für einen endlichen planaren Graphen auf der Kugel (ohne Loch!) oder einen Polyeder gilt mit $e := |\mathbb{E}|$ und $k := |\mathbb{K}|$ und $f :=$ Anzahl der Flächen, in welche die Kugel damit eingeteilt oder wovon der Polyeder begrenzt ist: $e + f - k = 2$. Dieselbe Formel gilt in der Ebene auch, wenn man die äußere Fläche mitzählt, so dass an jeder Kante genau zwei Flächen sitzen.