

Physik

Basiswissen



lizenziert nur zur Nutzung im Schuljahr 2026/27

Übergangslösung
2026/27
7-10 (G9)

Physik in den verschiedenen Zweigen am Gymnasium (Stand 2026, ohne Gewähr)

Ph-Stunden in Klasse 7 :		2	2	2
in Klasse 8 :		4	2	2
in Klasse 9 :		3	2	-
in Klasse 10 :		3	2	2
In welcher Klasse wird dieses (↓) Thema in diesen (→) Zweigen unterrichtet?		NW	In	SP
Mechanik:	Grundlagen: Dichte, Geschwindigkeit	7	7	7
	Kraft, Federhärte, Ortsfaktor, Vektor-Addition, Reibung, Hebel	8	8	8
	Kraft zerlegen, Reibung quantitativ, Flaschenzug	8		
	Arbeit, Energie, Leistung	8	8	8
	Druck, Schweredruck, Luftdruck, Auftrieb	8	8	8
	Schweredruck und Auftrieb quantitativ	8		
Strahlung:	Grundlagen Optik: Schatten, Lochkamera, Reflexion, Spiegel	7	7	7
	Brechung, Linsen, reelle Bilder	9	8 ¹	8 ¹
	virtuelle Bilder, Linsengleichung, optische Geräte	9		
	Farben, Spektrum el.magn.Wellen	9	10	10
	Strahlungsbilanz Erde (Klimaphysik)	10	10	10
	Astronomie	(10)		
	Ionisierende Strahlung	10	10	10
	Akustik	(8)		
Elektro:	Grundlagen: Magnete, Stromkreise, Wirkungen, ohmsches G.	7	7	7
	Ladung, QUIRP, Schaltungen seriell/parallel	9	9	10
	elektrisches Feld	9	9	
	spezifischer Widerstand	(9)	9	
	Innenwiderstand	(9)		
	Motor, Generator, Transformator	10	10	10
	Halbleiter, Dioden	9	9	
	Transistoren, als Schalter, mit Sensoren	(9)	9	
	Feldeffekttransistoren		9	
	Transistoren mit zeitlicher Dynamik	(9)	10	
	logische Schaltungen	(10)	10	
Wärme:	Temperatur, Celsius, Gay-Lussac, Kelvin	8	8	8
	Gasgleichung, Wärmekapazität, Wärmeausbreitung	8		
	Thermodynamik, Stirling-Motor, Wärmepumpe	10		
	Energieversorgung, Kraftwerke, Energiespeicher, Entwertung	10	10	10

Zweige: NW=NaturWissenschaftlich, In=Informatik, SP=SPrachlich und alle anderen eingeklammerte Klassenstufe: als Wahlthema im Lehrplan für diese Stufe genannt

¹: für das eine Schuljahr 2026/27 müsste hier eine **10** stehen, dafür ionisierende Strahlung fakultativ

Inhalt

Erste Größen&Mechanik	1
Struktur der Physik	1
1.1 Gegenstand der Physik	1
1.2 Die beiden Seiten der Physik	2
1.3 Größen allgemein	2
1.4 Geometrische Größen	2
1.5 Zeit	4
1.6 Masse	5
Größen für Dreisatz-Künstler	7
2.1 Dichte	7
2.2 Geschwindigkeit	12
2.3 Rechnen mit Größen	18
Kraft	20
3.1 Erste Erfahrungen	20
3.2 Dynamische Definition von Kraft	24
3.3 Kraftmessung	24
3.4 Hookesches Gesetz	26
3.5 Reaktionsprinzip	28
3.6 Gewichtskraft	28
3.7 Kraft als Vektor	30
3.8 Reibung	33
Druck	35
4.1 Stempeldruck	35
4.2 Schweredruck	41
4.3 Auftrieb	47
Kraftwandler	49
5.1 Hebel	49
5.2 Flaschenzüge	54
Mechanische Energie	57
6.1 Arbeit	57
6.2 Leistung	63
6.3 Energie	64
6.4 Reibungsarbeit	66
6.5 Wirkungsgrad	66
6.6 Weg-Zeit-Gesetz bei Beschleunigung	67

Optik&Akustik	68
Licht und Schatten	68
7.1 Erfahrungen mit Licht	68
7.2 Lichtausbreitung	68
7.3 Lichtgeschwindigkeit	70
7.4 Schatten	72
7.5 Mond	73
Lochkamera	76
8.1 Aufbau	76
8.2 Abbildungsgleichung	77
8.3 Bildqualität	79
Spiegelbilder	80
9.1 Reflexion an einem ebenen Spiegel	80
9.2 Streuung	81
9.3 Entstehung von Spiegelbildern	81
9.4 Konstruktion von Spiegelbildern	83
9.5 Reflexion an gekrümmten Flächen	84
9.6 Vergleich: Spiegel- und Lochkamerabilder	84
Brechung	85
10.1 Strahlengang durch Grenzflächen	85
10.2 Totalreflexion	90
Abbildung durch Linsen	91
11.1 Strahlengang durch Linsen	91
11.2 Bildentstehung bei Konvexlinsen	95
11.3 Geräte zur Bildaufnahme	99
11.4 Geräte zur optischen Vergrößerung	103
Farben	105
12.1 Spektren	105
12.2 Farbwahrnehmung des Menschen	107
12.3 Farbmischung	109
Akustik	111
13.1 Das Wesen von Schall	111
13.2 Elektrik rund um Schall	113
13.3 Beschreibende Größen	115
13.4 Musikalische Aspekte	118
13.5 Eigenschwingungen	119
Wärmelehre	122
Temperatur	122
14.1 Temperaturempfindung	122
14.2 Einfluss der Temperatur	123
14.3 Thermometer	127
14.4 Längenausdehnung fester Körper	131

Ideale Gase	133
15.1 Erwärmung unter konstantem Druck	133
15.2 Die absolute Temperatur	135
15.3 Gase bei konstanter Temperatur	136
15.4 Zusammenführung beider Gesetze	137
15.5 Die ideale Gasgleichung	138
Innere Energie	141
16.1 Wärme	141
16.2 Änderung von Aggregatzuständen	145
16.3 Übertragung von Wärme	152
Verbrennungsmotoren	157
17.1 Dampfmaschine	157
17.2 Ottomotor	158
17.3 Dieselmotor	159
Elektrizität	160
Magnetismus	160
18.1 Erste Erfahrungen	160
18.2 Innerer Aufbau von Magneten	161
18.3 Magnetfelder	162
Vorstellung von elektrischem Strom	166
19.1 Erste Erfahrungen	166
19.2 Elektrischer Strom	167
19.3 Schaltsymbole und Schaltpläne	168
19.4 Elektronen	169
19.5 Stromrichtung	171
19.6 Leiter und Nichtleiter	172
19.7 Steckdose	173
Wirkungen des elektrischen Stroms	174
20.1 Wärmewirkung	174
20.2 Magnetische Wirkung	175
20.3 Chemische Wirkung	176
Messgrößen des elektrischen Stroms	177
21.1 Die Stromstärke	177
21.2 Die Spannung einer Quelle	182
21.3 Der Widerstand eines Verbrauchers	183
Sicherer Umgang mit Elektrizität	188
22.1 Wirkung von Strom im Menschen	188
22.2 Das Schuko-System	189
22.3 Sicherungen	190
22.4 FI-Schalter	190
Elektrische Ladung	191
23.1 Elektrostatische Kräfte	191
23.2 Verteilung von Ladung	193
23.3 Elektrischer Strom	195
23.4 Die Hauptakteure: Elektronen	197
23.5 Elektrisches Feld	200

Elektrische Größen	203
24.1 <i>QUIRP</i> – Steckbriefe	203
24.2 Spannung vs. Stromstärke	205
24.3 Bestätigende Experimente	207
Elektrische Schaltungen	209
25.1 Serien- und Parallelschaltung	209
25.2 Komplexe Schaltungen	213
25.3 Anwendungen	214
Elektromagnetismus	218
26.1 Motoren	218
26.2 Generatoren	224
26.3 Transformatoren	229
26.4 Unsere Elektrizitätsversorgung	235
Energienutzung	239
27.1 Techniken zur Energiegewinnung	239
27.2 Der Energieverbrauch eines Landes	242
27.3 Wertigkeit von Energieformen	243
27.4 Referate	244
27.5 Wahrscheinlichkeit von Zuständen	245
Elektronik	250
28.1 Grundlegendes über Atomhüllen	250
28.2 Halbleiter	251
28.3 Grenzschichten	255
28.4 Digitalelektronik	264

Vorwort

Dieses Buch hat eine Vergangenheit, eine Gegenwart und – vorgesehen – eine Zukunft.

Seine Vergangenheit sind Arbeitsbücher, die geschrieben wurden, um einen effektiven Physikunterricht an saarländischen Gymnasien zu erleichtern. Passend zu den Lehrplänen gab es je ein Arbeitsbuch für die Klassenstufen 7, 8 und 9. Die vorgesehenen Stoffgebiete pro Klassenstufe waren recht stabil. Sie überstanden sogar den Wechsel vom neunjährigen Gymnasium (G9) zum achtjährigen (G8). Der Stoff im sprachlichen Zweig war stets Teilmenge des Stoffes im naturwissenschaftlichen Zweig, an dem sich die Arbeitsbücher folglich orientiert haben.

Beim Wechsel zurück zu G9 mit Ausbau eines Informatik-Zweiges wurde es unübersichtlicher. Es gibt jetzt drei Zweig-spezifische Lehrpläne, gemäß derer einige Gebiete nicht nur unterschiedlich weit vertieft, sondern auch in unterschiedlichen Klassenstufen unterrichtet werden. Das Prinzip „eine Klassentufe – ein Buch“ kommt an seine Grenzen. Für die Zukunft möchte ich daher ein Buch anbieten, das den gesamten Stoff in rein fachlicher Anordnung vermitteln soll, unabhängig von Zweigen und Klassenstufen. Außerdem will ich es in mancherlei Hinsicht so gründlich überarbeiten, dass diese Zukunft noch in einiger Ferne liegen dürfte.

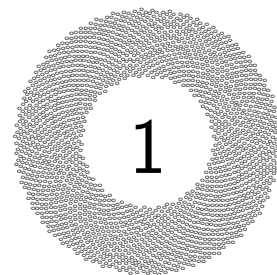
Dazwischen liegt das nun vorliegende Werk. Es ist aus den letzten Fassungen der drei Arbeitsbücher für Klassen 7 bis 9 zusammengesetzt. Deren Kapitel wurden einfach nur thematisch sortiert aneinandergehangen. Es fehlen zum angestrebten Stoff der Klassen 7-10 daher noch die Themen Klimaphysik, Thermodynamik, ionisierende Strahlung und ein wenig Elektronik.

Rückmeldungen aller Art an den Autor sind willkommen: auf kurzem Weg über e-mail an pbarbian@t-online.de oder über den Verlag.

Beckingen, im August 2026

Patrick Barbian

© Softfrutti Verlag, Saarbrücken 2026
P. Barbian:
Physik-Basiswissen
271 E-Book



Struktur der Physik

1.1 Gegenstand der Physik

Der Name „Physik“ stammt vom griechischen Wort *physis*=Natur ab. Die Physik ist schon daran als *Naturwissenschaft* zu erkennen. Bei anderen Naturwissenschaften wie der Biologie, der Geologie, der Astronomie und der Chemie grenzt der Name den (zumindest ursprünglichen) Forschungsgegenstand einigermaßen klar ein: griechisch *bios* (das Leben) und *gā* (die Erde), lateinisch *astra* (die Sterne) und arabisch (*al-*)*kimiya* (Goldmacher, Schwarzkünstler).

Die Physiker hingegen befassen sich nicht mit einem speziellen Bereich der Natur, sondern mit den fundamentalen Grundgesetzen der gesamten Natur. Alle Vorgänge im Universum laufen nämlich nach solchen Naturgesetzen ab. Hinter jeder Reaktion, die ein Chemiker untersucht, hinter jedem Erdbeben, das ein Geologe erforscht, und natürlich auch hinter jedem Gerät, das ein Ingenieur entwickelt, steckt letzten Endes immer dieselbe Physik. Physik ist daher *die* fundamentale Naturwissenschaft und wird ebendarum teilweise als Hauptfach unterrichtet.

Der Trend, Erscheinungen aus verschiedensten Wissenschaftsbereichen als Ausdrücke derselben Naturgesetze zu verstehen, setzt sich dabei auch innerhalb der Physik fort: Während andere Wissenschaften Fortschritte machen, indem sie immer *mehr* Einzelheiten über ihren Forschungsgegenstand erkennen, schreitet die Physik fort, wenn sie mit immer *weniger* Theorien auskommt, um alle Beobachtungen zu erklären.

Obwohl dementsprechend die Grenzen zwischen verschiedenen Arten von Naturerscheinungen immer mehr verwischen, lassen sich doch noch vor allem folgende Teilbereiche der Physik unterscheiden, aufgelistet mit ein paar ihrer Inhalte, wie sie als Stichworte auch im Alltag vorkommen können:

Mechanik: Bewegung von Körpern, Antrieb von Maschinen, Kraft, Arbeit, Druck

Optik: Licht (nicht nur solches, das man sehen kann), Helligkeit, Farben

Akustik: Schall (nicht nur solchen, den man hören kann), Lautstärke, Tonhöhen

Elektrizitätslehre: elektrischer Strom, Volt, Ampere, Generatoren, Ladegerät

Wärmelehre: Wärme, Bewegung von (Gas-)Teilchen, Temperatur, Kelvin

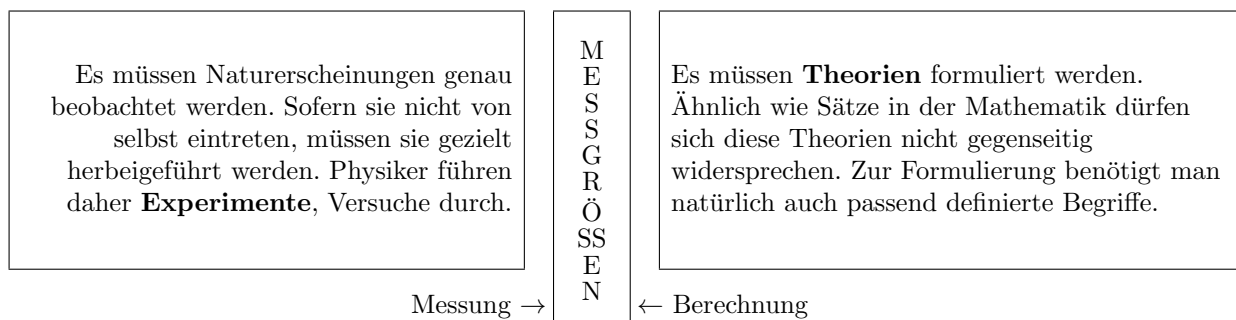
Atomphysik, Kernphysik Atome, mit Schwerpunkt auf deren Hülle bzw. Kern

Festkörperphysik: Aufbau und Eigenschaften fester Körper wie Kristallen

Dabei kommt der Mechanik eine besondere Rolle zu, weil es letzten Endes oft um die Bewegung irgendwelcher Teilchen geht. Beispiele: Was Du hörst, beruht auf der Bewegung von Luftteilchen. Elektrischer Strom ist meist Bewegung von Elektronen. Wenn Du etwas erwärmst, bewegen sich die Bestandteile davon heftiger.

1.2 Die beiden Seiten der Physik

Die Arbeit eines Physikers hat grob betrachtet zwei Seiten:

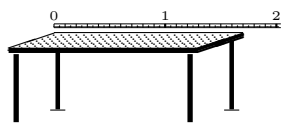


Hat man einen neuen Zusammenhang theoretisch hergeleitet, muss er in einem geeigneten Experiment überprüft werden. Macht man neue Beobachtungen, müssen eventuell bestehende Theorien angepasst oder neue entwickelt werden, um die Beobachtung zu erklären.

1.3 Größen allgemein

Unter „Beobachtungen“ sind dabei in der Regel Messwerte zu verstehen. In jedem Fall müssen also nachher für alle Messgrößen die nach der Theorie berechneten Werte mit den im Experiment gemessenen übereinstimmen. (im Rahmen der Messgenauigkeit)

Eine Größe ist jede Eigenschaft eines Körpers oder Vorgangs, für die man ein Messverfahren angeben kann.



Ein Beispiel ist die Länge ℓ eines Tisches. Man kann sie messen, indem man ein Maßband an seiner Kante entlang anlegt und abliest. Letzten Endes vergleicht man auf diese Weise den Tisch mit einem Meterstab.

Allgemein misst man eine Größe, indem man den interessierenden Körper oder Vorgang bezüglich dieser Eigenschaft mit einem Vergleichspartner vergleicht, den man für diesen Zweck vereinbart hat. Ergebnis des Messvorgangs ist die Angabe, wie viel mal mehr als am Vergleichspartner die Eigenschaft am vermessenen Körper oder Vorgang vorliegt. Typischerweise erfolgt diese Angabe in Form einer Gleichung wie $\ell = 1,70 \text{ m}$, was soviel bedeutet wie:

„Die gemessene Länge ist 1,70-mal so groß wie die Länge des Meterstabes.“

Solche Größenangaben setzen sich also aus einer **Maßzahl** und einer **Einheit** zusammen. Aus der Einheit geht die verwendete Vorschrift hervor, und aus der Maßzahl das mit dieser Vorschrift erhaltene Ergebnis. Als mathematischen Ausdruck bilden Maßzahl und Einheit ein Produkt.

Physikalische Theorien machen typischerweise Aussagen über derart zahlenmäßig erfassbare Größen und werden allermeist in Form mathematischer Gleichungen (seltener Ungleichungen) formuliert. Mathematik zu verwenden ist dabei nicht als zusätzliche Schikane gedacht, sondern erleichtert das Lösen physikalischer Probleme. Übersetzt man nämlich ein Problem in entsprechende Gleichungen, lässt sich zur Lösung alles anwenden, was die Mathematik an Techniken zur Lösung von Gleichungen zu bieten hat.

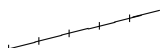
Maßzahl $\overbrace{1,70}^{\text{Maßzahl}} \text{ m} = 1,70 \cdot \underbrace{1 \text{ m}}_{\text{Einheit}}$ ↑ ↑ Einheit mal

Die Sprache der Physik ist die Mathematik.

1.4 Geometrische Größen

1.4.1 Länge

Unter einer **Länge** versteht man eine Abmessung entlang einer **Linie**. Je nach Zu-



sammenhang nennt man Längen auch Breiten, Höhen, Tiefen, Wegstrecken, Entfernungen, Dicken usw. Längen misst man in der Einheit **1 Meter**, die ursprünglich als Länge des „Urmeterstabes“ im «Bureau International des Poids et Mesures» in Paris festgelegt worden ist.

→ <http://www.bipm.fr/>

Zur Messung der Länge einer Linie stellt man fest, wie viele Meterstäbe man benötigen würde, um die Linie vollständig damit abzudecken.

Da man bei der Bewegung entlang einer Linie für die Bewegungsrichtung nur *eine* Wahlmöglichkeit hat – nämlich *vor* oder *zurück* – sagt man, eine Linie sei **eindimensional**.

Denkt man sich *ein* Maßband der Linie entlang ausgelegt, so lässt sich jeder Punkt der Linie durch den dort abzulesenden Wert auf dem Maßband beschreiben. Der Maßband-Wert wird dann als „Koordinate“ des Punktes bezeichnet.

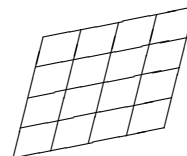
Als Koordinate könnte z.B. bei einer Kreislinie auch ein Winkel dienen: Auf dem Kreis, den die Spitze eines Minutenzeigers durchläuft, könnte die Position der „12“ als 0° festgelegt werden und als Koordinate der „3“ dann 90° angegeben werden. Bei einer Linie kommt man jedenfalls mit *einer* Koordinate aus.

1.4.2 Flächeninhalt

Mit dem **Flächeninhalt** gibt man die Größe einer **Fläche** an.

Man misst ihn in der Einheit **1 Quadratmeter**=1 m², indem man feststellt, wie viele 1 m²-„Kacheln“ man braucht, um die Fläche genau abzudecken.

Bei der Bewegung auf einer Fläche hat man zwei Wahlmöglichkeiten: die Wahl zwischen *vor* und *zurück*, aber auch die Wahl zwischen *nach links* oder *nach rechts*. Flächen sind **zweidimensional**.



Um Punkte einer Fläche zu beschreiben benötigt man jeweils *zwei* Koordinaten. Man könnte z.B. zwei sich kreuzende Maßbänder als x- und y-Achse auslegen. Als Koordinaten eines Punktes P könnte man dann auf jedem Band die dem Punkt P nächstgelegenen Stellen bestimmen und diese beiden Werte als Koordinatenpaar von P nehmen. Auf dem Ziffernblatt einer Uhr böte sich stattdessen neben dem Winkel als erster Koordinate als zweite Koordinate der Abstand des Punktes von der Mitte an. Weiteres Beispiel: Stellen der Erdoberfläche werden über Längen- und Breitengrade angegeben.

Bei der Umrechnung von Flächeneinheiten hat man zu beachten, dass wegen der zweiten Potenz z.B. gilt:

$$1 \text{ cm}^2 = 1 \cdot (10 \text{ mm})^2 = 10^2 \text{ mm}^2 = 100 \text{ mm}^2$$

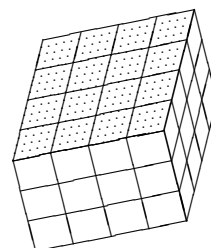
Weiterhin zu beachten: Für 100 m² sagt man auch 1 Ar= 1 a,
für 100 a auch 1 Hektar= 1 ha.

1.4.3 Volumen

Das **Volumen** heißt zu deutsch **Rauminhalt**.

Die Grundeinheit dafür ist **1 Kubikmeter**=1 m³. Das Volumen eines Raumes misst man, indem man feststellt, wie viele 1 m³-Würfel man bräuchte, um den Raum zu füllen.

Bei der Bewegung in einem Raum kommt eine dritte Wahlmöglichkeit hinzu: die Wahl zwischen *nach oben* und *nach unten*. Räume sind **dreidimensional**.



Zur Ortsangabe im Raum kommt man nicht mit weniger als *drei* Koordinaten aus. Es könnte in Fortführung der vorigen Anmerkungen ein drittes Maßband als z-Achse hinzugenommen werden. In der Geographie könnte zu Längen- und Breitengrad noch die Entfernung vom Erdmittelpunkt dazugenommen werden.

Bei der Umrechnung von Volumeneinheiten hat man zu beachten, dass wegen der dritten Potenz z.B. gilt:

$$1 \text{ cm}^3 = 1 \cdot (10 \text{ mm})^3 = 10^3 \text{ mm}^3 = 1000 \text{ mm}^3$$

Weiterhin zu beachten: Statt der Bezeichnung 1 dm^3
ist auch die Bezeichnung 1 Liter = 1 l gebräuchlich.

Berechnung von Volumina Bei einem Körper einfacher geometrischer Form kann man das Volumen aus den Abmessungen des Körpers berechnen. Mit Mathematikkenntnissen des 7. Schuljahres lassen sich immerhin Volumina von Körpern berechnen, die sich aus Quadern zusammensetzen. Zur Erinnerung: Ein Quader mit den Kantenlängen a, b, c hat das Volumen $V = a \cdot b \cdot c$.

Für anders, insbesondere unregelmäßig geformte Körper wie z.B. Gesteinsbrocken ist eine ähnlich einfache Berechnung nicht möglich.

weitere Formeln
für Volumina:

Kugel: $\frac{4}{3} \pi \cdot (\text{Radius})^3$

Prisma:

Grundfläche $G \cdot \text{Höhe } h$
Kegel: $\frac{1}{3} \cdot G \cdot h$

Volumen von Flüssigkeiten Bei Körpern ohne feste Form, wie z.B. einer Wassermenge, ist es hingegen wieder recht leicht, das Volumen zu bestimmen: Man füllt das Wasser in ein Gefäß, dessen Form man rechnerisch gut im Griff hat. Das Wasser passt sich dieser Form an und ist damit ebenfalls gut vermessbar.

Beispiel: Um das Volumen einer Wassermenge zu bestimmen, schüttet man sie in ein quaderförmiges Gefäß, dessen Grundfläche ein Quadrat mit Seitenlänge 4 cm ist. Das Wasser steht im Gefäß 12 cm hoch. Es hat also das Volumen eines Quaders mit Kantenlängen 4 cm, 4 cm und 12 cm, also $V = 4 \text{ cm} \cdot 4 \text{ cm} \cdot 12 \text{ cm} = 192 \text{ cm}^3$

An ein solches Gefäß könnte man leicht Markierungen anbringen, die anzeigen, welche Füllhöhen zu welchen Volumina gehören.

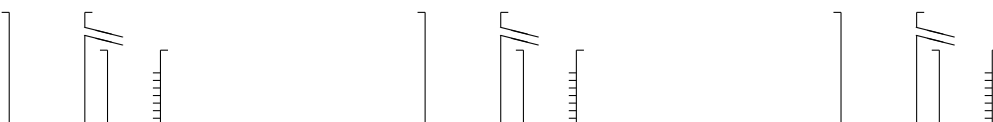
Man erhielte damit einen **Messbecher**.

Als nächstes könnte man sich sogar Messbecher beliebiger Form herstellen: Mit einem quaderförmigen Messbecher misst man eine bestimmte Flüssigkeitsmenge ab, schüttet diese in den neuen Becher und markiert an diesem die Füllhöhe dem eingefüllten Volumen entsprechend.

Verdrängung Zur Volumenbestimmung an unregelmäßig geformten Körpern nutzt man aus, dass sich Volumina von Flüssigkeiten leicht messen lassen: Man taucht den Körper in ein zunächst randvolles Gefäß. Es wird dann ebenso viel Volumen Flüssigkeit überlaufen, wie der Körper aufgrund seines eigenen Volumens Platz braucht. Man muss also nur messen, wie viel Flüssigkeit übergelaufen ist. Dazu genügt es, das übergelaufene Wasser in einem zunächst leeren Messbecher aufzufangen.

zu ergänzen:

- Körper
- Flüssigkeit



1.5 Zeit

Wie alle physikalischen Größen lässt sich auch die Zeit nur durch eine Messvorschrift erklären, sinngemäß wie in Albert Einsteins scheinbar scherzhafter Erklärung:

Die Zeit ist das, was man auf der Uhr abliest!

Je nach Zusammenhang nennt man eine Zeitspanne auch die Dauer eines Vorganges, ein Zeitintervall oder ein Alter. Die Grundeinheit für Zeiten ist die **Sekunde**, abgekürzt mit s. Die Dauer einer Sekunde kann man z.B. festlegen als Schwingungsdauer eines bestimmten (Uhr-) Pendels. Die offizielle Zeit liefert in Deutschland eine Atomuhr der Physikalisch-Technischen Bundesanstalt. → <http://www.ptb.de>

Allerdings stimmt wohl auch:
*Ein Mann mit einer Uhr
weiß, wie spät es ist. Ein
Mann mit zwei Uhren ist
sich nie sicher.*

1.6 Masse

1.6.1 Ohne Masse kein Gewicht

Die Masse eines afrikanischen Elefanten kann zum Beispiel 6300 kg betragen. Allerdings würden die meisten Leute eher vom *Gewicht* des Elefanten statt von seiner *Masse* sprechen. Warum benutzen die Physiker ein anderes Wort?

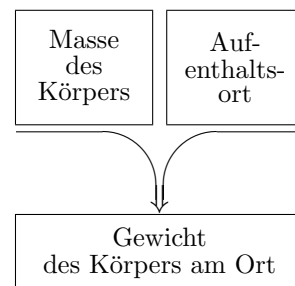
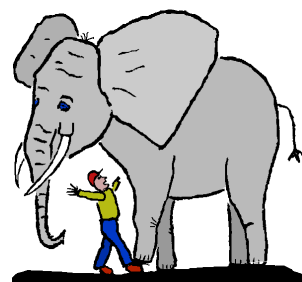
Natürlich weist der Elefant auch für einen Physiker ein Gewicht auf. Spätestens, wenn der Elefant dem Physiker auf den Fuß tritt, wird der Physiker dies bestätigen. Was genau der Physiker dann bestätigt, ist eine *Kraft*, die auf seinen Fuß ausgeübt wird, eine Kraft, die der Physiker „Gewichtskraft“ oder kurz „Gewicht“ nennt, – und deren Ausmaß vermutlich mit dem Elefanten zu tun hat.

Allerdings ist nicht der Elefant alleine die Ursache für diese Gewichtskraft. In einer weit abgelegenen Weltraumstation könnte nämlich derselbe Elefant denselben Physikerfuß gar nicht belasten, weil der Elefant dort ja in Schwerelosigkeit schweben würde. Das Gewicht kann also keine Eigenschaft des Elefanten alleine sein, denn bei diesem Experiment im All wäre zwar der Elefant da, aber kein Gewicht. Schon auf dem Mond würde der Elefant den Fuß deutlich weniger belasten als auf der Erde. Dort hat nämlich jeder Körper nur ein Sechstel so viel Gewicht wie auf der Erde. Offensichtlich entsteht ein Gewicht des Elefanten erst in der Nähe der Erde, des Mondes oder anderer Himmelskörper als Anziehungskraft zu diesem Himmelskörper hin.

Sicherlich hängt das Gewicht dann aber mit davon ab, wie der Elefant beschaffen ist. (wie groß, wie schlank, ...) Fügt man dem Elefanten irgendwelches Material (z.B. eine Decke oder Trinkwasser) hinzu, so wächst das Gesamtgewicht – und zwar an jedem Aufenthaltsort. Es ist daher sinnvoll, dem Elefanten an sich eine Eigenschaft zuzuordnen, aus der sich im Zusammenspiel mit dem Aufenthaltsort das Gewicht ergibt.

Ein Körper unterliegt der Anziehung von Himmelskörpern je nach seiner materiellen Zusammensetzung verschieden stark. Diese Eigenschaft eines Körpers, von Himmelskörpern angezogen zu werden, beschreibt man als die Masse m des Körpers. Die Masse eines Körpers beschreibt damit gleichzeitig, aus wie viel Materie der Körper besteht.

Wenn wir bisher das Gewicht für eine feste Eigenschaft eines Körpers gehalten haben, liegt das nur daran, dass wir in unserem Leben noch nicht so weit von der Erdoberfläche weggekommen sind. Nun wissen wir es besser: Unser Beispiel-Elefant hat als Verbindung bestimmter Mengen an Fleisch, Knochen und anderer Materie eine Masse von 6300 kg, egal, wo der Elefant sich gerade befindet. Aufgrund dieser Masse erfährt er auf der Erde eine gewisse Anziehungskraft, die man Gewicht nennt. Auf einem anderen Himmelskörper hätte er im Allgemeinen ein anderes Gewicht.



Als Einheit für Massen hat man das Gramm festgelegt. Man hat dazu einfach ein Stück Metall hergestellt und festgesetzt: *Dieses Stück Metall hat eine Masse von 1000 g, also von einem Kilogramm.* Dieses Metallstück nennt man deshalb das **Ur-Kilogramm**. Bei seiner Herstellung hat man sich an einem Liter Wasser orientiert. Wie das Ur-Meter wird auch das Ur-Kilogramm in Paris aufbewahrt. Um große und kleine Massen bequem angeben zu können, kann man – wie bei jeder Einheit – auch bei der Einheit Gramm bestimmte Vorsilben benutzen, um größere oder kleinere Einheiten zu bilden. Die folgende Tabelle gibt an, welche Vorsilbe eine Einheit um welchen Faktor verändert.

Beispiel: „Kilo-“ steht für „1000-“, also: $\text{kg} = 1000 \text{ g}$.

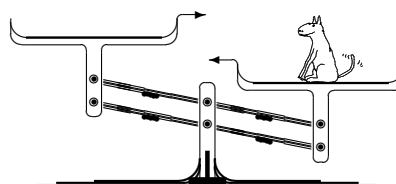
Faktor	Vorsilbe	Abkürzung
1 000 000 000 000 000 000	Exa-	E
1 000 000 000 000 000	Peta-	P
1 000 000 000 000	Tera-	T
1 000 000 000	Giga-	G
1 000 000	Mega-	M
1 000	Kilo-	k
100	Hekto-	h
10	Deka-	da
0,1	Dezi-	d
0,01	Zenti-	c
0,001	Milli-	m
0,000 001	Mikro-	μ
0,000 000 001	Nano-	n
0,000 000 000 001	Piko-	p
0,000 000 000 000 001	Femto-	f
0,000 000 000 000 000 001	Atto-	a

Neben der Einheit Gramm und ihren Vielfachen sind im Alltag und bei bestimmten Berufsgruppen weitere Masseneinheiten gebräuchlich:

$1 \text{ t} = \text{eine Tonne} = 1000 \text{ kg}$
 $1 \text{ Ztr} = \text{ein Zentner} = 50 \text{ kg}$
 $1 \text{ Pfd} = \text{ein Pfund} = 500 \text{ g}$
 $1 \text{ Kt} = \text{ein Karat} = 0,2 \text{ g}$

1.6.2 Balkenwaage als Messgerät für Massen

Eine **Balkenwaage** ist ähnlich wie eine Wippe aufgebaut. An jedem Ende des Waagebalkens ist eine Schale befestigt. Belastet man beide Schalen gleich stark, so stellt sich der Balken „waagrecht“ ein, ansonsten neigt er sich mit der stärker belasteten Seite nach unten.



Zum Wiegen legt man das Wägegut in die eine Schale. In die andere Schale legt man verschiedene Gewichtsstücke mit bekannten Massen. Man versucht dabei durch geeignete Auswahl an Gewichtsstücken, den Waagebalken wieder ins „Gleichgewicht“, d.h. in die waagerechte Stellung zu bringen.

Gelingt dies, so sind die Massen in beiden Schalen gleich groß. Die zu bestimmende Masse des Wägegutes in der einen Schale ist dann gleich der Summe der Massen aller Gewichtsstücke in der anderen Schale.

Projekt-Vorschlag:
Bau einer Balkenwaage
mit Wägesatz

Größen für Dreisatz-Künstler

2.1 Dichte

Durch Angabe einer Masse kann man ausdrücken, *wie viel* von einem Stoff man meint: Man kauft z.B. 100 g Leberwurst, verarbeitet 2 g Gold zu einem Schmuckstück oder erntet 4 Zentner Kartoffeln.

Zur Angabe, wie viel Stoff gemeint ist, benutzt man jedoch auch andere Größen. Statt von 200 g Speiseöl könnte man auch von 225 ml Öl sprechen. Man würde dann nicht die *Masse*, sondern das *Volumen*, zu deutsch: den *Rauminhalt* der Ölportion angeben. Für jemanden, der Ölfaschen entwirft, ist zunächst das Volumen interessanter, weil davon die Abmessungen der Flasche abhängen, und zwar unabhängig davon, wie viel das eingefüllte Öl wiegt. Für jemanden, der die Ölfasche nach Hause tragen will, spielt eher die Masse eine Rolle.

2.1.1 Definition von Dichte

Aufgaben wie diese:

1 dm³ Schaumstoff hat eine Masse von 3 g.
Welche Masse haben 2 dm³ Schaumstoff?

erledigst Du „mit links“: Natürlich haben 2 dm³ Schaumstoff zweimal so viel Masse wie 1 dm³, also $3 \text{ g} \cdot 2 = 6 \text{ g}$. Wenn Du von einem Stoff weißt, wie viel Masse 1 dm³ dieses Stoffes hat, brauchst Du nur diese Masse mit der Anzahl der interessierenden dm³ zu multiplizieren, um deren Masse zu berechnen. Es handelt sich hier um nichts weiter als eine Dreisatz-Aufgabe, eigentlich sogar nur um Zweisatz.

$$\left. \begin{array}{lcl} 1 \text{ dm}^3 & \mapsto & 3 \text{ g} \\ 2 \text{ dm}^3 & \mapsto & 6 \text{ g} \end{array} \right\} \cdot 2$$

Im Folgenden wirst Du wohl zunächst den Eindruck gewinnen, dass derselbe Zusammenhang unnötig komplizierter formuliert wird. Wozu tun wir Physiker das? Erstens natürlich, um damit Schüler besser quälen zu können. Zweitens aber auch, um den physikalischen Zusammenhang bei der Lösung auch schwierigster Probleme mit Hilfe der Mathematik nutzen zu können. Damit die Mathematik als Werkzeug richtig gut „greifen“ kann, müssen die Zusammenhänge zwischen Größen in Form von Gleichungen formuliert sein. Wir suchen nun also eine Gleichung, die den Zusammenhang zwischen dem Volumen V und der Masse m einer Stoffportion beschreibt. Wir verallgemeinern zunächst die bekannten Dreisatz-artigen Überlegungen mit Hilfe einer Hilfsvariablen x :

beteiligte Größen:	Volumen V	$\mapsto$	Masse m
bekannter Wert:	1 dm ³	$\mapsto$	3 g
Rechenbeispiel:	2 dm ³	$\mapsto$	$2 \cdot 3 \text{ g}$
allgemein:	$x \text{ dm}^3$	$\mapsto$	$x \cdot 3 \text{ g}$
als Gleichungen:	$V = x \text{ dm}^3$	$\mapsto$	$m = x \cdot 3 \text{ g}$

Die beiden Gleichungen für V und m bringen wir nun zusammen, indem wir die Gleichung für V nach x auflösen und den erhaltenen Term in die Gleichung für m einsetzen. Für das Auflösen nach x ist zu beachten, dass $x \text{ dm}^3$ als $x \cdot 1 \text{ dm}^3$ oder kürzer als $x \cdot \text{dm}^3$ angesehen werden kann. Um x alleine auf einer Seite der Gleichung zu erhalten, sind also beide Seiten der Gleichung durch dm^3 zu dividieren.

$$\begin{array}{llll} \text{die Gleichungen:} & V & = & x \cdot \text{dm}^3 \quad \text{und} \quad m = x \cdot 3 \text{ g} \\ V \text{ aufgelöst nach } x: & \frac{V}{\text{dm}^3} & = & x \\ x \text{ in } m \text{ ersetzt:} & & & m = \frac{V}{\text{dm}^3} \cdot 3 \text{ g} \\ \text{ein wenig umgestellt:} & & & m = 3 \frac{\text{g}}{\text{dm}^3} \cdot V \end{array}$$

$$m = 3 \frac{\text{g}}{\text{dm}^3} \cdot V$$

Und schon haben wir die gesuchte Formel für den Zusammenhang zwischen m und V gefunden. Störe Dich in der Herleitung nicht daran, dass die Division durch dm^3 nichts Anschauliches ist. Freilich hast Du die Division in der Grundschule als Rechenoperation kennengelernt, die dem Ver- oder Aufteilen entspricht: Man kann eine Anzahl Gummibärchen in 5 Haufen aufteilen, was der Division durch die Zahl 5 entspricht. Division durch die Einheit dm^3 kann aber sicher nicht mit einem „Aufteilen in dm^3 Haufen“ erklärt werden.

Um die Vorzüge¹ der Mathematik auch im Zusammenhang mit physikalischen Größen direkt nutzen zu können, überträgt man einfach Regeln der Mathematik aus der Welt der reinen Zahlen auf andere Dinge wie Größen und Einheiten. Es genügt, davon auszugehen, dass es auch für Größen und Einheiten eine (sonst nicht näher definierte) Multiplikation und Division nach den gängigen Regeln gibt. Konkret heißt dies, dass die Symbole für Größen und Einheiten wie gewöhnliche Variablen behandelt werden. Und so funktioniert's:

$$\begin{array}{ll} \text{unsere Formel:} & m = 3 \frac{\text{g}}{\text{dm}^3} \cdot V \\ \text{Anwendungsbeispiel: } V = 5 \text{ dm}^3 & \mapsto m = 3 \frac{\text{g}}{\text{dm}^3} \cdot 5 \text{ dm}^3 = 15 \text{ g} \end{array}$$

Wirklich etwas ausgerechnet wird natürlich immer nur mit den Maßzahlen, nicht mit anderen Dingen wie den Einheiten. Einheiten können aber je nach Stellung in einem Term z.B. weggelürzt oder zu Potenzen zusammengefasst werden.

Der Bruchstrich in der Einheit $\frac{\text{g}}{\text{dm}^3}$ wird meist als *pro* oder *je* gelesen, womit man deutlicher ausdrückt, dass *jeder* dm^3 jeweils 3 g wiegt.

Der Faktor $3 \frac{\text{g}}{\text{dm}^3}$ in unserer Formel ist typisch für den bestimmten Schaumstoff. Bei anderen Stoffen hat dieser Faktor i.A. andere Werte: Blei wiegt *mehr* Gramm pro Kubikdezimeter. Mehr Gramm pro dm^3 erhält man auch, wenn man den Schaumstoff zusammendrückt, also *verdichtet*. Dies erklärt auch, warum man diesen Faktor **Dichte** nennt. Die Dichte gibt an, wie dicht in dem betrachteten Stoff der Raum mit Masse gefüllt ist. Das übliche Formelsymbol für die Dichte ist der griechische Kleinbuchstabe ρ („rho“), gerne auch mit einem Index zur Angabe, von welchen Stoff die Dichte gemeint ist, z.B. ρ_{Blei} . Zur formalen Definition der Dichte genügt es im Wesentlichen, die Formel $m = \rho \cdot V$ nach ρ aufzulösen:

$$m = \rho \cdot V$$

zu ergänzen:

- Formelterm

Den Quotienten aus der Masse m und dem Volumen V einer Portion eines Stoffes nennt man die **Dichte** ρ („rho“) des Stoffes.

$$\rho :=$$

¹ja, solche gibt es:
Übersichtlichkeit, Eindeutigkeit, Effizienz, Widerspruchsfreiheit, Automatisierbarkeit.

Weiß man z.B. nach Messungen, dass ein $V = 2 \text{ dm}^3$ großes Stück Blei $m = 22600 \text{ g}$ wiegt, so lässt sich der Faktor für Blei so errechnen:

$$\rho = \frac{m}{V} = \frac{22600 \text{ g}}{2 \text{ dm}^3} = 11300 \frac{\text{g}}{\text{dm}^3}$$

Einheiten für die Dichte sind alle Quotienten aus einer Masseneinheit und einer Volumeneinheit. Sehr häufig wird $\frac{\text{g}}{\text{cm}^3}$ benutzt. Gleichbedeutend – weil einfach mit 1000 erweitert – ist $\frac{\text{kg}}{\text{dm}^3}$.

$$\frac{\text{kg}}{\text{dm}^3} = \frac{\text{kg}}{(\text{dm})^3} = \frac{1000 \text{ g}}{(10 \text{ cm})^3} = \frac{1000 \text{ g}}{10^3 \text{ cm}^3} = \frac{1000 \text{ g}}{1000 \text{ cm}^3} = \frac{\text{g}}{\text{cm}^3}$$

Jeder Stoff hat seine, für ihn im Normalzustand typische Dichte. Man bezeichnet die Dichte daher als **stoffspezifische** Größe. „spezifisch“ kommt von lateinisch *species*=Art und bedeutet daher etwa „art-typisch“.

Mit Hilfe der Formel $\rho = \frac{m}{V}$ kannst Du von den drei Größen ρ , m und V jede berechnen, wenn Dir die beiden anderen bekannt sind. Mit mathematischen Umformungen² formst Du auch physikalische Formeln so um, dass die gesuchte Größe nur alleine auf einer Seite der Gleichung steht. Behandle dazu die Symbole für Größen und Einheiten wie Variablen!

- Wenn Du m bestimmen willst:

$$\begin{array}{l|l} \rho = \frac{m}{V} & \cdot V \\ \rho \cdot V = \frac{m \cdot V}{V} & \text{kürzen} \\ \rho \cdot V = \frac{m}{1} & \text{Seitentausch, } \frac{m}{1} = m \\ m = \rho \cdot V & \text{das war's schon!} \end{array}$$

- Für V geht es z.B. so:

$$\begin{array}{l|l} \rho = \frac{m}{V} & \cdot V \\ \rho \cdot V = m & : \rho \\ V = \frac{m}{\rho} & \text{et voilà!} \end{array}$$

Nach dem Umformen musst Du nur noch die bekannten Größen auf der rechten Seite der Gleichung einsetzen, die Maßzahlen verrechnen und eventuell noch die Einheit vereinfachen. (Doppelbrüche beseitigen, kürzen)

2.1.2 Bestimmung einiger Dichten

Zur Bestimmung der Dichte eines Stoffes gemäß $\rho := \frac{m}{V}$ muss man von einer Portion des Stoffes die Masse m und das Volumen V bestimmen. Um die Dichte eines Stoffes möglichst genau zu bestimmen, kann man mehrere Portionen vermessen und den Mittelwert der Quotienten $\frac{m}{V}$ bestimmen. Messfehler bei einer Portion beeinflussen dann das Gesamt-Ergebnis nicht so sehr, als wenn nur diese eine Portion untersucht worden wäre. Fehler bei mehreren Portionen können sich sogar bei der Mittelung gegenseitig aufheben.

Fehler, durch die der wahre Wert immer in derselben Richtung verfälscht gemessen wird (immer zu groß oder immer zu klein), nennt man **systematische** Fehler. Ein systematischer Fehler wird durch Mittelung nicht verringert: Sind z.B. alle Einzelergebnisse um 2 Einheiten zu groß, ist auch der Mittelwert um diese 2 Einheiten zu groß. Fehler, die gleichermaßen Verfälschungen in beide Richtungen verursachen, nennt man **zufällige** Fehler. Solche Fehler lassen sich grundsätzlich beliebig verringern, indem man über ausreichend viele Einzelmessungen mittelt.

Flüssigkeiten

Wie oben dargelegt, lässt sich mit Hilfe eines Messbechers leicht das Volumen einer Flüssigkeitsmenge messen oder ein bestimmtes Volumen der Flüssigkeit abmessen. Bei der Bestimmung der Masse taucht das kleine Problem auf, dass man eine Flüssigkeit schlecht ohne Gefäß abwiegen kann. Wir müssen also beachten, dass wir nicht die Flüssigkeit alleine wiegen, sondern die Flüssigkeit plus Gefäß.

Manche Waagen kann man zu Beginn der Messungen so einstellen, dass beim Wiegen des leeren Gefäßes eine Null angezeigt wird. Bei gefüllten Gefäßen wird dann nur die Masse der Füllung angezeigt. Ansonsten muss man zunächst das leere Gefäß wiegen und bei späteren Wägungen die notierte Leer-Masse vom Messwert abziehen.

²auf beiden Seiten denselben Ausdruck addieren, beide Seiten mit demselben Ausdruck $\neq 0$ multiplizieren, von beiden Seiten $\neq 0$ den Kehrwert bilden, ...

V 2.1 Dichte verschiedener Flüssigkeiten

Durchführung: In einem Messbecher werden von verschiedenen Flüssigkeiten verschiedene Volumina abgemessen und gewogen.

Messwerte: /Auswertung: (mehrere Spalten pro Stoff vorgesehen)

Stoff											
$V[\text{cm}^3]$											
$m[\text{g}]$											
$\rho[\frac{\text{g}}{\text{cm}^3}]$											
$\bar{\rho}[\frac{\text{g}}{\text{cm}^3}]$											

Feste Stoffe

V 2.2 Dichte bei würfelförmigen Materialproben

Durchführung: Würfel aus unterschiedlichen Materialien werden gewogen.

Messwerte: /Auswertung:

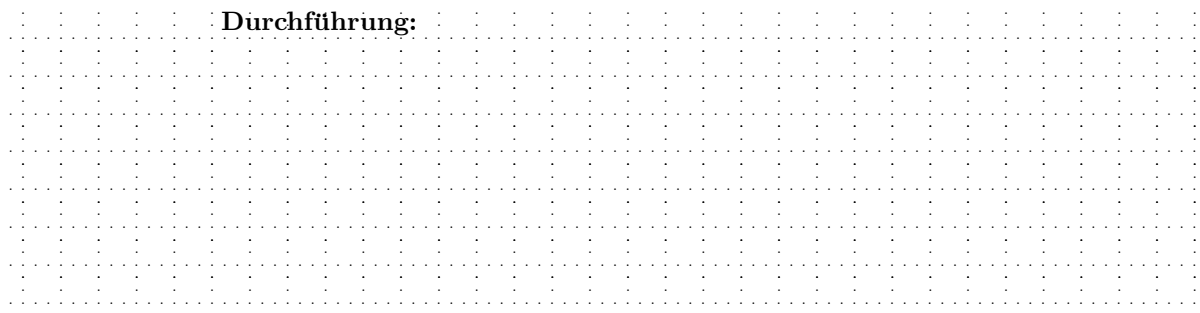
Die Würfel haben eine Kantenlänge von

und damit ein Volumen von

Material						

V 2.3 Dichte bei unregelmäßig geformten Körpern

Durchführung:



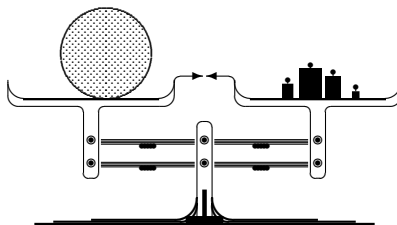
Messwerte: /Auswertung:

Gase

Außer in festem und in flüssigem Zustand können Stoffe auch gasförmig sein. Ein weit verbreitetes Beispiel für einen gasförmigen Stoff ist Luft. Vielleicht fragst Du Dich: Hat Luft überhaupt ein Gewicht, eine Masse, eine Dichte? Zur Beantwortung dieser Frage müssen wir ein bestimmtes Volumen an Luft wiegen. Luft lässt sich jedoch nicht so einfach auf eine Waage legen. Wie Flüssigkeiten können wir Luft nur zusammen mit einem Gefäß wiegen.

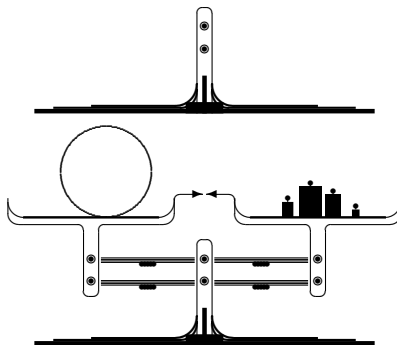
V 2.4 Dichte von Luft

Durchführung: Auf die eine Seite einer Balkenwaage wird ein offener – und daher luftgefüllter – Glaskolben gelegt. Mit Gewichtsstücken auf der anderen Seite wird die Waage ins Gleichgewicht gebracht. (→ Bild) Dann wird der Glaskolben von der Waage genommen, luftleer gepumpt („evakuiert“), verschlossen und auf die Waage zurückgelegt.



Beobachtung:

Die Waage stellt sich so ein:



Durchführung: Es werden zum Glaskolben solche – noch einzuzzeichnende! – Gewichtsstücke dazugelegt, dass die Waage wieder ins Gleichgewicht kommt. Anschließend wird noch festgestellt, wie viel Wasser in den Glaskolben passt.

zu ergänzen:

- Skizze der beobachteten Waage
- Gewichtsstücke
- Werte

Messwerte: /Auswertung: /Ergebnis:

nachgelegte Gewichtsstücke: $m =$

eingefüllte Wassermenge: $V =$

$\Rightarrow$ Dichte der Luft: $\rho = \frac{m}{V} =$

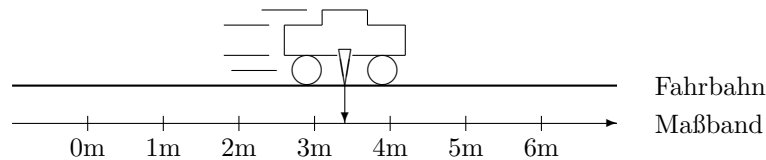
Tabellen weiterer Stoff-Dichten

Stoff	$\rho[\frac{\text{g}}{\text{cm}^3}]$	Stoff	$\rho[\frac{\text{g}}{\text{cm}^3}]$	Stoff	$\rho[\frac{\text{g}}{\text{cm}^3}]$	Stoff	$\rho[\frac{\text{g}}{\text{cm}^3}]$
Aluminium	2,70	Wasser	1,00	Stahlbeton	2,2-2,5	Fichtenholz	0,5
Blei	11,34	Wassereis	0,92	Granit	um 2,8	Eichenholz	0,9
Eisen	7,86	Quecksilber	13,6	Glas	um 2,5	Knochen	1,7-2,0
Gold	19,32	Alkohol	0,79	Baumwolle	1,47-1,5	Papier	0,7-1,2
Kupfer	8,96	Honig	1,45	Kork	um 0,15	Porzellan	2,2-2,5
Silber	10,50	Olivenöl	0,92	Styropor	0,015	Wachs	0,94-1,04
Zinn	7,31	Benzin	0,7	Luft	0,0013	Sand	1,5
Silizium	2,33	Diesel	0,85	Braunkohle	1,2-1,5	Mehl	0,40-0,55

2.2 Geschwindigkeit

2.2.1 Bezugssystem

Wenn ein Körper, z.B. ein Wagen, sich fortbewegt, so ändert er seinen Aufenthaltsort. Um diese Orte beschreiben zu können, denken wir uns der Bahn des Körpers entlang ein Maßband ausgelegt. Wir können dann Positionen am Maßband ablesen und z.B. sagen: „Der Wagen befindet sich an der **Position** 3,4 m.“ Statt „Position“ können wir natürlich auch **Stelle**, **Punkt** oder **Ort** sagen.



In Gleichungen werden wir für „Stelle“ den Buchstaben s schreiben. Das Maßband wird damit zu einem Ausschnitt aus einer „ s -Achse“. Für das letzte Beispiel würden wir also schreiben: $s = 3,4 \text{ m}$

Wenn sich der Wagen z.B. von der Stelle $s_1 = 3,4 \text{ m}$ zu der Stelle $s_2 = 5,7 \text{ m}$ bewegt, so *ändert* er seine Position um 2,3 m. Er legt eine **Strecke** von 2,3 m zurück. Statt „Strecke“ ist auch der Begriff **Weg** üblich.

Die Länge der zurückgelegten Strecke ergibt sich natürlich als Differenz der End- zur Startposition. Solche Änderungs-Differenzen bezeichnet man mit Hilfe des großen griechischen Buchstabens Δ (Delta). Die Änderung der Position s wird Δs abgekürzt.

$$\Delta s = s_2 - s_1$$

Halte Stellen und Strecken auseinander! Auch, wenn wir eine Stelle mit der Länge der Strecke zwischen ihr und dem Nullpunkt bezeichnen. Und auch, wenn manchmal Wegstrecken mit s statt Δs bezeichnet werden. Dies kommt schon mal vor, wenn aus dem Zusammenhang klar wird, dass ein Weg und keine Position gemeint ist.

Bewegungen dauern immer eine gewisse **Zeit**. Das Formelsymbol für die Zeit ist t von englisch *time*, lateinisch *tempus* oder französisch *temps* oder oder Auch bei der Zeit müssen wir unterscheiden zwischen **Zeitpunkten** wie z.B. 4 Uhr und **Zeitspannen** wie z.B. 4 Stunden. Zur Angabe von Zeitpunkten legt man eine Zeitskala mit einem Zeit-Nullpunkt willkürlich fest. Beispiele:

- Den Beginn eines Versuches legt man als Nullpunkt fest und Zeitpunkte werden als „sowasviel Sekunden vor oder nach diesem Nullpunkt“ angegeben.
- Die „Mitte“ der Nacht legt man als Nullpunkt für Zeitangaben im Alltag fest und meint dann mit „13 Uhr“ 13 Stunden nach Mitternacht.

Zeitspannen liegen zwischen Zeitpunkten.

Zeitspannen heißen auch **Zeitintervalle** oder **Zeiträume**.

Zwischen einem Zeitpunkt t_1 und einem Zeitpunkt t_2 liegt die Zeitspanne Δt .

$$\Delta t = t_2 - t_1$$

2.2.2 Gleichförmige Bewegung

Aufgaben wie diese hier kennst Du wohl bereits aus dem Mathematik-Unterricht. Eine solche Aufgabe macht nur Sinn, wenn man unterstellt, dass das Schiff in jeder Stunde gleich viele Kilometer zurücklegt.

*Ein Schiff legt in 2 Stunden 90 km zurück.
Wie weit fährt es in 5 Stunden?*

Eine Bewegung, bei der in gleich langen Zeitspannen gleich lange Wege zurückgelegt werden, nennen wir **gleichförmig**.

Theoretische Untersuchung

Charakteristisch für eine gleichförmige Bewegung ist also die Angabe, wie viel Weg in wie viel Zeit zurückgelegt wird. Gibt man z.B. an, dass in einer Stunde ein Weg von 45 km zurückgelegt wird, bedeutet dies zunächst, dass in *jeder* Stunde 45 km zurückgelegt werden. Man sagt dann auch, es werden 45 km *pro* Stunde zurückgelegt.

In 2 Stunden werden natürlich 2 mal 45 km (also 90 km) zurückgelegt, in 3 Stunden 3 mal 45 km, in einer halben Stunde $\frac{1}{2}$ mal 45 km usw. Auch hier können wir uns aus dem Dreisatz-Schema heraus eine richtige Formel zur Berechnung des zurückgelegten Weges herleiten.

beteiligte Größen:	Zeit Δt	$\mapsto$	Weg Δs
bekannter Wert:	1 h	$\mapsto$	45 km
Rechenbeispiel:	5 h	$\mapsto$	5 · 45 km
allgemein:	x h	$\mapsto$	$x \cdot 45$ km
als Gleichungen:	$\Delta t = x$ h	$\mapsto$	$\Delta s = x \cdot 45$ km
Δt aufgelöst nach x :	$\frac{\Delta t}{\text{h}} = x$		
x in Δs ersetzt:			$\Delta s = \frac{\Delta t}{\text{h}} \cdot 45$ km
ein wenig umgestellt:			$\Delta s = 45 \frac{\text{km}}{\text{h}} \cdot \Delta t$

Der in der Herleitung vorkommende Term $\frac{\Delta t}{\text{h}}$ entspricht der verbalen Anleitung: „Nimm von der Zeitangabe in der Einheit Stunden nur die Maßzahl ...“ Mit der so erhaltenen Anzahl x an Stunden weiß man dann ja, wie viel mal 45 km gefahren werden.

In dem zusammengefassten Term $45 \frac{\text{km}}{\text{h}}$ erkennst Du wohl sofort eine Geschwindigkeitsangabe. Der Zahlenwert ist nämlich größer, wenn man schneller fährt. Als neue Größe bezeichnen wir die Geschwindigkeit mit dem Symbol v , vom lateinischen Wort *velocitas* für Geschwindigkeit. Demnach haben wir die nebenstehende Formel gefunden, aus der sich folgende Definition ergibt:

$$\Delta s = v \cdot \Delta t$$

Legt ein Körper in der Zeitspanne Δt den Weg Δs zurück, so heißt der Bruch $\frac{\Delta s}{\Delta t}$ die Geschwindigkeit v des Körpers. Kurzfassung:

$$v :=$$

Nach dieser Definition ergibt sich für v immer derselbe Wert, wenn man Zeiten und Wege derselben gleichförmigen Bewegung einsetzt:

$$v = \frac{45 \text{ km}}{1 \text{ h}} = \frac{90 \text{ km}}{2 \text{ h}} = \frac{135 \text{ km}}{3 \text{ h}} = \dots = \frac{22,5 \text{ km}}{\frac{1}{2} \text{ h}} = \dots = \frac{n \cdot 45 \text{ km}}{n \cdot 1 \text{ h}}$$

mit beliebigem Wert für den Faktor n (außer 0): Er kürzt sich eh weg!

zu ergänzen:

- Formelterm

Experimentelle Untersuchung

Für folgende Untersuchung der Bewegung eines Körpers machen wir uns das Leben etwas leichter: Wir setzen den Beginn der Messung als Zeit-Nullpunkt fest und die dann eingenommene Position des Körpers als Nullpunkt für Positionsangaben. Wenn der Körper dann z.B. zum Zeitpunkt $t = 3 \text{ s}$ an der Position $s = 15 \text{ cm}$ ist, hat er seit dem Start der Messung den Weg $\Delta s = s = 15 \text{ cm}$ in der Zeitspanne $\Delta t = t = 3 \text{ s}$ zurückgelegt. Wir können dann also die Werte für s und t direkt als Δs bzw. Δt nehmen.

Aufbau: /Durchführung:

Messwerte: /Auswertung:

Ergebnis:

Graphische Darstellung

Im vorigen Versuch haben wir Messwert-Paare aus jeweils einem Zeitpunkt t und der dann eingenommenen Position s erhalten. Die beiden Maßzahlen eines solchen Paares kann man als Koordinaten eines Punktes in einem Koordinatensystem ansehen. Zu jedem Messwertpaar gehört dann ein Punkt, auch Messpunkt genannt.

Die Koordinatenachsen beschriftet man mit den Symbolen für die entsprechenden Größen, hier also mit t und s . Die für eine Größe verwendete Einheit schreibt man entweder in eckigen Klammern dahinter, z.B. $s[\text{cm}]$, oder als Nenner unter das Größensymbol und einen Bruchstrich.

So ist gewährleistet, dass einerseits die Einheit nicht ganz aus der Darstellung verschwindet, andererseits aber nur die Maßzahlen zu Koordinaten werden. $s[\text{cm}]$ soll verstanden werden als s , *gezählt in* cm; der Bruch $\frac{s}{\text{cm}}$ stellt s *ohne die „wegdividierte“ Einheit* cm dar.

Ein so angelegtes Schaubild nennt man ein t-s-Diagramm.



Es fällt in unserem Diagramm auf, dass die Messpunkte (etwa)

.....

Aus welchen Gründen liegen unsere Messwerte nicht exakt auf einer Geraden?

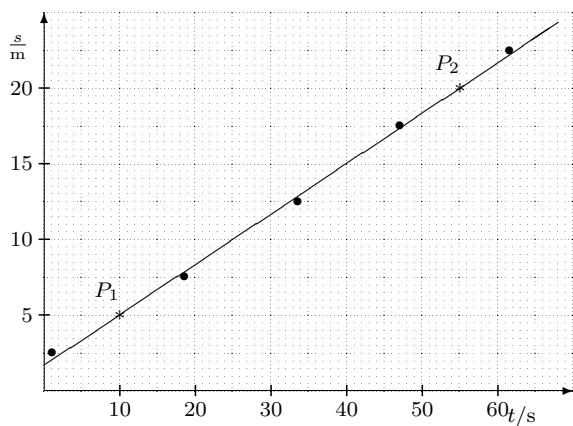
Wie Du vielleicht schon aus der Mathematik weißt, ist ein wesentliches Merkmal einer Geraden ihre Verlaufs-Richtung, ihre **Steigung**. Die Steigung einer Geraden lässt sich über ein Steigungsdreieck ermitteln. Man wählt dazu zunächst zwei nicht allzu dicht beieinander liegende Punkte der Geraden aus, im Folgenden von links nach rechts P_1 und P_2 genannt. Diese Punkte könnten Messpunkte sein. Die Koordinaten von P_1 sind (bis auf die Einheiten) $(t_1|s_1)$, bei P_2 entsprechend $(t_2|s_2)$.

Man zeichnet durch die beiden Punkte Parallelen zu den Achsen, so dass ein rechtwinkliges Dreieck entsteht. Die Steigung der Geraden ist definiert als das Längen-Verhältnis der vertikalen zur horizontalen Dreiecksseite. In unserem Falle entspricht die Länge der vertikalen Seite jedoch gerade $s_2 - s_1$, also einer Wegstrecke Δs , und die Länge der horizontalen Seite beschreibt $t_2 - t_1$, also ein Zeitintervall Δt .

Die Steigung beschreibt damit $\frac{\Delta s}{\Delta t}$, die Geschwindigkeit v des Körpers!

In einem t-s-Diagramm ist die Steigung (bei Berücksichtigung der Einheiten) die Geschwindigkeit der dargestellten Bewegung.

Verwechsle nicht, was mit „Steigung“ hier gemeint ist: Es geht hier nicht darum, dass der untersuchte Körper sich auf einer ansteigenden Bahn aufwärts bewegt. Sondern: Die Bewegung des Körpers wird durch eine Gerade in einem t-s-Diagramm beschrieben, deren Verlauf mit dem entsprechenden Verhältnis-Rechenausdruck $\frac{\Delta s}{\Delta t}$ namens „Steigung“ beschrieben wird.



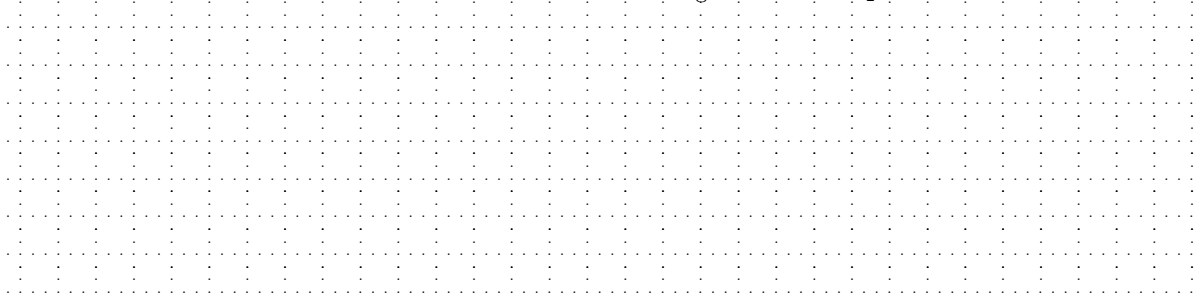
Ausführliche Bestimmung einer Steigung:

Entlang eines 25 m langen Gleises einer Modelleisenbahn liegt ein Maßband aus, auf dem man die Position s einer Lokomotive auf diesem Gleis als Maßangabe sofort ablesen kann. Auf einer Stoppuhr wurden bei einer Fahrt der Lokomotive die Zeitpunkte t abgelesen, zu denen die Lok bestimmte Positionen erreichte. Das nebenstehende t-s-Diagramm zeigt die so erhaltenen Messwertpaare als Punkte (•). Da man eine Gerade dazuzuichnen kann, von der diese Punkte kaum abweichen, betrachten wir diese Gerade als Beschreibung der vollständigen Fahrt. Die einzelnen Messpunkte interessieren nicht mehr.

1. Ergänze die folgende Aussage, die der Punkt $P_1(t_1|s_1)$ macht! Setze dazu unter anderem die Worte *Lok*, *Maßband*, *Position*, *Stoppuhr*, *Zeitpunkt* an der richtigen Stelle ein!

Zum, als die (= t_1) anzeigte, befand sich
 die an der, die auf
 dem als (= s_1) abgelesen wurde.

2. Formuliere nach diesem Muster die Aussage des Punktes P_2 !



3. Wie viel Zeit liegt zwischen den durch P_1 und P_2 beschriebenen Situationen? Ergänze die Berechnung:

$$\Delta t = t_2 - t_1 = \dots \text{ s} - \dots \text{ s} = \dots$$

Stelle diese Zeitspanne im Diagramm durch geeignete Strecken in einer Farbe Deiner Wahl dar!

4. Wie viel Weg legt die Lokomotive in dieser Zeitspanne zurück?
 Berechne wie oben:

$$\dots = \dots = \dots = \dots$$

Stelle diesen Weg im Diagramm durch geeignete Strecken in einer anderen Farbe dar!

5. Berechne die Geschwindigkeit v der Lokomotive!

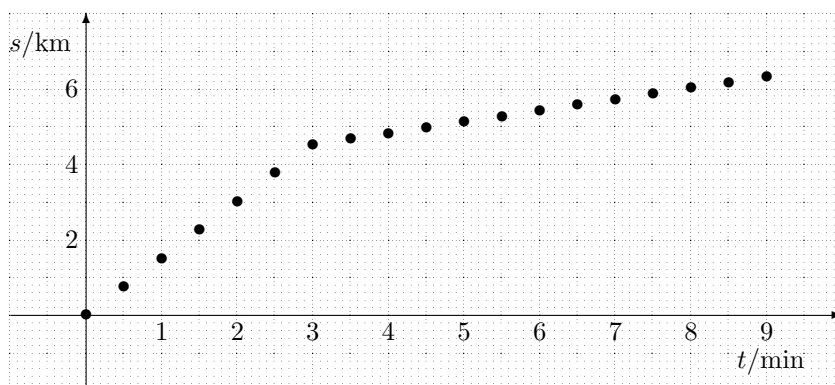
$$v = \frac{\Delta s}{\Delta t} = \dots = \dots$$

2.2.3 Ungleichförmige Bewegung

Mittlere Geschwindigkeit

Dieses Diagramm zeigt Messwerte der Bewegung eines Radrennfahrers, der sich für einen hervorragenden Start ziemlich verausgabt und anschließend nur noch im Schongang fahren kann. Der Radfahrer legt vom Start an in 8 Minuten 6 Kilometer zurück. Bei gleichförmiger Bewegung hätte er dabei

eine Geschwindigkeit von $\bar{v} =$



Tatsächlich fährt er aber anfangs schneller (nämlich $1,5 \frac{\text{km}}{\text{min}}$) und später langsamer (nämlich $0,3 \frac{\text{km}}{\text{min}}$) als mit dieser Geschwindigkeit $\bar{v}$. Man nennt $\bar{v}$ deshalb nur die **mittlere Geschwindigkeit** des Radfahrers in den ersten 8 Minuten.

Legt ein Körper in der Zeitspanne Δt den Weg Δs zurück, so hat er in dieser Zeitspanne die mittlere Geschwindigkeit $\bar{v} = \frac{\Delta s}{\Delta t}$.

Momentangeschwindigkeit

Die Angabe einer mittleren Geschwindigkeit bezieht sich immer auf eine *Zeitspanne*. Überhaupt können wir bisher Geschwindigkeiten immer nur auf eine Zeitspanne beziehen. Den Begriff „Geschwindigkeit“ haben wir ja mit Hilfe eines Δt definiert. Dennoch halten wir es auch für sinnvoll, von einer Geschwindigkeit zu einem bestimmten *Zeitpunkt* zu sprechen: Zu jedem Zeitpunkt (=in jedem Moment) hat ein Körper eine bestimmte Geschwindigkeit, die sogenannte **Momentangeschwindigkeit**.

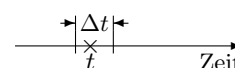
Im Auto oder auf dem Fahrrad ist die Momentangeschwindigkeit zu einem Zeitpunkt diejenige Geschwindigkeit, die der Tachometer (griechisch: *tachys*=schnell, *metron*=Maß) zu diesem Zeitpunkt anzeigt. Wie Du weißt, hat diese Momentangeschwindigkeit auch eine größere physikalische Bedeutung als irgendwelche mittleren Geschwindigkeiten: Wenn ein Auto gegen einen Baum prallt, ist für das Ausmaß des Schadens die Momentangeschwindigkeit beim Aufprall entscheidend,

und nicht z.B. die mittlere Geschwindigkeit in der letzten halben Stunde vor dem Aufprall.

Um mit Momentangeschwindigkeiten in der Physik umgehen zu können, benötigen wir nun noch eine richtige Definition. Dabei haben wir jedoch ein Problem: Wenn wir die Bewegung eines Körpers untersuchen wollen, muss sich der Körper natürlich bewegen. Für eine Bewegung braucht er allerdings immer eine *Zeitspanne*. Wir wollen jedoch die Geschwindigkeit zu einem *Zeitpunkt* ermitteln. In einem einzelnen Zeitpunkt bewegt sich der Körper aber überhaupt kein Stück weit!

Um trotzdem eine brauchbare Definition aufzustellen nutzen wir aus, dass sich Momentangeschwindigkeiten nicht abrupt ändern können: Ein Fahrrad, das mit $10 \frac{\text{km}}{\text{h}}$ fährt, kann nicht plötzlich $15 \frac{\text{km}}{\text{h}}$ schnell fahren, ohne in einer Zwischenzeit nacheinander alle Geschwindigkeiten zwischen 10 und $15 \frac{\text{km}}{\text{h}}$ gehabt zu haben.

Wir können demnach davon ausgehen, dass die Momentangeschwindigkeit zu einem Zeitpunkt t ungefähr gleich ist der mittleren Geschwindigkeit in einer sehr kurzen Zeitspanne um den Zeitpunkt t herum. Je kürzer die Zeitspanne, desto besser. Ob die kleine Zeitspanne bis zum Zeitpunkt t , ab t oder von kurz vor t bis kurz nach t liegt, spielt dabei eigentlich keine Rolle, wenn die Spanne nur kurz genug ist.



Die Momentangeschwindigkeit v eines Körpers zu einem Zeitpunkt t ist die mittlere Geschwindigkeit des Körpers in einer winzigst kleinen Zeitspanne um den Zeitpunkt t herum.

Ausblick auf kommende Schuljahre:

Die Änderung der Momentangeschwindigkeit eines Körpers erfordert stets, dass eine Zeit lang von außerhalb des Körpers eine Kraft auf ihn ausgeübt wird. („Trägheitssatz“)

Um bei einem Körper der Masse m eine gewisse Beschleunigung a zu bewirken ist eine Kraft F mit $F = m \cdot a$ erforderlich.

Dabei ist a definiert als Geschwindigkeitsänderung pro Zeit, also $a := \frac{\Delta v}{\Delta t}$.

2.3 Rechnen mit Größen

2.3.1 Multiplikation

Die Multiplikation hast Du in der Grundschule als Abkürzung für wiederholte Addition desselben Summanden kennengelernt. Diese Sichtweise macht natürlich für einen Ausdruck wie $3 \frac{\text{g}}{\text{dm}^3} \cdot 2 \text{ dm}^3$ überhaupt keinen Sinn: Man kann nicht „ $3 \frac{\text{g}}{\text{dm}^3}$ -fach“ das Volumen 2 dm^3 addieren. Eine Summe von Volumina wäre außerdem selbst wieder ein Volumen, niemals aber eine Masse.

Es ist nützlich, die Multiplikation auch als eigenständige Verknüpfung von Größen zu sehen, losgelöst von der bloßen Abkürzung für Summen. Die Multiplikation zweier Größenangaben besteht aus der

gewöhnlichen Multiplikation der Maßzahlen und dem Zusammenfassen der Einheiten gemäß den Regeln zum Multiplizieren von (Bruch-)Termen aus Variablen. So können Größen beliebiger Art sinnvoll miteinander verknüpft werden, wobei mit einer solchen Verknüpfung i.A. eine Größe neuer Art entsteht. Summen aus verschiedenartigen Größen sind niemals sinnvoll.

Multiplikation als Abkürzung für Summation ist so gesehen nur der eine Sonderfall, dass mindestens einer der beiden Faktoren den Charakter einer „Anzahl“ hat. (mit der Einheit „Stück“, dem „neutralen Element“ der Multiplikation von Einheiten)

2.3.2 Gleichheit

Das **Gleichheitszeichen** ist *nicht* – wie es in der Grundschule zuweilen scheinen mag – eine Art Doppelpunkt der Mathematik, mit dem man hinter einer Aufgabe ihre Lösung anschließt.

In $3 \cdot 5 = 15$ ist $3 \cdot 5$ auch keine „Aufgabe“, sondern nur eine von vielen Möglichkeiten, die Zahl 15 auszudrücken. (Ebenso wie $10+5$, *quinze* oder **XV**) Eine Gleichung besagt immer nur, dass die Ausdrücke links und rechts vom $=$ das gleiche bedeuten; hier im Beispiel, dass überall, wo $(3 \cdot 5)$ steht, ebensogut (15) stehen könnte und umgekehrt.

Eine Anwendung: Bisher haben wir Geschwindigkeiten manchmal in $\frac{\text{km}}{\text{h}}$, manchmal in $\frac{\text{m}}{\text{s}}$, manchmal in $\frac{\text{cm}}{\text{s}}$ angegeben. Die Einheit für eine Geschwindigkeitsangabe hat sich in allen Fällen daraus ergeben, in welcher Einheit Weglängen und in welcher Einheit Zeiten gemessen worden sind. Es ergibt sich immer eine brauchbare Einheit für Geschwindigkeiten, wenn man irgendeine Einheit für Längen durch irgendeine Einheit für Zeiten dividiert. Natürlich verwenden wir nur äußerst selten eine Einheit wie Seemeilen pro Schulstunde, aber korrekt wäre auch diese.

Zum Vergleichen und zum Verrechnen von Geschwindigkeiten muss man manchmal die Einheit wechseln, Angaben von einer Einheit in eine andere umrechnen.

Beispiel:

Rechne $90 \frac{\text{km}}{\text{h}}$ um in $\frac{\text{m}}{\text{s}}$!

1. Methode: Überlegungen wie bei Dreisatz-Aufgaben:

in 1 h : 90 km

d.h in 3600 s : 90000 m

in 1 s : 90000 m / 3600

also in 1 s : 25 m

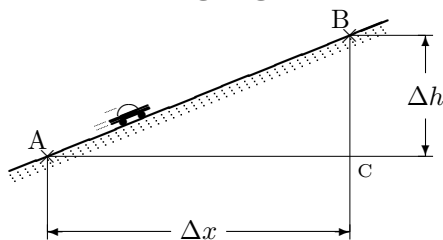
2. Methode: sture Benutzung der Mathematik: Ersetze Terme durch gleichwertige und vereinfache! Du darfst hemmungslos alles, was Du in Mathematik gelernt hast, in der Physik anwenden! Dafür ist Mathematik ja da!

Da $1 \text{ km} = 1000 \text{ m}$ und $1 \text{ h} = 3600 \text{ s}$, gilt:

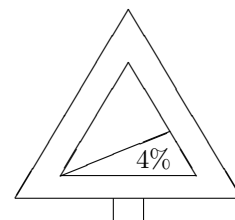
$$90 \frac{\text{km}}{\text{h}} = 90 \cdot \frac{1000 \text{ m}}{3600 \text{ s}} = \frac{90 \cdot 1000 \text{ m}}{3600 \text{ s}} = 25 \frac{\text{m}}{\text{s}}$$

Der Vorteil der ersten Methode ist, dass sie in Fällen wie diesem Beispiel evtl. leichter erscheint und anschaulicher ist. Die zweite Methode erfordert dafür überhaupt keine Anschauung und funktioniert deshalb auch für Umrechnungen, bei denen man sich die Zusammenhänge nicht so gut vorstellen kann.

2.3.3 Steigungen



Zur Steigung einer Straße:
Ein Auto fährt auf einer ansteigenden Straße. Zwischen den Punkten A und B, die auf einer Landkarte um Δx voneinander entfernt eingezeichnet wären, gewinnt das Auto um Δh an Höhe.



Beträgt die Steigung einer Straße z.B. 4%, so steigt die Straße pro 100 m waagerechter „Landkarten-Entfernung“ Δx um $\Delta h = 4$ m an. Natürlich steigt sie dann z.B. auf $\Delta x = 200 \text{ m} = 2 \cdot 100 \text{ m}$ um $\Delta h = 2 \cdot 4 \text{ m} = 8 \text{ m}$ an, oder als weiteres Beispiel auf $\Delta x = 10 \text{ m}$ um 0,4 m.

Immer gleich ist dabei jedenfalls das Verhältnis $\frac{\Delta h}{\Delta x}$, nämlich $\frac{4 \text{ m}}{100 \text{ m}} = \frac{8 \text{ m}}{200 \text{ m}} = \frac{0,4 \text{ m}}{10 \text{ m}} = 0,04$ (ja eben die 4%)

Und genauso geben wir auch die Verlaufsrichtungen von Geraden in irgendwelchen x-y-Diagrammen als deren „Steigung“ an: Zwischen zwei Punkten A und B der Geraden gibt es eine Änderung Δy der y-Koordinate und eine Änderung Δx der x-Koordinate. Deren Verhältnis $\frac{\Delta y}{\Delta x} = \frac{y_B - y_A}{x_B - x_A}$ ist die Steigung.

Dabei kommt auch automatisch ein negativer Wert für fallende Geraden heraus. Die Beträge von Δy und Δx kannst Du als Längen der Katheten des rechtwinkligen „Steigungs-Dreiecks“ $\triangle ABC$ (siehe obiges Bild) ansehen.

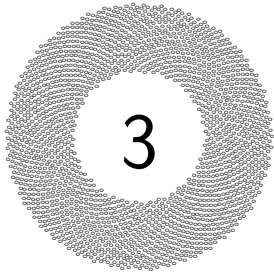
Wörtlich genommen „steigt“ eine Gerade natürlich nicht wirklich. Auf Papier gezeichnet kann sie sich ja überhaupt nicht bewegen. Du würdest dich sehr wundern, wenn eine von dir gezeichnete Linie vom Blatt emporsteigen oder sich überhaupt irgendwie bewegen würde! Bestenfalls steigt der Arm eines Lehrers, während er die Gerade an der Wand-Tafel zeichnet. (So wie auf einer steigenden Straße eigentlich nur das fahrende Auto an Höhe gewinnt, nicht die Straße.) Wir beschreiben die Richtung einer Geraden lediglich mit demselben mathematischen Ausdruck, der auch zur Beschreibung von Straßensteigungen geeignet ist, und der deshalb „Steigung“ heißt.

2.3.4 Proportionalität

Zusammenhänge wie der zwischen der Zeitspanne t und dem Weg s bei gleichförmiger Bewegung nennt man Proportionalitäten.

- Man sagt: Bei gleichförmiger Bewegung ist der zurückgelegte Weg s proportional zur dafür benötigten Zeit t .
- Man schreibt dafür symbolisch: $s \sim t$
- Man erkennt Proportionalität gleichermaßen daran, dass ...
 - ... alle aus je einem Messwertpaar (gleichartig) gebildete Quotienten den gleichen Wert haben,
 - ... zum n-fachen Wert der einen Größe auch der n-fache Wert der anderen Größe gehört,
 - ... bei graphischer Auftragung die Messpunkte auf einer Ursprungsgeraden liegen,
 - ... zu bestimmtem Zuwachs in der einen Größe immer auch ein bestimmter Zuwachs in der anderen Größe gehört und beide Größen gemeinsam den Wert Null annehmen.

Einen gemeinsamen Nullpunkt und damit einen Graphenverlauf durch den Ursprung kann man mit geeigneter Wahl oder Verschiebung von Nullpunkten immer erzwingen, sogar im Nachhinein: Hat z.B. s zur Zeit $t = 0 \text{ s}$ den Wert $s_0 \neq 0$, so haben zwar s und t keinen gemeinsamen Nullpunkt, wohl aber die Größen t und die (Differenz-)größe $\Delta s = s - s_0$.



Kraft

3.1 Erste Erfahrungen

Beim Begriff „Kraft“ denkt man im Alltag wohl am ehesten an Muskelkraft. Mit ihr sind wir schließlich von Geburt an bestens vertraut und haben reichlich (– wenn auch nie genug –) davon. Die physikalische Größe „Kraft“ deckt allerdings nicht dieselbe Vorstellung ab wie der Alltagsbegriff. Dennoch können wir die neue

Größe kennenlernen, indem wir von ganz alltäglichen Erfahrungen mit Muskelkraft ausgehen. Belaste Dich aber nicht damit, Deine bisherige, eventuell ungünstige Vorstellung von Kraft in der Physik beibehalten zu wollen! Nimm die Dinge einfach so, wie sie im Folgenden entwickelt werden!

?

Was kann man mit Muskelkraft machen?

3.1.1 Tauziehen

...geht bekanntlich so, dass zwei Menschen/Mannschaften an jeweils einem Ende eines dicken Seils (Tau=dickes Seil) ziehen. Wir halten schon einmal fest, dass die beiden Mannschaften dabei

in entgegengesetzte ziehen müssen.

Zweck eines Tauziehens ist natürlich, die stärkere Mannschaft zu ermitteln. Gelingt es keiner Mannschaft, das Seil zu sich herüber zu ziehen, so sind die Mannschaften offenbar gleich stark. Es greifen dann am Seil zwei gleich große, aber einander entgegengesetzt gerichtete Kräfte an.

Dies nennt man ein

Deine Alltagsvorstellung könnte Dir nahelegen, dass bei unterschiedlich starken Mannschaften kein Gleichgewicht am Seil herrsche. Das muss aber nicht stimmen! Bis auf ganz kleine Ungleichgewichte herrscht am Seil praktisch immer Kräftegleichgewicht, selbst wenn 100 erwachsene Männer gegen 5 Kindergartenkinder antreten! Dazu später mehr.

Um in eine Abbildung eine Kraft vollständig beschrieben einzuzeichnen, zeichnet man einen Pfeil nach den drei folgenden Regeln:

- Der Pfeil *beginnt* (!) am **Angriffspunkt** der Kraft. Kann man keinen einzelnen Angriffspunkt angeben, (z.B. weil eine Tauzieh-Mannschaft ja ihre gesamte Zugkraft über ein längeres Seilstück verteilt angreifen lässt,) so beginnt man den Pfeil wenigstens innerhalb des Körpers, auf den die Kraft ausgeübt wird.
- Als *Richtung* des Pfeiles nimmt man die **Richtung** der Kraft.
- Die *Länge* des Pfeiles soll der **Stärke** der Kraft entsprechen: Größere Kräfte werden durch längere Pfeile dargestellt.

3.1.2 Heben von Lasten

...ist eine weitere Anwendung von Muskelkraft. Bereits das *Halten* einer Last erfordert, dass man eine Kraft nach oben auf die Last ausübt. Um einen Koffer in gleichbleibender Höhe über dem Boden zu halten kann man zum Beispiel an seinem Griff nach oben ziehen. Alternativ kann man mit derselben Kraft von unten gegen den Koffer nach oben drücken.

Würde man keine Kraft auf den Koffer ausüben, so würde er nach unten fallen. Den Grund für diesen („freien“) Fall kennst Du schon: Die Erde übt auf jeden Körper eine gewisse Anziehungskraft aus. (Du erinnerst Dich vielleicht an den Elefanten, der im Skript für das 7. Schuljahr einem Physiker auf den Fuß getreten ist.)

Am gehaltenen Koffer greifen also zwei Kräfte an: die Erdanziehungskraft nach unten und die Haltekraft des Kofferträgers nach oben. Bleibt der Koffer trotz dieser Kräfte in Ruhe, herrscht offensichtlich auch hier ein Kräftegleichgewicht.



3.1.3 Verformung

Erfahrungsgemäß lassen sich mit Kräften auch Verformungen hervorrufen. Man unterscheidet **elastische** Verformungen, die nach dem Ende der Krafteinwirkung wieder verschwinden, und **plastische** Verformungen, die auch nach der Krafteinwirkung bleiben. Ein Gummiband wird also durch (nicht allzu große) Kräfte

elastisch verformt, ein Knetklumpen hingegen plastisch. Die Adjektive elastisch und plastisch werden auch als entsprechende Materialeigenschaften verwendet: elastisches Gummi und plastische Knetmasse. (Warum heißt Plastik „Plastik“?)

3.1.4 Beschleunigung

Ein losgelassener Koffer zeigt uns, was passiert, wenn einmal kein Kräftegleichgewicht herrscht. Auf den Koffer wirkt ja nach dem Loslassen nur noch die Erdanziehungskraft nach unten. Der Koffer fällt dann bekanntlich nach unten, bewegt sich also immer schneller in Richtung der wirkenden Kraft. Für dieses Anwachsen der Geschwindigkeit sagt man: Der Koffer wird *beschleunigt*.

Das Verb „beschleunigen“ kennst Du sicher schon im Zusammenhang mit Fahrzeugen: Wenn z.B. ein Auto seine Geschwindigkeit von $60 \frac{\text{km}}{\text{h}}$ auf $80 \frac{\text{km}}{\text{h}}$ steigert, sagen wir auch, dass es *beschleunigt*.

Um nicht unnötig viele Begriffe einzuführen, sprechen wir in der Physik nicht nur dann von einer **Beschleunigung**, wenn die Geschwindigkeit *zunimmt*, sondern auch, wenn die Geschwindigkeit *abnimmt*. Die Beschleunigung ist dann einfach nur negativ: Eine Abnahme um $20 \frac{\text{km}}{\text{h}}$ wird eben einfach als eine Zunahme

um $-20 \frac{\text{km}}{\text{h}}$ angesehen.

Für Physiker liegt eine Beschleunigung sogar schon dann vor, wenn sich nur die Bewegungs*richtung* ändert. Die Geschwindigkeit ist nämlich eine Größe, die nicht nur einen Betrag, sondern auch eine Richtung hat. Jede Änderung der Bewegungsrichtung ist deshalb auch eine Änderung der Geschwindigkeit.

Unter der Beschleunigung versteht man nun aber nicht einfach die Geschwindigkeitsänderung. Wenn ein Auto seine Geschwindigkeit von 0 auf $100 \frac{\text{km}}{\text{h}}$ ändert, sagt das nämlich noch gar nichts über das aus, was wir als Beschleunigung bezeichnen: Bei einem Rennwagen, der diese Änderung in 5 Sekunden schafft, sprechen wir von einer höheren Beschleunigung als bei einem Lastwagen, der für dieselbe Geschwindigkeitsänderung 150 Sekunden benötigt. Der Rennwagen schafft ja in jeder Sekunde einen größeren Geschwindigkeitszuwachs als der Lastwagen.

Unter der Beschleunigung a verstehen wir die Geschwindigkeitsänderung Δv pro Zeitspanne Δt . Kurzfassung:

$$a :=$$

zu ergänzen:

- Formel

Der Buchstabe a wurde gewählt, weil Beschleunigung auf lateinisch *acceleratio* heißt.

Der Rennwagen von eben beschleunigt mit

$$a_{\text{Renn}} = \frac{\Delta v}{\Delta t} = \frac{100 \frac{\text{km}}{\text{h}}}{5 \text{ s}} = 20 \frac{\frac{\text{km}}{\text{h}}}{\text{s}} = 20 \frac{\text{km}}{\text{h} \cdot \text{s}}$$

Der Lastwagen beschleunigt nur mit

$$a_{\text{LKW}} = \frac{\Delta v}{\Delta t} = \frac{100 \frac{\text{km}}{\text{h}}}{150 \text{ s}} = \frac{2}{3} \frac{\frac{\text{km}}{\text{h}}}{\text{s}} \approx 0,667 \frac{\text{km}}{\text{h} \cdot \text{s}}$$

In 7 Sekunden von $260 \frac{\text{m}}{\text{s}}$ auf $50 \frac{\text{m}}{\text{s}}$ abbremsen entspricht einer Beschleunigung von

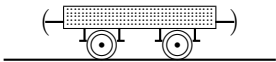
$$a_{\text{Brems}} = \frac{\Delta v}{\Delta t} = \frac{50 \frac{\text{m}}{\text{s}} - 260 \frac{\text{m}}{\text{s}}}{7 \text{ s}} = \frac{-210 \frac{\text{m}}{\text{s}}}{7 \text{ s}} = -30 \frac{\frac{\text{m}}{\text{s}}}{\text{s}} = -30 \frac{\text{m}}{\text{s}^2}$$

Ein fallender Koffer verdankt seine Beschleunigung der Erdanziehungskraft, ein positiv beschleunigendes Auto verdankt sie seiner Motorkraft, ein negativ beschleunigendes (=bremsendes) Auto seiner Bremskraft, ein Geschoss der Druckkraft des explodierenden Schießpulvers. Es gilt sogar generell:

Die Ursache jeder Beschleunigung eines Körpers ist eine Kraft, die auf den Körper ausgeübt wird.

Die Erkenntnis der engen Beziehung zwischen Kraft und Beschleunigung war ein großer Durchbruch in der Physik. Lange Zeit hatte man zu dieser Thematik nämlich ziemlich falsche Vorstellungen, die den wissenschaftlichen Fortschritt jahrhundertlang hemmten.

Betrachten wir dazu einen Modelleisenbahn-Waggon, der auf ebenem Gleis steht. Wir üben nun eine Kraft entlang des Gleises auf den Waggon aus, indem wir ihn anschieben. Solange wir kräftig schieben, wird der Waggon immer schneller. Hören wir auf zu schieben, so hört auch diese Beschleunigung auf. Der Waggon rollt anschließend aus und kommt wieder zum Stehen.



Heute erklärt man das Langsamerwerden so: Durch den Luftwiderstand, die Reibung der Achsen in ihren Lagern und der Räder auf den Schienen erfährt der rollende Waggon Kräfte entgegen der Fahrtrichtung, wodurch der Waggon negativ beschleunigt, also abgebremst wird.

Bis ins 17. Jahrhundert n.Chr. hielt sich jedoch die folgende Sichtweise, die auf **Aristoteles** im 4. Jahrhundert v.Chr. zurückgeht:

- Auf der Erde bewegen sich Lebewesen aufgrund ihrer Seele von sich aus.
- Himmelskörper bewegen sich nach einer ewigen Ordnung und von sich aus (Sie müssen folglich eine Seele haben!) Da ihre Bewegungen aus perfekten Kreisbewegungen bestehen, sind Himmelskörper göttlich.
- Andere Körper können sich bewegen, um eine gestörte Ordnung wiederherzustellen. Diese Ordnung sieht vor, dass sich schwere Körper unten und leichte oben befinden. Außerdem ist Ruhe der ordnungsgemäße Zustand.
- Die Bewegung eines Körpers kann ansonsten erzwungen werden, indem ein anderer Körper über unmittelbaren Kontakt eine Kraft ausübt. Dabei ist die Geschwindigkeit höher, wenn die Kraft größer ist.

Nach dieser Sichtweise fällt der Koffer nach unten, weil er als schwerer Körper dort seinen Platz hat. Der Waggon bewegt sich erzwungen, solange jemand oder etwas ihn anschiebt. Ansonsten kommt er zur Ruhe, weil die Ordnung dies als seinen Grundzustand vorsieht. Nur Lebewesen und Himmelskörper bewegen sich einfach so. Zugegebenermaßen passen diese Deutungen recht gut zu allen damaligen und heutigen Alltagserfahrungen. Eine über den Alltag hinausgehende Entwicklung der Physik ist mit einem solchen Weltbild jedoch kaum möglich.

PHILOSOPHIÆ
NATURALIS
P R I N C I P I A
MATHEMATICA

Erst **Newton** formulierte in seinem 1687 erschienenen Werk → (kurz *Principia* genannt) die bis heute bewährte Sichtweise von Kräften. Von entscheidender Bedeutung ist dabei die Erkenntnis:

Bewegung ist kein *Prozess*, dessen *Fortdauer* durch eine Kraft gewährleistet werden muss, sondern ein *Zustand*, zu dessen *Änderung* eine Kraft nötig ist.

Sir Isaac Newton,
* 4.1.1643 in
Woolsthorpe/Grantham
† 31.3.1727 in
Kensington (London)

Ist – wie im Waggon-Beispiel – zur Aufrechterhaltung einer Bewegung erfahrungsgemäß doch eine Kraft nötig, so gleicht diese nur andere, meist durch Reibung bedingte Kräfte aus, die die Bewegung sonst hemmen würden. Bei Berücksichtigung tatsächlich *aller* Kräfte gilt stets der

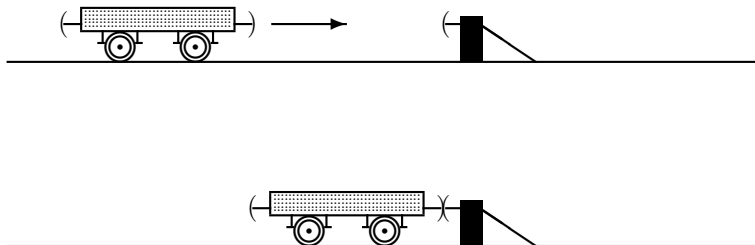
Trägheitssatz: Ein Körper, auf den keine Kraft einwirkt, behält seine Geschwindigkeit bei. Er bewegt sich also geradlinig gleichförmig weiter oder bleibt in Ruhe.

Die Eigenschaft jedes Körpers, von sich aus seine Geschwindigkeit beizubehalten, solange keine Kraft auf den Körper einwirkt, nennt man **Trägheit**.

V 3.1 Aufprall eines beladenen Wägelchens

Aufbau: Auf einem Wägelchen liegt lose ein Klotz auf.

Durchführung: Das Wägelchen wird angeschoben und gegen ein festes Hindernis prallen gelassen.



zu ergänzen:

- Klotz

Beobachtung:

Erklärung:

3.2 Dynamische Definition von Kraft

?

Welcher rechnerische Zusammenhang besteht zwischen einer Beschleunigung eines Körpers und der verursachenden Kraft?

Man weiß aus dem Alltag bereits Folgendes:

1) Um einem Körper eine größere Beschleunigung zu erteilen, ist eine größere Kraft erforderlich. Angenommen, Du willst im Supermarkt einen zunächst stehenden, vollen Einkaufswagen (mit leichtgängigen Rollen) innerhalb von 4 Sekunden auf eine bestimmte Geschwindigkeit beschleunigen. Dazu musst Du eine bestimmte Kraft aufwenden. Willst Du dieselbe Geschwindigkeit bereits nach 2 Sekunden erreichen, brauchst Du sicher mehr Kraft.

2) Eine größere Kraft benötigt man andererseits auch, um einem schwereren Körper die gleiche Beschleunigung zu erteilen. Ist der Einkaufswagen im vorigen Beispiel doppelt so schwer, so wirst Du für dieselbe Beschleunigung wohl auch mehr Kraft brauchen. Schiebst Du einen doppelt schweren Wagen mit derselben Kraft an, wirst Du nur die halbe Beschleunigung erreichen.

Vor diesem Hintergrund legt man fest, eine Kraft x -mal so groß wie eine Vergleichskraft zu nennen, wenn sie einen x -mal so schweren Körper gleich stark beschleunigt, oder einen gleich schweren Körper x -mal so stark beschleunigt. Dies ist mit folgender Definition erfüllt:

Eine Kraft, die einen sonst kräftefreien Körper der Masse m mit a beschleunigt, hat den Betrag

$$F =$$

Hierzu passend definierte Einheit:

Eine Kraft beträgt 1 Newton (abgekürzt 1 N), wenn sie einen Körper der Masse 1 kg in einer Sekunde um die Geschwindigkeit $1 \frac{\text{m}}{\text{s}}$ beschleunigt. Kurz:

$$1 \text{ Newton} = 1 \text{ N} :=$$

zu ergänzen:

- Formeln

Die Gleichung $F = m \cdot a$ wird wegen ihrer zentralen Bedeutung in der Mechanik „*Grundgleichung der Mechanik*“ genannt. Sie wird auch als 2. Newton'sches Axiom bezeichnet. Das 1. Newton'sche Axiom ist der Trägheitssatz. Die Bezeichnung „Axiom“ ist dabei streng genommen falsch, weil ja mit $F = m \cdot a$ aus $F = 0$ sofort $a = 0$ folgt. (Zumindest solange man über reale Körper redet, die ja immer $m \neq 0$ erfüllen) Das dritte und letzte Newton'sche Axiom ist das Reaktionsprinzip. (→ Seite 28)

3.3 Kraftmessung

Es ist aufwändig, eine Kraft gemäß ihrer Definition über Beschleunigungen zu messen: Man benötigt eine ebene Fahrbahn, auf der z.B. ein 1 kg schwerer Wagen reibungsarm rollen kann. Man muss die zu messende Kraft eine Zeit lang auf den Wagen einwirken lassen. Man muss die Momentangeschwindigkeiten des Wagens vorher und nachher bestimmen, wozu jeweils eine Weg- und eine Zeitmessung erforderlich sind. Schließlich muss man die Messwerte noch miteinander verrechnen. Leichter geht es, indem man die zu messende Kraft z.B. eine Schraubenfeder elastisch verformen lässt.

Wie verläuft eine solche Messung im Prinzip?

?

Idee: Die zu messende Kraft zieht an einem Ende der Feder, so dass diese sich dehnt. Dadurch entsteht eine Spannkraft in der Feder, die die Feder in die entspannte Ausgangslage zurückzustellen sucht. Mit zunehmender Dehnung wird diese Spannkraft (auch Rückstellkraft genannt) immer größer.

Bei einer gewissen Dehnung sind die zu messende Kraft und die Rückstellkraft gleich stark. Unterhalb dieser Dehnung überwiegt die zu messende Kraft, was zugunsten einer Vergrößerung der Dehnung wirkt. Oberhalb überwiegt die Rückstellkraft, was die Dehnung zu verringern strebt. Die Feder kann also nur mit derjenigen Dehnung zur Ruhe kommen, bei der Kräftegleichgewicht herrscht.

Für eine größere zu messende Kraft ist diese Dehnung natürlich auch größer. An der Dehnung kann man also erkennen, wie groß die dehnende Kraft ist. Es fehlen jetzt nur noch Markierungen, welche Dehnung welchen Kraftbetrag anzeigen. Die erste Markierung ist dabei sehr leicht zu setzen: Bei entspannter Feder zeigt ihr unteres Ende die Stelle für die „0 N“-Markierung an.

Die weiteren Markierungen kann man grundsätzlich so setzen: Wenn ich wissen will, welche Dehnung z.B. 47,11 N anzeigt, lasse ich 47,11 N auf die Feder einwirken und markiere die sich dann einstellende Dehnung. Natürlich fangen wir nicht mit 47,11 N an, sondern etwas systematischer mit 1 N. Wir stehen damit vor der Frage:

Wie können wir eine Kraft von 1 N erzeugen?

?

Ideal wäre es, ein Gewichtsstück zu finden, das genau 1 N als Gewichtskraft hat. Dieses Gewichtsstück könnte man dann nämlich einfach an die Feder hängen. Nach der Definition von 1 Newton müsste die Gewichtskraft des gesuchten Gewichtsstückes einen 1 kg schweren Körper in einer Sekunde um $1 \frac{\text{m}}{\text{s}}$ beschleunigen. Doch welches Gewichtsstück leistet dieses? Zum Ausprobieren verschiedener Gewichte könnte man anstelle von Gewichtsstücken einen mit verschieden viel Wasser gefüllten Becher verwenden. Damit könnte man „tropfengenau“ beliebige Gewichte herstellen. Man würde feststellen:

1 N ist etwa die Gewichtskraft von 102 ml Wasser.

Da 1 ml Wasser etwa 1 g wiegt, können wir anstelle des Wassers auch ohne großen Fehler ein 100 g-Gewichtsstück verwenden.

1 N ist etwa die Gewichtskraft eines 100 g-Gewichtsstückes.

Also: 100 g-Stück an die Feder hängen, sich einstellende Dehnung markieren, mit „1 N“ beschriften, fertig.

Wie setzen wir weitere Markierungen?

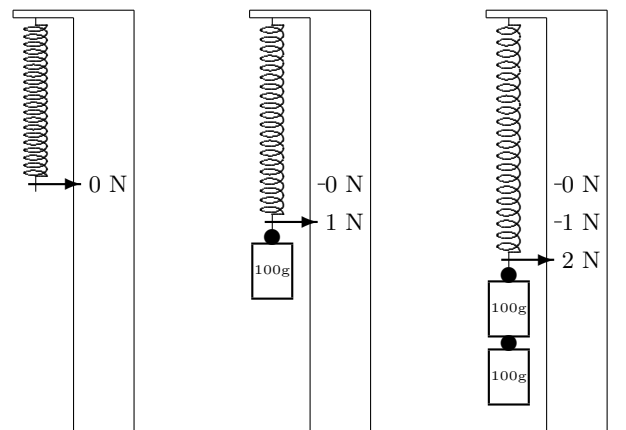
?

Aus unserer alltäglichen Erfahrung mit Gewichten erwarten wir zu Recht, dass wir eine Kraft von 2 N erhalten, wenn wir 200 g verwenden.

Also: Zwei 100 g-Stücke an die Feder hängen, Dehnung markieren und mit „2 N“ beschriften, fertig.

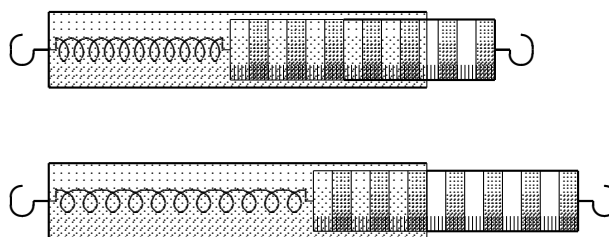
In dieser Weise können wir beliebige Kräfte erzeugen und entsprechende Markierungen setzen: Der benötigte Gewichtskörper ergibt sich nach der „Formel“: *Gewünschte Kraft in Newton mal 100 g*.

Den ganzen Vorgang zum Anbringen aller Markierungen nennt man eine **Eichung**. Hier nebenan noch einmal die Vorgehensweise bildlich. Mit der so geeichten Feder kann man nun Messungen ausführen: Man lässt die zu messende Kraft die Feder dehnen und liest der Dehnung entsprechend die Kraft ab.



Die Kraftmesser, die in der Schule meist benutzt werden, enthalten ebenfalls eine Schraubenfeder. Die folgende Abbildung zeigt einen solchen Kraftmesser, zum Teil durchsichtig dargestellt.

Die Feder ist innerhalb einer Hülse zwischen der Hülse (links) und einem skalierten Zylinder (rechts) eingespannt. Zieht man Zylinder und Hülse auseinander, so dehnt man die Feder. Der offene Rand der Hülse zeigt auf der Skala am Zylinder die Zugkraft an.



3.4 Hookesches Gesetz

An unseren Kraftmessern fällt auf, dass die Markierungen für 1 N, 2 N, usf. in gleichen Abständen liegen. Jedes zusätzliche Newton Zugkraft bewirkt immer den gleichen Zuwachs an Dehnung. Die n -fache Kraft bewirkt die n -fache Dehnung. Kurz: Kraft und Dehnung sind bei unseren Kraftmessern zueinander proportional. Wie der nächste Versuch bestätigt, gilt näherungsweise für jeden festen Körper:

Hookesches Gesetz: Bei nicht allzu großer Dehnung s und die Kraft F zueinander proportional: $F \sim s$

Robert Hooke, * 18.7.1635 in Freshwater/Isle of Wight, † 3.3.1703, London

Wie groß „allzu groß“ ist, hängt vom jeweiligen Körper ab. Bei Schraubenfedern kann die Dehnung durchaus ein Mehrfaches der entspannten Federlänge betragen. Bei einem weichen Gummiseil kann bereits ein kleiner Bruchteil der Ausgangslänge „allzu groß“ sein.

Bei Proportionalität zwischen zwei Größen ist deren Quotient bekanntlich konstant. Im Gültigkeitsbereich des Hookeschen Gesetzes sind bei einer Feder also die Quotienten $\frac{F}{s}$ und natürlich ebenso $\frac{s}{F}$ jeweils konstant.

Meist wird der Quotient $\frac{F}{s}$ betrachtet. Er gibt an, wie viel Kraft F pro Dehnung s erforderlich ist. Ist dieser Quotient bei einer Feder A größer als bei Feder B, so ist demnach an A mehr Kraft für die gleiche Dehnung nötig als an B. Feder A ist *härter* als Feder B. Daher definiert man:

Den Quotienten $\frac{F}{s}$ nennt man die **Federhärte** oder einfach **Federkonstante**. Symbol: D

$$D :=$$

Sinngemäß gleiches wie beim Dehnen gilt auch beim Stauchen, Biegen, Verdrillen, Scheren eines Körpers.

An einem Gummiseil, aber auch an einer Schraubenfeder, die fast zu einem geraden Draht auseinandergezogen wird, können wir außerdem erkennen:

Bei stärkerer Belastung kann ein Körper eine dauerhafte Verformung erleiden, d.h. er geht nach Entlastung nicht wieder in seine ursprüngliche Form zurück.

Zur Beschreibung dieser Erscheinung benutzen wir folgende Begriffe:

Eine Verformung heißt **elastisch**, wenn sie bei Entlastung wieder vollständig zurückgeht, sonst **plastisch**.

Solange wir mit den Zugkräften in dem Wertebereich bleiben, in dem Proportionalität zwischen F und s besteht, bewirken wir auch nur elastische Verformungen.

V 3.2 Dehnung verschiedener Körper

Aufbau: Verschiedene Körper hängen neben einer Messlatte, so dass ihre Dehnungen abgelesen werden können. Die Körper sind:

① _____, ② _____, ③ _____

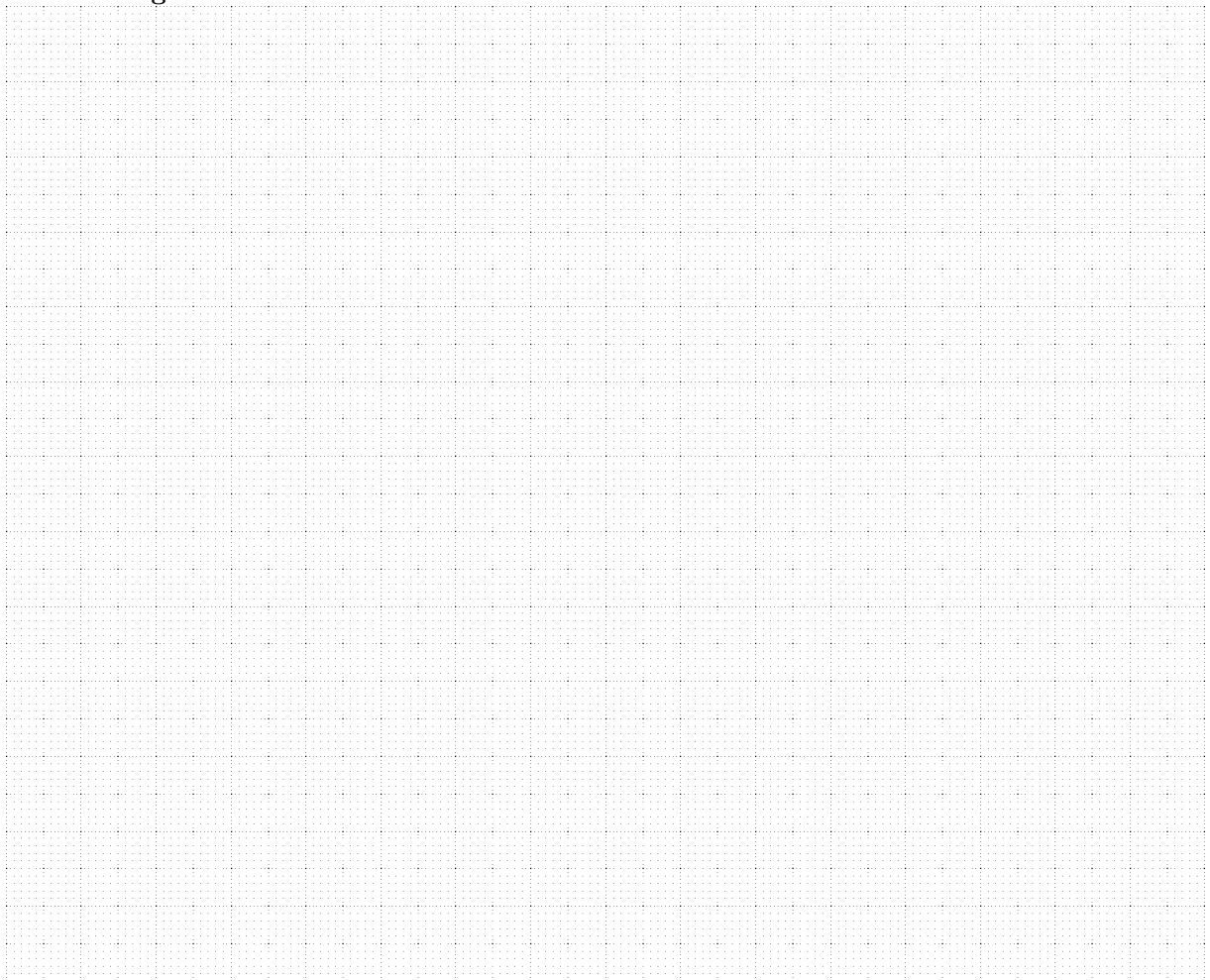
Durchführung: Wir belasten die Körper mit immer schwereren Gewichten und lesen jeweils die Dehnung ab.

Messwerte:

①

②

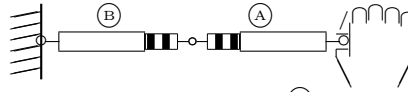
③

Auswertung:

3.5 Reaktionsprinzip

V 3.3 Über einen Kraftmesser an einem Kraftmesser ziehen

Aufbau:



Durchführung: Wir ziehen über einen Kraftmesser A an einem Kraftmesser B, der an der Wand befestigt ist.

Ergebnis: Dort, wo die Kraftmesser ineinander verhakt sind, übt jeder der Kraft-

messer auf den anderen eine Kraft aus. Diese beiden Kräfte

Nicht nur für Kraftmesser, sondern ganz allgemein gilt als drittes Newton'sches Axiom das sogenannte

Reaktionsprinzip: Übt ein Körper A auf einen Körper B eine Kraft aus, so übt auch Körper B auf Körper A eine gleich große, jedoch entgegengesetzt gerichtete Kraft („Gegenkraft“) aus.
(Kurze Merkhilfe: „*actio=reactio*“)

Kräfte treten immer in Kraft/Gegenkraft-Paaren auf, niemals einzeln. Jedes solche Paar ist Ausdruck einer Wechselwirkung zwischen *zwei* Objekten, die in gegenseitiger Anziehung oder gegenseitiger Abstoßung besteht. Das Reaktionsprinzip hat nichts zu tun mit einem Kräftegleichgewicht, bei dem *ein* Körper aufgrund *mehrerer* Wechselwirkungen mehrere Kräfte erfährt, die sich gegenseitig aufheben.

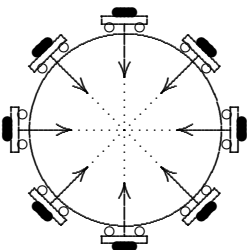
Beispiel: Wenn Du auf dem Boden stehst, befindest

Du Dich im Gleichgewicht der beiden Kräfte, mit denen die Erde Dich nach unten zieht bzw. der Boden Dich nach oben drückt. Die Gegenkräfte zu diesen Kräften sind diejenigen, mit denen Du die Erde anziehst bzw. Du den Boden belastest. Es finden zwei Wechselwirkungen statt, an denen Du beteiligt bist, nämlich die gegenseitige Anziehung zwischen Dir und dem Planeten Erde sowie das gegenseitige Aneinanderdrücken zwischen Dir und dem Boden.

3.6 Gewichtskraft

3.6.1 Eigenschaften

Gewichtskräfte erlebt ein Mensch in seinem Alltag ständig: Er selbst hat ein Gewicht und alle Gegenstände um ihn herum haben Gewicht. („ein Gewicht haben“ sagen wir kurz für „einer Gewichtskraft unterliegen“) Wir haben sogar schon Gewichtskräfte benutzt, um Kräfte zu untersuchen und zu messen. Für praktische Zwecke benutzen wir die Gewichtskraft von gut 100 ml Wasser als „Ur-Newton“. Fassen wir einmal zusammen, was wir über Gewichtskräfte wissen:



- Jeder Gegenstand in unserer Umwelt unterliegt einer Gewichtskraft.
- Die Gewichtskraft wirkt zum Boden hin, nach unten. Dies ist überall auf der runden Erde so. Wir sagen also genauer: Die Gewichtskraft wirkt zum Erdmittelpunkt hin.
Die Richtung der Gewichtskraft ergibt sich nämlich eigentlich nicht aus der Richtung, wo aus irgendwelchen Gründen „unten“ ist, sondern „unten“ ist jeweils dort, wohin die Gewichtskraft wirkt!
- Unserer Erfahrung nach hat jeder Körper ein für ihn typisches Gewicht, solange man ihm nichts hinzufügt oder abtrennt.
- Wie alle Kräfte kann man Gewichtskräfte grundsätzlich mit einem Feder-Kraftmesser messen.

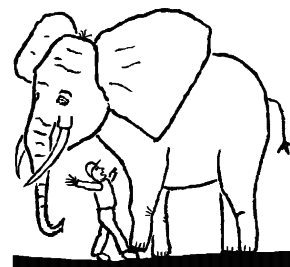
Manche Menschen (Astronauten) haben zusätzlich Folgendes erfahren:

- Derselbe Körper kann an verschiedenen Orten (z.B. verschiedenen Planeten) verschiedene Gewichtskräfte erfahren.
Beispiel: Eine Tafel Schokolade erfährt auf der Erde etwa 1 N Gewichtskraft, auf dem Mond nur etwa $\frac{1}{6}$ N.
- Erfahren zwei Körper an einem Ort die gleiche Gewichtskraft, so erfahren sie an jedem anderen Ort auch gleiche Gewichtskräfte.
Beispiel: Erfährt ein Apfel auf der Erde die gleiche Gewichtskraft wie eine Tafel Schokolade, nämlich 1 N, so erfährt er auch auf dem Mond die gleiche Gewichtskraft wie die Schokolade, nämlich $\frac{1}{6}$ N.
- Weit weg von Sternen, Planeten und anderen Himmelskörpern sind gar keine Gewichtskräfte bemerkbar.

Das Gewicht eines Körpers ist also eigentlich keine Eigenschaft des *Körpers alleine*, sondern ergibt sich aus der *Kombination aus Körper und Aufenthaltsort*. Wenn wir bisher das Gewicht für eine feste Eigenschaft eines Körpers gehalten haben, liegt das nur daran, dass wir in unserem Leben noch nicht so weit von der Erdoberfläche weggekommen sind.

Als Physiker möchte man sich das Leben natürlich nicht unnötig schwer machen. Selbst wenn man nur 10 Körper an 5 möglichen Orten betrachten wollte, müsste man für 10·5, also für 50 Kombinationen aus Körpern und Orten die jeweilige Gewichtskraft festhalten.

Günstiger wäre es, wenn jeder Körper eine ortsunabhängige Eigenschaft und jeder Ort eine Körper-unabhängige Eigenschaft hätte, aus denen man bei Bedarf für jede Kombination die Gewichtskraft ausrechnen könnte. Für 10 Körper und 5 Orte würden dann 10 + 5, also 15 Daten genügen, wo man andernfalls 50 bräuchte.



... auf dem Mond
(noch nicht einmal)
halb so dramatisch.

3.6.2 Masse

Da zwei Körper, die an einem Ort gleiches Gewicht haben, an jedem anderen Ort auch gleiches Gewicht haben, muss es offensichtlich eine Eigenschaft von Körpern geben, die ein Körper immer und überallhin „mitnimmt“, und die für das Gewicht des Körpers mit verantwortlich ist.

Diese Eigenschaft nennen wir die **Masse** eines Körpers. Die Masse kürzen wir meist mit m ab. Als Grund-Einheit für Massen hat man das **Gramm** mit der Abkürzung g festgelegt.

Man hat dazu einfach ein Stück Metall hergestellt und festgesetzt: *Dieses Stück Metall hat eine Masse von 1000 g, also von einem Kilogramm*. Dieses Metallstück nennt man deshalb das **Ur-Kilogramm**. Bei seiner Herstellung hat man sich am Gewicht von einem Liter Wasser orientiert. Wie das Ur-Meter wird auch das Ur-Kilogramm in Paris aufbewahrt.

Ein Körper hat nun eine Masse von einem Kilogramm, wenn er die gleiche Gewichtskraft erfährt, die das Ur-Kilogramm am selben Ort erfahren würde. Setzt man aus zwei Körpern mit jeweils 1 kg Masse einen neuen Körper zusammen, so hat dieser eine Masse von 2 kg. Alle Erfahrung bestätigt, dass der 2 kg-Körper die doppelte Gewichtskraft wie das Ur-Kilogramm am selben Ort erfährt. Ein Körper hat allgemeiner ausgedrückt x kg Masse, wenn er am selben Ort die x -fache Gewichtskraft wie das Ur-Kilogramm erfährt.

Mathematisch präziser ausgedrückt: Wir schreiben einem Körper eine Masse zu, die proportional zur Gewichtskraft ist, die der Körper an einem gleichbleibenden Ort erfährt. Diese Proportionalität haben wir übrigens schon bei der Eichung eines Kraftmessers vorausgesetzt.

Die im Alltag gebräuchliche Sprechweise ... *ein Gewicht von soundsoviel Gramm* ... müssen wir jetzt aus unserem Sprachgebrauch verbannen: Gewichte sind Gewichtskräfte und werden in der Einheit Newton angegeben, die Einheit Gramm bezeichnet Massen!

?

3.6.3 Ortsfaktor

Welche Gewichtskraft erfährt ein Körper der Masse 1 kg auf der Erde?

Diese Frage können wir leicht beantworten, indem wir einen Körper der Masse 1 kg an einen geeigneten Kraftmesser hängen und seine Gewichtskraft ablesen. Wir können uns diese Messung jedoch sparen, wenn wir bedenken, wie wir vor kurzem Kraftmesser geeicht haben. Wir haben für die 1 N-Marke an eine Feder einen 100 g-Stein gehangen, für die 2 N-Marke zwei 100 g-Steine usw.

Natürlich hat 1 kg ($= 10 \cdot 100 \text{ g}$) demnach ein Gewicht von 10 N.

„Auf der Erde“ müssen wir jedoch hinzufügen. Wie Astronauten uns berichten, würde derselbe Körper der Masse 1 kg nämlich auf dem Mond nur $\frac{1}{6}$ so viel Gewicht haben, also etwa $\frac{10}{6} \text{ N}$, rund 1,7 N.

Auf der Erde erfährt jedes kg Masse eine Gewichtskraft von 10 N. Das heißt auch: Auf der Erde erfährt jeder Körper eine Gewichtskraft von 10 Newton pro Kilogramm Körpermasse. Das Gewicht eines Körpers ergibt sich demnach als seine Masse in Kilogramm multipliziert mit 10 Newton pro Kilogramm. Für das Gewicht des Körpers auf dem Mond multiplizieren wir seine Masse nur mit 1,7 Newton pro Kilogramm.

Beispiel: Ein Körper der Masse 3,2 kg erfährt auf der Erde eine Gewichtskraft von $F_E = 3,2 \text{ kg} \cdot 10 \frac{\text{N}}{\text{kg}} = 32 \text{ N}$. Auf dem Mond erfährt er $F_M = 3,2 \text{ kg} \cdot 1,7 \frac{\text{N}}{\text{kg}} = 5,44 \text{ N}$ Gewichtskraft.

$$g_{\text{Saarland}} = 9,81 \frac{\text{N}}{\text{kg}}$$

Um das Gewicht eines Körpers mit bekannter Masse an einem bestimmten Ort zu berechnen, muss also nur noch dieser „Gewicht pro Masse“-Faktor für diesen Ort bekannt sein. Da dieser Faktor charakteristisch für den Ort ist, nennt man ihn den **Ortsfaktor** des Ortes. In Gleichungen kürzt man ihn meist mit g ab. Beispiel: $g_{\text{Erde}} = 10 \frac{\text{N}}{\text{kg}}$

Dieser Wert ist allerdings nur ein Näherungswert. Der genaue Wert hängt sogar ein wenig davon ab, wo genau auf der Erde man sich befindet. In größeren Höhen ist der Ortsfaktor wegen der größeren Entfernung zur Erde i.A. etwas kleiner als in geringerer Höhe. Außerdem: Da die Erde an den Polen ein wenig abgeplattet ist, ist der Ortsfaktor an den Polen auch etwas größer als am Äquator. Für unsere Zwecke merkst Du Dir einfach nebenstehenden Wert.

Unsere Vorgehensweise zur Berechnung von Gewichtskräften können wir nun auch in einer Gleichung ausdrücken:

Für die Gewichtskraft F_G , die ein Körper der Masse m an einem Ort mit dem Ortsfaktor g erfährt, gilt:

$$F_G =$$

zu ergänzen:

- Formel

Merkur	3,70
Venus	8,87
Erde	9,78
Mars	3,71
(Erdmond)	1,62
Jupiter	23,12
Saturn	8,96
Uranus	8,69
Neptun	11,00
Pluto	0,72

Die nebenstehende Tabelle zeigt für die Planeten unseres Sonnensystems den Ortsfaktor an ihrem Äquator in der Einheit $\frac{\text{N}}{\text{kg}}$.

3.7 Kraft als Vektor

3.7.1 Merkmale einer Kraft

Wenn Du an einem Stuhl ziehst, wirst Du ganz verschiedene Bewegungen des Stuhles bewirken, nicht nur wenn Du

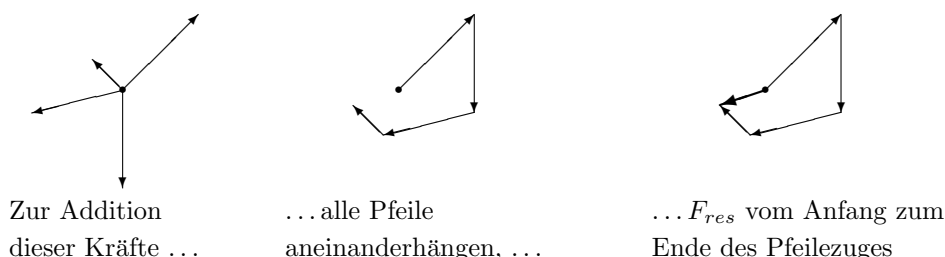
- verschieden stark ziehst oder
- in verschiedene Richtungen ziehst, sondern auch, wenn Du
- an verschiedenen Stellen des Stuhles angreifst.

Der Stuhl wird nämlich je nach Angriffspunkt sich verschieden bewegen, drehen, kippen.

zu ergänzen:

- F_{res} rot nachzeichnen
- die 4 einzelnen Kräfte in 4 anderen Farben

Vorgehensweise bei mehr als zwei Kräften: Ebenfalls alle Kraftpfeile in beliebiger Reihenfolge hintereinander hängen und F_{res} vom Anfangs- zum Endpunkt dieses Pfeilezuges zeichnen. Beispiel:



Andere Vorgehensweise: Resultierende von zwei Kräften konstruieren, dann diese Resultierende mit einer dritten Kraft zu neuer Resultierenden kombinieren, diese dann mit der vierten Kraft kombinieren usw.

Greifen an einem Körper mehrere Kräfte an *verschiedenen* Punkten an, wird die Sache komplizierter. Einige Gesetzmäßigkeiten darüber lernst Du bald im Kapitel über Drehmomente kennen. Einen wichtigen Sonderfall klären wir gleich hier: die Gewichtskraft eines Körpers!

Die Gewichtskraft eines Körpers ist nämlich streng genommen gar nicht *eine einzige* Kraft, sondern ergibt

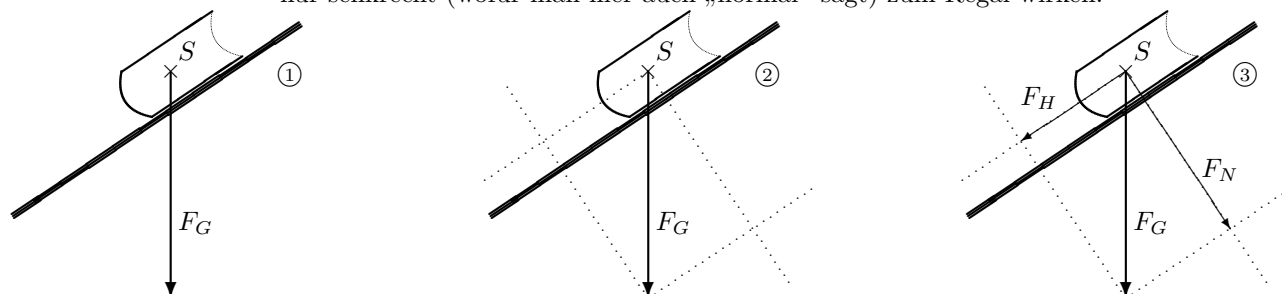
sich als Summe all der winzigen Gewichtskraftchen, die jedes einzelne winzige Teilstückchen (z.B. Atom) des Körpers erfährt. Jedes einzelne davon hat seinen eigenen Angriffspunkt, nämlich das betroffene Teilstückchen. (– sofern wir es uns erlauben, bei hinreichender Winzigkeit des Stückchens idealisierend von einem Punkt zu sprechen; ohne diese Erlaubnis wäre der Begriff „Angriffspunkt“ sowieso nie sinnvoll)

Wir müssen uns hier nicht weiter um diese Kraft-Winzlinge kümmern. Da nämlich (– zumindest für nicht allzu große Körper an der Erdoberfläche –) sämtliche Mini-Teilgewichtskräfte praktisch parallel wirken, hat ihre Resultierende einfach als Betrag die Summe der Beträge der Winzlinge, und sie wirkt in derselben Richtung wie die Winzlinge. Damit sich durch diese Gesamt-Gewichtskraft genau dieselbe Wirkung (– zumindest auf einen starren Körper –) errechnet wie durch die Winzlinge, muss man nur noch den richtigen Angriffspunkt für sie annehmen: den **Schwerpunkt** des Körpers, eine Art Mittelwert aus den Angriffspunkten aller Winzlinge.

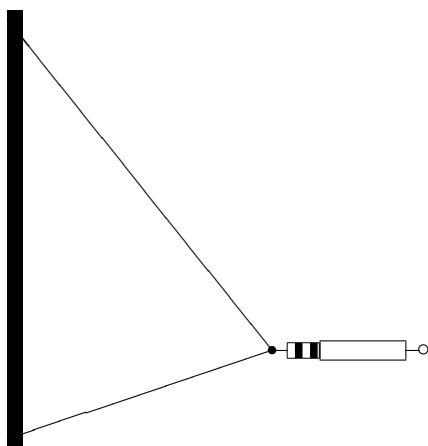
Angriffspunkt
der Gewichtskraft
==
Schwerpunkt
des Körpers

3.7.3 Zerlegung von Kräften

Auf einem schiefen Bücherregal liegt ein Buch über die Kulturgeschichte der Physik. (folgende Bilderserie, ①) Es unterliegt seiner Gewichtskraft F_G , die im Schwerpunkt S des Buches eingezeichnet ist. Wie stark presst F_G das Buch auf das Regal, und wie stark treibt sie es in Richtung des abschüssigen Regals zum Rutschen an? Zur Antwort wollen wir F_G in zwei Komponenten F_H und F_N zerlegen, von denen F_H als Hangabtriebskraft nur in Richtung des möglichen Rutschens und F_N als Normalkraft nur senkrecht (wofür man hier auch „normal“ sagt) zum Regal wirken.



Da F_G als Resultierende beider Kräfte der Diagonalen im Kräfteparallelogramm aus F_H und F_N entsprechen muss, zeichnet man zunächst Parallelen in den betrachteten Richtungen durch die Enden des Kraft-Pfeiles von F_G ein, so dass sich das Parallelogramm ergibt. (Hier ist es sogar ein Rechteck. ②) Die Pfeile der gesuchten Kräfte ergeben sich als Seiten des Parallelogramms. (③)



Anderes Problem: Wie gezeichnet zieht man in waagerechter Richtung über einen Kraftmesser an zwei Gummi-Seilen, die an einer Wand befestigt sind. Der Kraftmesser zeigt 3 N an. Mit welcher Kraft sind die Seile gespannt? Zeichne die bekannte Zugkraft am Knotenpunkt in der Mitte schon einmal in rot ein! (1 cm : 1 N)

Da der Knotenpunkt in Ruhe bleibt, herrscht an ihm offensichtlich ein Kräftegleichgewicht. Es wirken an ihm drei Kräfte, nämlich die bekannte Zugkraft nach rechts und die beiden unbekannten Zugkräfte von den beiden Seilen her.

Noch unbekannt an den beiden anderen Kräften sind nur die Beträge, die Richtungen kennen wir: Ein Seil kann nur eine Zugkraft in Richtung des Seiles ausüben. Außerdem bekannt: Da ein Kräftegleichgewicht herrscht, müssen die beiden Seilkräfte zusammen eine Kraft ergeben, die der „roten“ Kraft genau entgegengesetzt gleich stark ist. Zeichne diese Gesamtkraft

der Seile in grau ein!

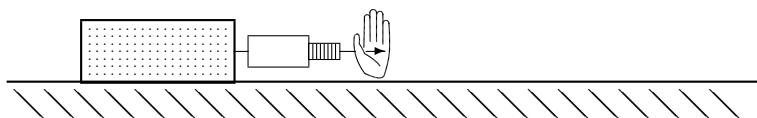
Der graue Pfeil muss eine Diagonale des Parallelogramms sein, das die Seilkräfte-Pfeile aufspannen. Zwei Seiten dieses Parallelogramms liegen auf den Seilen, die anderen erhält Du, indem Du durch die graue Pfeilspitze jeweils eine Parallele zu den beiden Seilen zeichnest. Dort, wo diese Parallelen die Seile schneiden, liegen die Spitzen der Seilkräfte-Pfeile. Zeichne diese Pfeile in zwei weiteren Farben ein. Aus ihren Längen kannst Du nun auch noch den Betrag der Seilkräfte errechnen. (Es sollten dabei etwa 2,50 N bzw. 1,01 N herauskommen.)

3.8 Reibung

3.8.1 Arten von Reibung

V 3.4 Holzklötz über den Tisch schleifen

Aufbau: Ein Holzklötz liegt auf dem Tisch. Er kann über einen Kraftmesser von Hand gezogen werden.



Durchführung: Wir ziehen mit langsam gesteigerter Kraft am Klotz.

Beobachtung:

Bei schwachem Zug bewegt sich der Klotz noch nicht. Offensichtlich hält eine Reibungskraft zwischen Klotz und Unterlage unserer Zugkraft das Gleichgewicht: Der Klotz **haftet**. Ebenso offensichtlich kann diese Haftkraft aber nur bis zu einem gewissen Wert mit der Zugkraft anwachsen.

Bei größerer Zugkraft wird der Klotz in Bewegung gesetzt: er **gleitet**. Sobald der Klotz gleitet, genügt eine

geringere Zugkraft, um das Gleiten mit konstanter Geschwindigkeit aufrecht zu erhalten, um also der Gleit-Reibungskraft das Gleichgewicht zu halten.

Man kann sich weiterhin gut vorstellen, dass ein runder „Klotz“, der **rollen** kann, mit noch geringerer Kraft in Bewegung gehalten werden könnte: Roll-Reibungskräfte sind also noch geringer als die bisher genannten.

Als Grund für das Auftreten von Reibungskräften kann man sich vorstellen, dass alle Oberflächen – auch sehr glatte wie z.B. die einer Glasplatte – mikroskopisch betrachtet doch mehr oder weniger rau sind. Solche rauhen Oberflächen verhaken sich bei gegenseitigem Kontakt ein wenig ineinander.

Gut vorstellbar ist, dass Folgendes einen Einfluss auf die Größe einer Reibungskraft hat:

- die Materialkombination, welcher Stoff auf welchem haftet, gleitet bzw. rollt
- die Größe der Berührungsfläche
- die Kraft, mit der die Körper aneinander gedrückt werden (im Folgenden immer die Gewichtskraft eines aufliegenden Körpers)
- beim Gleiten und Rollen: die Geschwindigkeit

3.8.2 Haftreibung

V 3.5 Haftreibung

Aufbau: wie zuvor, jedoch mit verschiedenen Materialien, verschiedenen Auflageflächen und verschieden schweren Klötzen.

Durchführung: Wir steigern die Zugkraft und lesen am Kraftmesser ab, bei welcher Kraft das Haften endet.

Ergebnis:

- Die maximale Haftkraft hängt in unseren Versuchen nicht merklich von der Auflagefläche ab.
- Sie ist etwa proportional zur Gewichtskraft des Klotzes. Den Proportionalitätsfaktor, $\frac{\text{Reibungskraft}}{\text{Gewichtskraft}}$, nennen wir **Haftzahl** f_H .
- Die Haftzahl hängt ab von der Materialkombination.

Die maximale Haftkraft beträgt bei einem Körper des Gewichtes F_G näherungsweise

$$F_H =$$

wobei die **Haftzahl** f_H von den aneinander haftenden Materialien abhängt.

zu ergänzen:

- Formelterm

3.8.3 Gleitreibung

In einem ähnlichen Versuch, bei dem die Zugkraft beim Gleiten mit konstanter Geschwindigkeit untersucht wird, erkennt man:

Die Gleitreibungskraft beträgt bei einem Körper des Gewichtes F_G weitgehend unabhängig von der Geschwindigkeit näherungsweise

$$F_R =$$

wobei die **Reibungszahl** f_R von den aneinander reibenden Materialien abhängt.

zu ergänzen:

- Formelterm

3.8.4 Rollreibung

Schon ein Verdacht? Ja, richtig geraten: Beim Rollen ist die Reibungskraft wieder näherungsweise nur proportional zur Gewichtskraft, mit einer eigenen, relativ kleinen, aber auch materialabhängigen Reibungszahl.

Druck

4.1 Stempeldruck

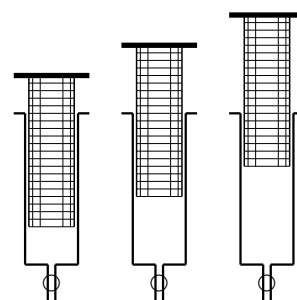
4.1.1 Kolbenprober

„Vorsicht! Behälter steht unter Druck!“ können wir auf mancher Spraydose lesen. Was versteht man eigentlich genau unter *Druck*?

Zur Untersuchung von Druck benutzen wir **Kolbenprober**. Ein Kolbenprober ist ein oben offener Glaszylinder, in dem ein Kolben auf und ab bewegt werden kann. Der Kolben ist genau passend zum Zylinder geschliffen, so dass er in jeder Stellung den Zylinder nach oben abdichtet. Durch den Boden des Zylinders führt ein Röhrchen mit Absperrhahn.

Zwischen dem Kolben und dem Boden des Zylinders befindet sich also ein absperrbarer Hohlraum. Eine darin eingeschlossene Flüssigkeit (oder Gas) kann man leicht *unter Druck setzen*, indem man von oben mit einer Kraft auf den Kolben drückt. Für unsere Untersuchungen ist nun folgende Überlegung wichtig:

Wenn der Kolben still steht, wird er sicher auch nicht beschleunigt. Wenn der Kolben nicht beschleunigt wird, herrscht an ihm Kräftegleichgewicht. Es muss dann am Kolben eine Kraft von unten nach oben wirken, die genauso stark ist, wie die Kraft, mit der wir von oben auf den Kolben drücken. Diese Kraft kann offenbar nur die unter Druck stehende Flüssigkeit (bzw. das Gas) von unten gegen den Kolben ausüben. Man nennt diese Kraft daher **Druckkraft**.



Bei unseren Untersuchungen sehen wir generell vom Gewicht des Kolbens ab.

Wenn wir die Kraft bestimmen, mit der der Kolben niedergedrückt wird, kennen wir demnach auch die Kraft, die aufgrund des Druckes im Hohlraum auf den Kolbenboden wirkt. Mit einem Kolbenprober können wir daher nicht nur Druck erzeugen, sondern auch noch

Druckkräfte messen. Den Kolben nennt man in diesem Zusammenhang auch Stempel. Den Druck, den wir mit Kolbenprobern erzeugen können, nennt man daher **Stempeldruck**. In Kürze wirst Du erfahren, dass es auch andere Ursachen von Druck gibt.

4.1.2 Allseitigkeit

Die vorige Überlegung zum Kräftegleichgewicht bleibt sinngemäß richtig, wenn der Kolbenprober nicht aufrecht gehalten wird. Drückt man selbst z.B. von links gegen den waagrecht gehaltenen Kolben, so drückt die Flüssigkeit natürlich von rechts ebenso stark dagegen. Die eingespernte Flüssigkeit kann auch sicher nicht den gläsernen Kolbenboden vom gläsernen Zylinderboden

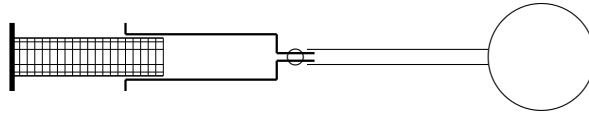
oder der gläsernen Zylinderwand unterscheiden. Da sie folglich auch nicht gezielt *nur* gegen den Kolbenboden drücken kann, müssen wir davon ausgehen, dass sie gleichermaßen gegen *alle* ihre Begrenzungswände drückt, also gegen alle Flächen, die die Flüssigkeit berührt. Der folgende Versuch bestätigt uns, dass die eingespernte Flüssigkeit tatsächlich nach allen Richtungen gleichermaßen drückt.

V 4.1 Spritzkugel

Aufbau: eine Hohlkugel mit Löchern, wassergefüllt, verbunden mit einem Kolbenprober

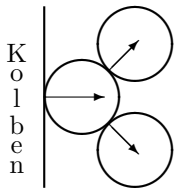
zu ergänzen:

- Löcher
- Spritzer



Durchführung:

Beobachtung:



Es mag zunächst überraschen, dass die Druckkräfte tatsächlich in alle Richtungen gleichermaßen wirken, obwohl doch z.B. der Kolben im Spritzkugelversuch nur nach rechts gedrückt wird. Wie kommt es, dass Wasser sogar genau entgegengesetzt nach links aus der Kugel gedrückt wird?

Wie jeden Körper muss man sich eine Wassermenge aus sehr vielen sehr winzigen Teilchen aufgebaut vorstellen. (In einem Milliliter Wasser befinden sich über 30 000 000 000 000 000 000 Wasserteilchen.) Diese Teilchen liegen dicht aneinander, können sich aber leicht gegeneinander verschieben. (Gerade deshalb ist Wasser flüssig.) Vereinfacht sind solche Teilchen hier am Rand als Kugeln gezeichnet.

Drückt man nun mit dem Kolben von links nach rechts gegen ein Teilchen, das gerade am Kolben anliegt, so wird dieses Teilchen diese Kraft auf seine Nachbarn übertragen. Allermeist werden diese Nachbarn gerade nicht ganz genau rechts vom ersten Teilchen liegen, sondern z.B. so wie gezeichnet etwas nach oben oder unten versetzt. Und schon kann die Kraft nur schräg vom ersten Teilchen auf seine Nachbarn übertragen werden.

Natürlich gibt es auch solche Richtungsänderungen, wenn die Nachbarteilchen die Kraft weiter auf ihre Nachbarn übertragen. Über einige Teilchen übertragen können die Kräfte demnach alle möglichen Richtungen haben. (Zumal für jede Kraft eines Teilchens A auf B ja auch noch eine entgegengesetzte Gegenkraft von B auf A ins Spiel kommt.)

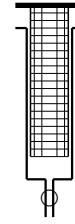
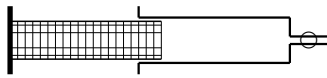
Die Teilchen an einer Wand drücken schließlich immer senkrecht zur Wandfläche auf die Wand, egal, wie die Wand gerichtet ist.

4.1.3 Definition

Um zu einer zahlenmäßigen Beschreibung des Druckes zu kommen, untersuchen wir nun die Kräfte auf Begrenzungsflächen genauer.

V 4.2 Fernbediente Hebebühne

Aufbau: Zwei Kolbenprober sind bei offenen Hähnen über einen langen Schlauch miteinander verbunden. Prober und Schlauch sind mit Wasser gefüllt. (Wir haben hier also wieder *eine* zusammenhängende Flüssigkeitsmenge eingesperrt.) Ein Kolbenprober wird geeignet befestigt, und es wird ein Gewichtsstück auf seinen Kolben gestellt. Der andere Prober wird vom Experimentator in der Hand gehalten.



Durchführung: Der Experimentator drückt an seinem Prober den Kolben ein.

Beobachtung:

(Zeichne neben die Prober dünn, wie die Prober nun aussehen!)

Ergebnis:

zu ergänzen:

- Schlauch
- Befestigung
- Hand
- Daumen
- Gewichtsstück

Man fragt sich nun, wie viel Kraft der Experimentator im vorigen Versuch aufwenden muss, um das Gewichtsstück gehoben zu halten. Braucht er mehr oder weniger als die Gewichtskraft der zu hebenden Last? Oder braucht er genau diese Gewichtskraft? Dies klärt der folgende Versuch.

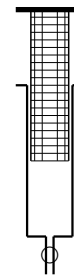
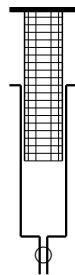
Da die Muskelkraft schwer zu messen ist, drückt nun anstelle des Experimentators ein weiteres Gewichtsstück den Kolben ein.

V 4.3 Messung der Kraftübertragung

Aufbau: wie zuvor, jedoch beide Prober aufrecht stehend und fest montiert.

zu ergänzen:

- Schlauch
- Befestigung
- Gewichtsstücke



Durchführung: Wir probieren aus, welches Gewichtsstück auf dem linken Prober ein Stück auf dem rechten Prober hoch halten kann, ihm also das Gleichgewicht halten kann.

Ergebnis:

Wenn auf baugleiche Kolbenprober gleich große Gewichtskräfte von oben wirken und dabei beide Kolben jeweils im Kräftegleichgewicht sind, muss in beiden Probern das Wasser ebenfalls mit der gleichen Kraft von unten gegen den Kolben drücken. Es liegt daher nahe zu sagen, dass in beiden Kolbenprobern der gleiche Druck herrscht.

Eigentlich gehen wir ja schon seit dem Spritzkugel-

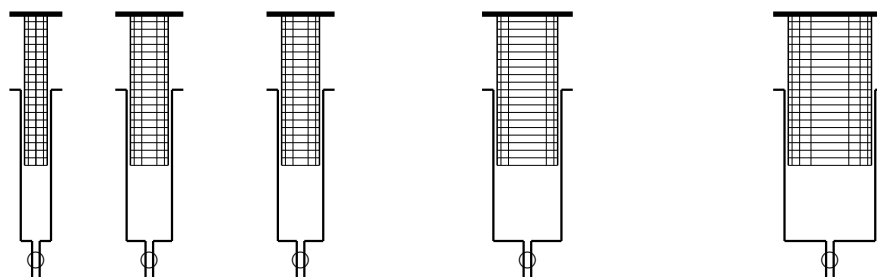
Versuch davon aus, dass eine eingespernte Flüssigkeit überall gleichermaßen drückt. Da in diesem Versuch beide Kolbenprober *eine* zusammenhängende Flüssigkeitsmenge begrenzen, war das Versuchsergebnis also schon zu erwarten. Während es aber bei der Spritzkugel nur *augenscheinlich danach ausgesehen* hat, dass die Flüssigkeit überall gleichermaßen drückt, konnten wir bei dem letzten Versuch ebendies *messen*.

?

Welche Druckkräfte wirken auf *verschieden* große Kolbenflächen?

V 4.4 Kraftübertragung bei verschiedenen Kolbenflächen

Aufbau: und **Durchführung:** wie zuvor, jedoch mit Kolbenprobern unterschiedlicher Kolbenfläche. Der zweite Prober wird mit 100 g belastet.



zu ergänzen:

- Schläuche
- Gewichtsstücke

Ergebnis: Die folgende Tabelle zeigt, bei welchen Gewichtsstücken auf welchen Probern Gleichgewicht herrscht:

Nr. d. Probers					
Kolbenfl. $A[\text{cm}^2]$					
Gew.kraft $F[\text{N}]$					

Auswertung: Wir vermuten ität zwischen F und A und bilden daher:

--	--	--	--	--	--

Ergebnis:

Der Quotient $\frac{F}{A}$ eignet sich gut zur Beschreibung des Druckes in einer Flüssigkeit oder einer Gasmenge.

Wenn nämlich $\frac{F}{A}$ größer ist, d.h. pro bestimmter Wandfläche mehr Druckkraft ausgeübt wird, sprechen wir sicher von einem höheren Druck. Wir definieren daher:

Der Buchstabe p wurde als Abkürzung von z.B. engl. *pressure* gewählt.

Der Quotient aus der Kraft F , die auf ein Stück Wandfläche ausgeübt wird, und dem Flächeninhalt A dieses Wandstückes heißt Druck p . Kurz:

$$p :=$$

Einheiten für Drücke sind demnach alle Quotienten aus einer Kraft- und einer Flächeneinheit. Als neue Einheit führt man ein:

$$1 \text{ Pascal} = 1 \text{ Pa} =$$

Außerdem darf noch die ältere Einheit **bar** benutzt werden.

$$1 \text{ bar} = 10 \frac{\text{N}}{\text{cm}^2}$$

Nachdem wir nun die Größe *Druck* definiert haben, dürfen wir damit auch als das gemeinsame Ergebnis aller unserer Druck-Versuche formulieren:

In einer zusammenhängenden Flüssigkeits- oder Gasmenge ist der Stempeldruck an allen Stellen gleich groß.

Damit erklärt sich nämlich, ...

- ...dass es bei der Spritzkugel aus allen Löchern gleich stark spritzt.
- ...dass bei gleichen Kolbenprobern gleiche Kräfte wirken.
- ...dass z.B. bei einer doppelt so großen Kolbenfläche die doppelte Kraft wirkt:
Es wirkt ja auf jede Hälfte der Kolbenfläche die gleiche, einfache Kraft.

4.1.4 Anwendungen

Die Tatsache, dass der Stempeldruck innerhalb einer Flüssigkeitsmenge überall gleich ist, lässt sich technisch ausnutzen, z.B. für Fernsteuerungen, Bremsen, Hebebühnen oder Pressen.

Entsprechende Anlagen bestehen im Wesentlichen aus zwei Bauteilen ähnlich wie Kolbenprober, die über einen Schlauch oder ein Rohr verbunden sind. Das System ist mit einer Flüssigkeit gefüllt. Vom griechischen Wort *hýdraulis* (=Wasserorgel) stammt das Adjektiv **hydraulisch** für solche Systeme, die mit dem Stempeldruck in Flüssigkeiten arbeiten. Moderne Systeme arbeiten allerdings nicht mit Wasser, sondern meist mit (Hydraulik-)Ölen. Bei Wasser würde es langfristig zu größeren Schäden (z.B. Rost) an metallenen Bauteilen kommen, bei Öl nicht. Außerdem schmiert das Öl noch die Reibungsflächen. Mit dem Einsatz hydraulischer Systeme verfolgt man zwei Ziele:

1. Es lassen sich Kräfte leicht über weite Entfernungen, auch über verwinkelte Wege, übertragen. Man muss ja nur einen Schlauch verlegen. Stangen, Hebel und Seilzüge sind bei weitem nicht so unproblematisch und flexibel zu installieren.
2. Durch geeignete Wahl der Kolbenflächen lassen sich Kräfte bei dieser Übertragung sogar vervielfachen.

Wir betrachten ein System aus zwei Kolben 1 und 2 mit den Kolbenflächen A_1 bzw. A_2 , die den Drücken p_1 bzw. p_2 ausgesetzt sind und daher die Kräfte F_1 und F_2 ausüben. (Zeichne diese Größen in die vorige Abbildung ein!) Es gelten nach Definition des Druckes

$$p_1 = \frac{F_1}{A_1} \quad \text{und} \quad p_2 = \frac{F_2}{A_2}$$

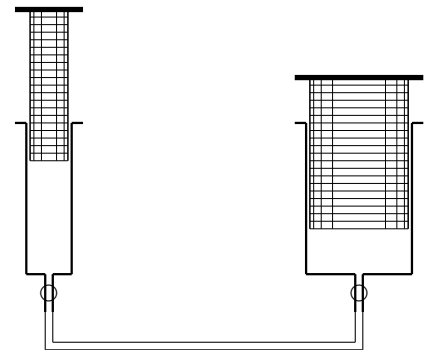
Dabei gilt

$$p_1 = p_2$$

weil es sich um Stempeldrücke innerhalb derselben Flüssigkeitsmenge handelt. Es gilt damit

$$\frac{F_1}{A_1} = \frac{F_2}{A_2} \quad \text{und folglich auch} \quad \frac{F_1}{F_2} = \frac{A_1}{A_2}$$

Die Kräfte verhalten sich also wie die Kolbenflächen: An der doppelten Fläche greift eine doppelte Kraft an, wie wir ja eigentlich bereits wissen.



Beispiel: Auf Kolben 2 mit $A_2 = 2000 \text{ cm}^2$ steht ein Auto mit der Gewichtskraft $F_2 = 8000 \text{ N}$ (Auto der Masse 800 kg). Kolben 1 habe eine Fläche $A_1 = 10 \text{ cm}^2$. Um das Auto auf Kolben 2 hoch zu halten, genügt daher an Kolben 1 eine Kraft

$$\begin{aligned} F_1 &= F_2 \cdot \frac{A_1}{A_2} \\ &= 8000 \text{ N} \cdot \frac{10 \text{ cm}^2}{2000 \text{ cm}^2} = 8000 \text{ N} \cdot \frac{1}{200} = 40 \text{ N} \end{aligned}$$

entsprechend der Gewichtskraft einer Masse von nur 4 kg , ein 200stel der Gewichtskraft des Autos!

Es lassen sich also mit beliebig kleinen Kräften beliebig große Kräfte ausüben. Prima Sache! Aber hat sie nicht doch einen Haken? Den hat sie!

Ein Auto fährt man auf eine Hebebühne, natürlich um es hochzuheben. Dazu wird in obigem Beispiel Kolben 1 um die Strecke s_1 in seinen Zylinder hineingedrückt. Genausoviel Öl, wie dabei aus seinem Zylinder herausgedrückt wird, strömt in den Zylinder von Kolben 2, wodurch Kolben 2 um eine Strecke s_2 angehoben wird.

Wird Kolben 1 z.B. um $s_1 = 50 \text{ cm}$ eingedrückt, verringert sich gemäß der Formel für das Volumen eines Zylinders das mit Öl gefüllte Volumen unter dem

Kolben um

$$V = A_1 \cdot s_1 = 10 \text{ cm}^2 \cdot 50 \text{ cm} = 500 \text{ cm}^3$$

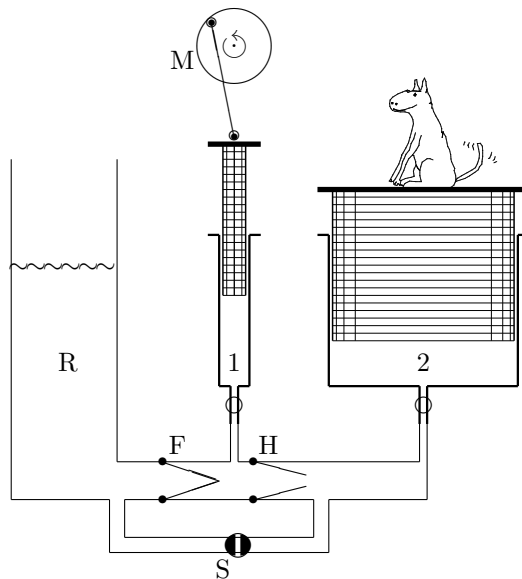
Diese 500 cm^3 Öl strömen unter Kolben 2, wo sie ein Volumen gemäß der Formel $V = A_2 \cdot s_2$ darstellen. Es ist daher

$$\begin{aligned} s_2 &= \frac{V}{A_2} = s_1 \cdot \frac{A_1}{A_2} \\ &= 50 \text{ cm} \cdot \frac{10 \text{ cm}^2}{2000 \text{ cm}^2} = 50 \text{ cm} \cdot \frac{1}{200} = 0,25 \text{ cm} \end{aligned}$$

Der Haken ist also, dass ich zwar das Auto mit nur 40 N heben kann, dafür aber meinen Kolben 1 um 50 cm bewegen muss, um das Auto um nur $2\frac{1}{2} \text{ mm}$ zu heben! Die 200-fache Kraftersparnis gegenüber direktem Hochheben erkaufe ich mir also mit einem 200-fachen Mehraufwand an zurückzulegenden Wegen.

Um das Auto um 2 m hochzuheben, müsste Kolben 1 also um 400 m gesenkt werden. Natürlich baut man keine so großen „Kolbenprober“. Hier hilft ein Trick:

Wenn in der folgend dargestellten Apparatur der kleine Kolben 1 vom Motor M nach unten gedrückt wird, wird das Ventil F zu gedrückt und Ventil H auf gedrückt, so dass Kolben 1 sein Öl unter Kolben 2 pumpt, wodurch dieser ein wenig angehoben wird.



Wird Kolben 1 wieder nach oben gezogen, drückt Kolben 2 das Ventil H zu. Aus dem Ölreservoir R saugt Kolben 1 durch das Ventil F wieder Öl in seinen Zylinder, der so wieder gefüllt wird. Nun kann Kolben 1 die neue Füllung wieder unter Kolben 2 drücken.

Zum Absenken des Kolbens 2 wird Kolben 1 arretiert und der Absperrhahn S geöffnet, so dass das Öl unter Kolben 2 zurück ins Reservoir fließen kann.

Natürlich lassen sich mit diesem Prinzip nicht nur Autos hochheben. (– sondern auch Hunde ;-). Ein paar Möglichkeiten außer dem Hochheben:

- In jedem PKW werden per Fußbremse die Bremsbeläge hydraulisch gegen die Radtrommeln oder die Bremscheiben gedrückt.
- In vielen starken Pressen werden zwei Formblöcke hydraulisch gegeneinander gedrückt, um z.B. ein dazwischenliegendes Blech zu einem Kotflügel zu formen.
- Greifarme von Robotern, Ausleger von Baggern und Ähnliches werden oft von hydraulisch ausfahrbaren Kolben bewegt.

4.1.5 Druckmessgeräte (Manometer)

Um einen Druck gemäß seiner Definition als Kraft pro Fläche zu messen, muss man messen, welche Kraft dieser Druck auf eine bestimmte Fläche ausübt. Bisher haben wir den Druck immer auf den Kolbenboden eines Kolbenprobers wirken lassen. Der Druckkraft auf die Bodenfläche haben wir mit Gewichtskräften das Gleichgewicht gehalten. So konnten wir mit den benötigten Gewichtskräften die Druckkräfte messen. Division der

Kraft durch die Kolbenfläche ergibt den Druck.

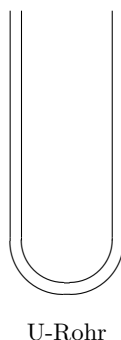
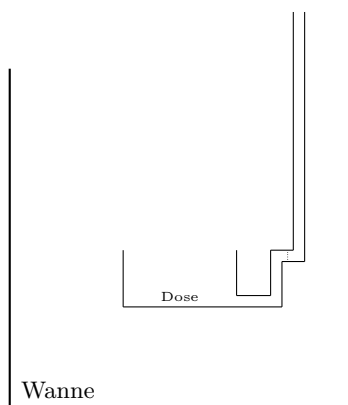
Statt Gewichtskräften kann man auch Federkräfte nutzen. Beispiel: Man lässt die Druckkraft, die auf eine bekannte Fläche wirkt, eine Feder dehnen. Mit zunehmender Dehnung wächst die Rückstellkraft der Feder und hält der Druckkraft schließlich das Gleichgewicht. An der Dehnung lässt sich die Rückstellkraft und damit auch die Druckkraft ablesen.

4.2 Schweredruck

4.2.1 Eigenschaften

Hast Du beim Tauchen schon einmal einen Druck gespürt?

Tauche doch beim nächsten Schwimmbad-Besuch einmal im Sprungbecken so tief Du kannst! (Vorsicht vor den Springern!)



Dieser Druck ist nicht ein Stempeldruck, weil ja offensichtlich schon die Wasserfüllung des Bades eingeschlossen ist, noch irgendjemand mit Stempeln darauf herumdrückt. Mangels Schwimmbad im Physiksaal untersuchen wir diesen Druck in einer Wanne. Wir benutzen dafür folgendes Messgerät, eine **Druckdose**: Das Wasser in der Wanne drückt die Membran je nach Druckkraft mehr oder weniger ein, und über die Luft im System wird die Wassermenge im U-Rohr mehr oder weniger in den rechten Schenkel verschoben. Je größer der Druck an der Membran, desto größer wird der Höhenunterschied der beiden Wassersäulen im U-Rohr.

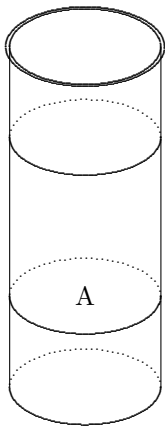
Um Druckkräfte aus verschiedenen Richtungen messen zu können, lässt sich die Dose um eine horizontale Achse drehen. Die Membran kann so auch nach unten, vorne oder hinten gedreht werden.

V 4.5 Druckdose

Aufbau: Wanne mit Wasser, Druckdose.

Durchführung: Wir bewegen die Druckdose in der Wanne umher, drehen auch an einigen Stellen die Dose in andere Stellungen

Beobachtung: /Ergebnis:



4.2.2 Entstehung

Dass der Druck bei jeder Richtung der Dose gleich stark auf die Membran wirkt, sind wir von Druck so gewohnt: Durch die große Anzahl an Wasserteilchen werden alle Kräfte in alle Richtungen gleichmäßig übertragen.

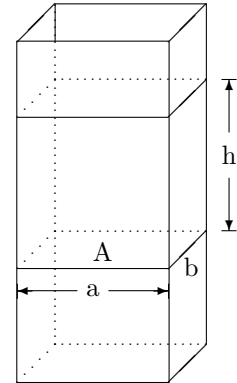
Anders als bisher ist jedoch, dass der Druck nicht an allen Stellen gleich ist, sondern mit der Tiefe zunimmt. Wie kommt das?

Betrachte nebenstehenden Glaszylinder, der mit Wasser gefüllt sein soll: Auf dem Teil der Wassersäule unterhalb der Querschnittsfläche A lastet der oberhalb von A liegende Teil der Wassersäule wie ein Stempel. Auf die Fläche A wirkt daher die Gewichtskraft dieser oberen Säule.

Mit dieser Überlegung lässt sich der Druck in der Tiefe von A berechnen. Um die Rechnung noch etwas zu erleichtern, betrachten wir eine Wassersäule mit rechteckigem Querschnitt und führen Bezeichnungen gemäß nebenstehender Abbildung ein. Die Tiefe, in der A liegt, wurde h genannt, weil dies auch die Höhe der Wassersäule über A ist. Das Volumen der Wassersäule über A ist das Volumen eines Quaders mit den Seitenlängen a , b und h . Also: $V =$.

Es gilt außerdem $A = a \cdot b$, so dass auch gilt:

$$V =$$



Diese Formel für V gilt für beliebige Querschnittsformen.

Die Wassersäule über A hat demnach eine Masse von

$$m = \rho \cdot V =$$

wobei mit ρ die Dichte von Wasser bezeichnet ist.

Die Gewichtskraft der Säule beträgt demnach

$$F = m \cdot g =$$

Diese Kraft verteilt sich auf die Fläche A. Es entsteht an dieser Fläche demnach der Druck

$$p = \frac{F}{A} =$$

In einer Flüssigkeit der Dichte ρ herrscht bei einem Ortsfaktor g in der Tiefe h aufgrund der Schwere der Flüssigkeit der sogenannte **Schweredruck**

$$p =$$

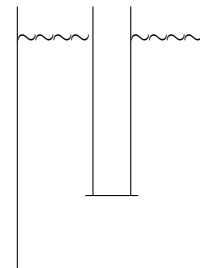
zu ergänzen:

- Formelterm

V 4.6 Wassersäule auf Glasplättchen

Aufbau: Eine Glasröhre wird unten mit einem losen leichten Glasplättchen dicht zugehalten und unter Wasser getaucht. (Der Schweredruck in der Eintauchtiefe drückt nun das Plättchen gegen die Röhre, so dass es nicht absinkt.) Bei mehreren Wiederholungen des Versuches wird verschieden tief eingetaucht.

Durchführung: Es wird vorsichtig Wasser in die Röhre gegossen.



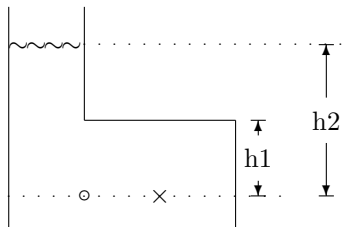
Beobachtung:

Erklärung:

4.2.3 Kommunizierende Röhren

Welcher Druck herrscht an der mit $\times$ markierten Stelle?

?



Ist es entscheidend, wie hoch über der Messstelle tatsächlich Wasser steht (h_1), oder ist es entscheidend, wie hoch über der Messstelle der Wasserspiegel des Gewässers liegt (h_2)?

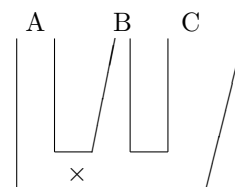
Wäre h_1 entscheidend, so müsste der Druck in der gepunkteten Wasserschicht an der Stelle $\circ$ von links nach rechts einen Sprung von $\rho \cdot g \cdot h_2$ zu dem kleineren Wert $\rho \cdot g \cdot h_1$ machen.

Ein Wassertropfen an dieser Stelle würden demnach von links stärker gedrückt als von rechts, würde sich entsprechend etwas nach rechts verschieben, so dass der Druck rechts steigen würde. Einen stabilen Zustand kann es also nur geben, wenn auch rechts von $\circ$ der höhere Druck entsprechend h_2 herrscht. Die Wasserteilchen übertragen den Druck also gleichmäßig in der ganzen Wasserschicht.

Ähnlich lässt sich auch die Zunahme des Druckes mit der Tiefe begründen: Auf ein Wassertropfen wirken in vertikaler Richtung 3 Kräfte: seine Gewichtskraft nach unten sowie die Druckkräfte von oben und von unten. In Ruhe bleiben kann das Tröpfchen nur bei Kräftegleichgewicht. Die Druckkraft von unten muss also um die Gewichtskraft des Wassertropfchens größer als die Druckkraft von oben sein.

Für den Schweredruck an einer Stelle ist nur die Tiefe unter dem Flüssigkeitsspiegel entscheidend.

Wir betrachten nun noch Gefäße, die oben mehrere Öffnungen haben, wie z.B. das am Rand gezeichnete mit den drei Schenkeln A, B und C. Man kann sich auch vorstellen, dass drei Gefäße A, B, und C unten über einen Schlauch verbunden sind. Wie hoch wird in diesen drei Schenkeln/Gefäßen das Wasser stehen, wenn man Wasser einfüllt?



Um den Druck z.B. an der Stelle $\times$ zu berechnen, muss man nach dem vorigen Merksatz nur berücksichtigen, wie hoch der Wasserspiegel über dieser Stelle steht. Problem: Es gibt hier in jedem Schenkel einen Wasserspiegel. Was ist nun *der* Wasserspiegel?

Da an der Stelle $\times$ natürlich nur *ein* Druck herrscht, darf es keine Rolle spielen, auf welchen der drei Spiegel man den Merksatz anwendet. Wenn der Merksatz so stimmt, müssen demnach alle Spiegel gleich hoch stehen.

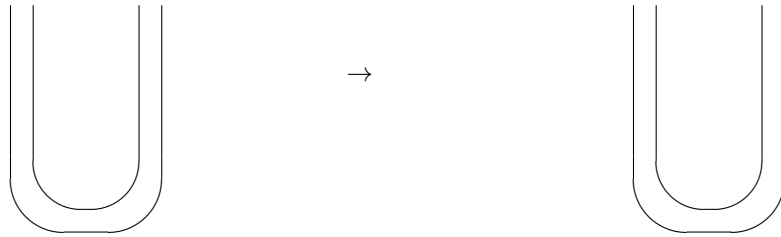
Stünde einer der Spiegel höher als die anderen, würde in den unteren Wasserschichten unter diesem Spiegel ein höherer Druck herrschen, als unter den anderen Spiegeln. Die Wasserteilchen zwischen diesen Druckgebieten würden sich dem Druckunterschied entsprechend verschieben und den Unterschied damit ausgleichen.

Das Prinzip kommunizierender Röhren: In verbundenen Gefäßen (in „kommunizierenden Röhren“) steht der Flüssigkeitsspiegel gleich hoch.

4.2.4 U-Rohr

V 4.7 Wasser und Öl in U-Rohr

Aufbau:



zu ergänzen:

- Wasser
- Öl
- $h_{L,R,O}$

Durchführung: Ein U-Rohr wird zunächst mit Wasser gefüllt.

Beobachtung: Wasserstände wie im linken Bild gezeichnet. (Nach Prinzip der kommunizierenden Röhren klar)

Durchführung: In den linken Schenkel wird zusätzlich etwas Öl gefüllt, das sich ja mit Wasser nicht mischt und wegen seiner geringeren Dichte ρ_O auf dem Wasser schwimmt.

Beobachtung: Wasser-/Ölstände wie im rechten Bild gezeichnet.

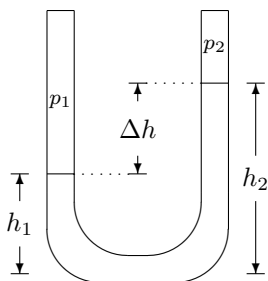
Erklärung: Am Boden des rechten Schenkels erzeugt die Wassersäule der Höhe h_R einen Schweredruck

$$p_R = \dots\dots\dots$$

Am Boden des linken Schenkels herrscht ein Druck, der von den Gewichten der Wasser- und der Ölsäule erzeugt wird: Zu dem Schweredruck, den die Ölsäule bereits an der Öl-Wasser-Grenzschicht erzeugt, kommt bis zum Boden des U-Rohres noch der Schweredruck, den die Wassersäule erzeugt: $p_L = \dots\dots\dots$

Solange $p_R \neq p_L$ ist, treibt der größere Druck die Flüssigkeiten in den Schenkel mit dem kleineren Druck. Dort steigt also der Füllstand und mit ihm auch der Druck. Zum Stillstand kann die Flüssigkeit nur kommen, wenn $p_R = p_L$ geworden ist:

$$\rho_W \cdot g \cdot h_R = \rho_W \cdot g \cdot h_L + \rho_O \cdot g \cdot h_O \quad \Leftrightarrow \quad \rho_W \cdot (h_R - h_L) = \rho_O \cdot h_O$$



Die Ölsäule im vorigen Versuch übt wie ein Kolben auf das darunter befindliche Wasser bereits einen Druck aus, zu dem dann im Wasser noch der Schweredruck des Wassers selbst kommt.

Ebenso könnte auch unter Druck stehende Luft über dem Wasserspiegel Druck auf die Wassersäule ausüben. Außerdem ist auch im rechten Schenkel recht, was links billig ist. Die nebenstehende Abbildung zeigt ein wassergefülltes U-Rohr, in dessen oben geschlossenen Schenkeln gewisse Luftmengen unter den Drücken p_1 bzw. p_2 eingeschlossen sind. (Färbe das Wasser bläulich!) Welches ist der größere Druck, p_1 oder p_2 ?

Auch dieses System kann in einem stabilen Zustand nur dann sein, wenn am Boden des Rohres links und rechts gleiche Drücke herrschen. Mit den Bezeichnungen aus der Abbildung bedeutet dies:

$$\begin{aligned} p_1 + \rho g h_1 &= p_2 + \rho g h_2 \\ p_1 - p_2 &= \rho g h_2 - \rho g h_1 \\ p_1 - p_2 &= \rho g \cdot (h_2 - h_1) \end{aligned}$$

$$\Delta p = \rho \cdot g \cdot \Delta h \quad \text{„U-Rohr-Formel“}$$

Aus dem leicht messbaren Höhenunterschied Δh kann also bei bekanntem ρg die Druckdifferenz Δp berechnet werden. Kennt man außerdem einen der beiden Drücke $p_{1,2}$, dann kann der andere berechnet werden. Markiert man (an zumindest einem Schenkel) die Wasserstandshöhen, die bestimmten Druckdifferenzen entsprechen, so hat man ein weiteres Druckmessgerät, ein **U-Rohr-Manometer**. ($\rightarrow$ Seite 41)

4.2.5 Luftdruck

Früher haben wir bei der Besprechung von „Dichte“ festgestellt, dass auch Luft Masse und folglich Gewicht hat. Demnach sollte auch in der Erdatmosphäre ein Schweredruck herrschen, der durch das Gewicht der Luft hervorgerufen wird. Wir leben auf dem Erdboden ja quasi am Grunde eines riesigen „Luftmeeres“.

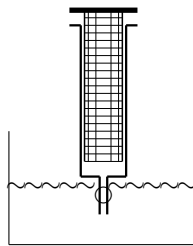
Lässt sich die Existenz dieses Luftdruckes experimentell belegen?

?

Paradoxerweise erkennt man die Existenz des Luftdruckes am besten in Versuchen, bei denen der Luftdruck stellenweise beseitigt oder zumindest verringert wird. Dies liegt einfach daran, dass der Luftdruck praktisch überall in unserer Umgebung wirkt, und folglich nur ein Wegnehmen dieses Druckes einen Unterschied zum gewohnten Normalfall herstellen kann.

V 4.8 Ansaugen

Aufbau: Ein vollständig eingeschobener Kolbenprober wird mit seiner Spitze bei geöffnetem Hahn in eine Wanne mit ausreichend viel Wasser getaucht.



Durchführung: Der Kolben wird nach oben ausgezogen.

Beobachtung:

Zur Erklärung dieser Beobachtung müssen wir angeben, warum das Wasser in den Hohlraum hinaufströmt, der beim Hochziehen des Kolbens im Prober entsteht. Offensichtlich muss eine Kraft nach oben auf das Wasser wirken, die stärker als die Gewichtskraft des Wassers ist.

Natürlich zieht der Experimentator den Kolben mit einer Kraft nach oben. Diese Kraft kommt jedoch nicht in Frage, weil sie nur auf den Kolben ausgeübt wird und nicht auf das Wasser übertragen werden kann. Wahrscheinlich war beim Hochziehen sogar noch ein wenig Luft zwischen dem Kolbenboden und dem aufströmenden Wasser, so dass Kolben und Wasser gar keinen direkten Kontakt hatten.

Und selbst bei direktem Kontakt könnten am Kolbenboden höchstens ein paar Tröpfchen Wasser haften und mit hoch gezogen werden, keinesfalls aber eine größere Menge Wasser, wie beobachtet. Davon kann man sich überzeugen, indem man versucht, mit dem Kolben alleine (aus der Hülse gezogen) Wasser aus der Wanne zu ziehen.

Es ist der Luftdruck, der das Wasser nach oben treibt: Der Luftdruck drückt *außerhalb* des Kolbenprobers auf die gesamte Wasseroberfläche in der Wanne.

Der beim Hochziehen des Kolbens entstehende Hohlraum *innerhalb* des Kolbens ist hingegen von der Umgebungsluft abgeschnitten. Von unbedeutenden Resten abgesehen enthält er gar keine Luft.

Wenn wir einmal annehmen, dass das Wasser nicht sofort in diesen Hohlraum hineingedrückt würde (z.B. weil wir den Hahn verschlossen hielten), wäre dieser Hohlraum absolut leer. Einen solchen völlig leeren Raum nennt man nach dem lateinischen Wort für *leer* ein **Vakuum**.

In diesem Vakuum kann natürlich auch kein Luft- oder sonstiger Druck herrschen. Auf die Wasseroberfläche in diesem Bereich kann dann auch keine Druckkraft wirken.

Drückt nun aber der Luftdruck außerhalb des Kolbenprobers von oben auf die Wasseroberfläche, und drückt bei geöffnetem Hahn innerhalb der Probers keine Kraft auf das Wasser, so wird das Wasser von der äußeren Luft in den Hohlraum hinein gedrückt. Das Wasser weicht dem äußeren Luftdruck in den Hohlraum aus. (Ähnlich wie in einem U-Rohr das Wasser in den rechten Schenkel hoch ausweicht, wenn man es im linken Schenkel nach unten drückt.)

Eigentlich sind alle Versuche zum Luftdruck so zu deuten. Es sind praktisch alle Vorgänge, die etwas mit Saugen zu tun haben: Zum Beispiel sind schon das Einatmen von Luft, das Trinken mit Strohalm oder Staubsaugen nur aufgrund des Luftdruckes möglich.

Die Schwierigkeit besteht höchstens darin, dass wir zwar im Alltag von einem „Sog“ sprechen, in allen diesen Fällen aber physikalisch gesehen keine anziehenden Sog-Kräfte existieren.

Wenn Du Luft einatmest, gibt es keine Kräfte, mit denen Du die Luft vor Deiner Nase in Dich hineinziehst, sondern der Luftdruck *drückt* diese Luft in Dich hinein, sobald Du Deine Lungen mit Muskelkraft (gegen

den äußeren Luftdruck) ausdehnst und damit Platz für zusätzliche Luft schaffst.

Ein Staubsauger *zieht* Luft und mit ihr Staubkörnchen nicht in sich hinein, sondern der Luftdruck *drückt* sie von vorne in das Gerät, sobald der Motor des Staubsaugers Luft hinten aus dem Gerät (gegen den äußeren Luftdruck) herausdrückt und damit Platz für neue Luft von vorne schafft.

Durch einen Strohalm *ziehst* Du auch keine Limonade in Deinen Mund, sondern der Luftdruck *drückt* von der anderen Seite des Halmes aus die Limo in Deinen Mund, sobald Du wie beim Einatmen Platz dafür schaffst.

Ähnliche Umkehrungen der Sichtweise kennst Du sicher bereits:

Loch im Papier=fehlendes Papier, umgeben von Papier

Schatten an der Wand=fehlendes Licht, umgeben von Licht

Nun kommt neu hinzu: Sog zu einer Seite hin

=fehlender Druck gegenüber Luftdruck von der anderen Seite her

Demnächst wird als weiteres Beispiel folgen: Kälte=fehlende Wärme

Wie groß ist der Luftdruck?

Wir haben vor kurzem als Formel für den Schweredruck $p = \rho \cdot g \cdot h$ kennengelernt. Für eine Flüssigkeit an der atmosphärischen Luft kommt, wie wir nun wissen, noch der Luftdruck dazu. Es lastet ja auf jeder Flüssigkeitsschicht nicht nur die Flüssigkeitssäule darüber, sondern auch die Luft über dem Flüssigkeitsspiegel.

Wenn wir den Luftdruck einmal p_0 nennen, herrscht also in der Tiefe h der Druck $p = p_0 + \rho \cdot g \cdot h$. Insbesondere herrscht an einer Wasseroberfläche, d.h. bei Tiefe $h = 0$ der Druck p_0 .

Hält man von einer Stelle der Oberfläche den Luftdruck fern (wie im vorigen Versuch innerhalb des Probers), so wird dort Wasser nach oben gedrückt. Das emporsteigende Wasser erzeugt dabei jedoch einen wachsenden Schweredruck unter sich. Nur solange dieser Schweredruck kleiner als der Luftdruck ist, kann die Luft weiterhin Wasser nach oben drücken. Das Wasser kann also höchstens so hoch gedrückt werden, dass sein Schweredruck gleich dem Luftdruck ist.

Wenn man also misst, wie hoch der Luftdruck Wasser halten kann, lässt sich der Luftdruck berechnen. Am einfachsten füllt man dazu einen durchsichtigen Schlauch mit Wasser und verschließt ein Ende so, dass an diesem Ende keine Luft ist. Das offene Ende hält man nach oben, das verschlossene hebt man so weit hoch, dass das Wasser darin einen Hohlraum freigibt. Dann sieht die Anordnung wie ein U-Rohr so aus, wie hier am Rand gezeichnet.

Über dem Wasserspiegel im linken Schenkel herrscht Vakuum und damit der Druck 0. Über dem rechten Wasserspiegel herrscht der Luftdruck p_0 . Nach der „U-Rohr-Formel“ auf Seite 44 gilt also mit $p_1 = 0$ und $p_2 = p_0$

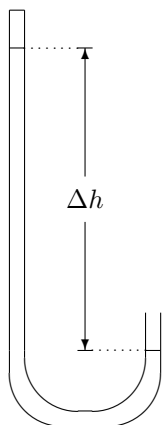
$$p_0 = \rho \cdot g \cdot \Delta h$$

Bei diesem Versuch misst man etwa $\Delta h = 10$ m. Man erhält grob gerechnet:

$$\begin{aligned} p_0 &= \rho_{\text{Wasser}} \cdot g \cdot \Delta h = 1 \frac{\text{g}}{\text{cm}^3} \cdot 10 \frac{\text{N}}{\text{kg}} \cdot 10 \text{ m} \\ &= \frac{0,001 \text{ kg}}{(0,01 \text{ m})^3} \cdot 10 \frac{\text{N}}{\text{kg}} \cdot 10 \text{ m} = \frac{0,001}{0,000001} \cdot 10 \cdot 10 \cdot \frac{\text{kg}}{\text{m}^3} \cdot \frac{\text{N}}{\text{kg}} \cdot \text{m} \\ p_0 &= 100\,000 \frac{\text{N}}{\text{m}^2} = 100\,000 \text{ Pa} = 1 \text{ bar} \end{aligned}$$

Man beachte: Auf jeden cm^2 einer Fläche drückt der Luftdruck mit etwa 10 N, was etwa dem Gewicht eines Körpers der Masse 1 kg entspricht!

?



Auf einen Menschen mit einer Körperoberfläche von 1 m^2 wirkt damit insgesamt eine Kraft von N !
Dass wir diese enormen Kräfte nicht bemerken, liegt an zweierlei:

Erstens wirken diese Kräfte (fast) von allen Seiten her auf unseren Körper, so dass der Körper als Ganzes praktisch keine resultierende Kraft erfährt und folglich nicht beschleunigt wird.

Zweitens: Zellen eines Organismus sind durch den äußeren Luftdruck so zusammengedrückt, dass im Zellinneren ein Druck derselben Größenordnung herrscht. (Siehe dazu die Überlegungen auf Seite 35 mit Zellhäuten statt Kolben) Da dies auch schon während der gesamten Evolution (=Entwicklung aller Lebensformen) so war, sind alle Lebewesen dem äußeren Druck in ihrem üblichen Lebensraum hervorragend angepasst.

Probleme gibt es erst dann, wenn der äußere Druck rasch vom gewohnten Druck stark abweicht, z.B. beim Abtauchen in große Wassertiefe, dem Auftauchen oder beim Erklimmen hoher Berge. Einige mit dem raschen Druckwechsel einhergehende Veränderungen im Körper

werden nicht schnell genug ausgeglichen und führen zu verschiedenen Funktionsstörungen bis hin zum Tode.

So spielt häufig nur die Abweichung vom Normaldruck eine Rolle: Auch in einem platten Fahrradreifen herrscht der gewöhnliche Luftdruck von etwa 1 bar, d.h. auf jeden cm^2 der Innenseite des Schlauches drücken 10 N. Da aber von außen genauso starke Kräfte wirken, ist der Schlauch völlig schlaff. Erst, wenn der Innendruck höher als der Außendruck wird, wird der Schlauch prall.

Man spricht in diesem Fall von einem **Überdruck** im Schlauch: Herrschen im Schlauch z.B. 1,3 bar, so ist dies ein Überdruck von 0,3 bar gegenüber dem äußeren Luftdruck von 1 bar. Nur diese 0,3 bar muss der Schlauch im Endeffekt aushalten, nur diese 0,3 bar machen die Härte des Reifens aus.

Bei Druckangaben muss immer klar sein, ob nur die Abweichung vom Normaldruck als Über- oder Unterdruck angegeben ist, oder der tatsächlich herrschende Absolutdruck.

In Fällen wie diesem benutzt man häufig die Adjektive **relativ** und **absolut**. *Relativ zu ...* heißt übersetzt *bezogen auf ...*. *Absolut* heißt *losgelöst* (von irgendwelchen Vergleichsgrößen, wie hier vom Luftdruck).

Der genaue Wert des Luftdruckes hängt unter anderem von der Höhenlage der Messstelle ab. Wie jeder Schweredruck nimmt auch der Luftdruck mit der Höhenlage ab: Je höher die Messstelle gelegen ist, desto weniger Luft lastet ja auf ihr.

Für Luft gilt die uns bekannte Formel für den Schweredruck jedoch nicht: Der Luftdruck nimmt nicht proportional zur Höhe der Luftsäule über der Messstelle zu. Da man Luft leicht zusammendrücken kann, und da der Luftdruck in der Atmosphäre unten größer als oben ist, ist die Luft in unteren Schichten der Atmosphäre dichter als in den darüberliegenden. Beim Eintauchen in die Atmosphäre nimmt der Luftdruck daher nicht gleichmäßig mit der Tiefe, sondern immer stärker zu.

Weiterhin hängt der Luftdruck mit dem Wetter zusammen. Je nach Wetter ist Luft verschieden feucht, verschieden warm und schon von daher auch verschie-

den dicht. Wegen der starken Zusammenhänge zwischen Wetter und Luftdruck spricht man in Wetterberichten so häufig von Hoch- und Tiefdruckgebieten.

Messgeräte für den Luftdruck nennt man **Barometer**. (von griechisch *barys*=schwer)

Welche enormen Kräfte der gewöhnliche Luftdruck ausüben kann, zeigte 1654 der Magdeburger Politiker und Naturwissenschaftler Otto von Guericke dem Reichstag zu Regensburg mit den „Magdeburger Halbkugeln“:

Zwei metallene Halbkugeln, die luftdicht aufeinander passen, bilden eine Kugel, die über einen Hahn luftleer gepumpt werden kann. In diesem luftleeren Zustand schaffen es zweimal 8 Pferde in einer Art „Tauziehen“ nicht, die Halbkugeln zu trennen, so stark werden sie vom äußeren Luftdruck zusammengedrückt. Lässt man wieder Luft in die Kugel strömen, fallen die Halbkugeln von selbst auseinander.

Im Mittel beträgt der Luftdruck auf Meereshöhe etwa 1013 hPa. Er schwankt etwa um 50 hPa.

4.3 Auftrieb

Wenn z.B. ein Styropor-Körper in einer Flüssigkeit nach oben steigt, muss offensichtlich eine Kraft nach oben auf ihn wirken, die stärker als seine Gewichtskraft ist. Man nennt sie die **Auftriebskraft**.

?

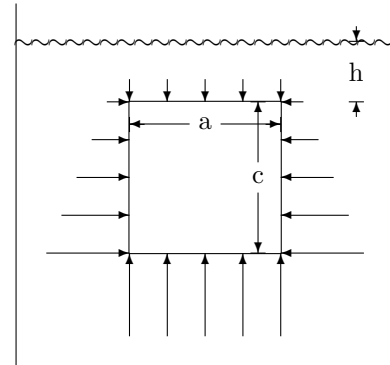
Welche Ursache hat die Auftriebskraft?

Wir betrachten einen Quader mit den Abmessungen a, b und c , der sich um h tief unter der Oberfläche einer Flüssigkeit befindet. Auf die 6 rechteckigen Begrenzungsflächen des Quaders wirken je nach Tiefe unterschiedliche Druckkräfte. Diese Kräfte sind im Bild

durch Pfeile angedeutet. (Entgegen der Regel, dass ein Kraftpfeil am Angriffspunkt beginnt, enden diese Pfeile mit ihrer Spitze am Quader, um verwirrende Überschneidungen zu vermeiden.)

Die horizontal wirkenden Kräfte, die seitlich am Quader angreifen, brauchen wir nicht weiter zu berücksichtigen: Erstens wollen wir ja die nur vertikal wirkende Auftriebskraft untersuchen. Zweitens heben sich die von links und rechts (sowie die von vorn und hinten) in gleicher Tiefe wirkenden Kräfte sowieso jeweils auf. Es bleiben also nur noch die Kräfte auf die obere und die untere Deckseite zu untersuchen, die jeweils den Flächeninhalt $A = a \cdot b$ haben. Im Folgenden seien noch ρ die Dichte der Flüssigkeit und g der Ortsfaktor.

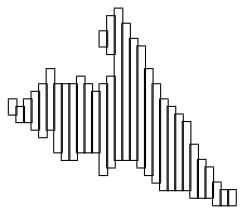
Oben herrscht dann der Druck $p_o = \rho \cdot g \cdot h$, so dass die obere Deckfläche die Kraft $F_o = A \cdot p_o = a \cdot b \cdot \rho \cdot g \cdot h$ erfährt. An der unteren Deckfläche herrscht der Druck $p_u = \rho \cdot g \cdot (h + c)$, so dass diese Fläche die Kraft $F_u = a \cdot b \cdot \rho \cdot g \cdot (h + c)$ erfährt. Diese Kräfte heben sich nur teilweise auf. Es bleibt als Auftriebskraft nach oben ein Überschuss vom Betrage $F_a = F_u - F_o = a \cdot b \cdot \rho \cdot g \cdot c$.



Berücksichtigt man, dass das Volumen V des Quaders gerade $V = a \cdot b \cdot c$ ist, so lässt sich schreiben $F_a = \rho \cdot V \cdot g$.

Es ist weiterhin $\rho \cdot V \cdot g$ die Masse der Flüssigkeit, die vom Quader verdrängt wird. $\rho \cdot V \cdot g$ – und damit die Auftriebskraft F_a – ist demnach das Gewicht der verdrängten Flüssigkeitsmenge.

Dieses Ergebnis gilt nicht nur für so ordentlich ausgerichtete Quader, wie es der Rechnung zugrunde gelegt wurde, sondern für jeden Körper. Zur theoretischen Herleitung denkt man sich dann den beliebig geformten Körper aus Quadern (wie aufrecht stehende extra dünne und möglichst lange pommes frites) zusammengesetzt.



Satz des Archimedes: Jeder Körper, der in eine Flüssigkeit getaucht ist, erfährt eine Auftriebskraft nach oben, deren Betrag gleich dem Gewicht der verdrängten Flüssigkeit ist.

Es gibt ebenso einen Auftrieb in Gasen.

So wird klar, warum Körper je nach ihrer Dichte schwimmen, schweben oder sinken: Ein Körper der Dichte ρ_K und Volumen V erfährt eine Gewichtskraft von $\rho_K \cdot V \cdot g$ nach unten. In einer Flüssigkeit mit der Dichte ρ_F erfährt er aber auch eine Auftriebskraft $\rho_F \cdot V \cdot g$ nach oben. Je nachdem, welche Dichte größer ist, überwiegt also die entsprechende, mit der größeren Dichte zusammenhängende Kraft. Überwiegt wegen $\rho_K > \rho_F$ die Gewichtskraft, wird der Körper nach unten beschleunigt. Überwiegt wegen $\rho_K < \rho_F$ die Auftriebskraft, ergibt sich eine nach oben wirkende resultierende Kraft, die den Körper nach oben treibt. Diese resultierende Kraft nach oben verschwindet, sobald der Körper so weit aus der Flüssigkeit ragt, dass er nur noch

so viel Flüssigkeit verdrängt, wie er selbst wiegt.

Mit ähnlichen Überlegungen lassen sich alle Aufgaben zum Auftrieb lösen, bei denen ein Körper in einem stabilen Zustand ist, genauer: nicht nach oben oder unten beschleunigt wird. Beispiele: Ein schwimmendes Schiff; ein unter Wasser festgehaltenes Stück Styropor; ein Stein, der an einer Schnur hängend ins Wasser gehalten wird.

In allen diesen Fällen muss die Summe aller nach oben wirkenden Kräfte (Auftriebskräfte, Haltekräfte) gleich sein der Summe aller nach unten wirkenden Kräfte (Gewichtskräfte, Haltekräfte): Kräftegleichgewicht!

Herrscht kein Kräftegleichgewicht, so wird der Körper in Richtung der resultierenden Kraft beschleunigt, gemäß $F = m \cdot a$.

Kraftwandler

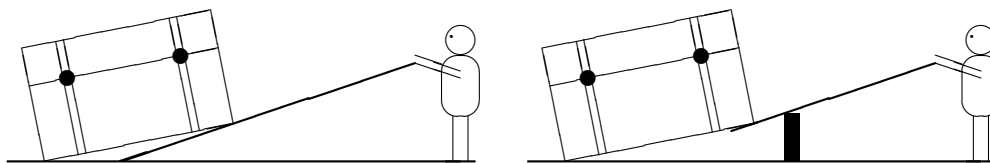
5.1 Hebel

5.1.1 Hebel in Stangenform

Was hat man davon, *am längeren Hebel zu sitzen*?

?

Hebel sind in der einfachsten Ausführung Stangen, die an einer Stelle drehbar gelagert sind. Man setzt Hebel ein, um z.B. eine schwere Schatzkiste anzuheben, damit man ein Wägelchen darunter schieben kann. Falls man keinen passenden Schlüssel findet, hilft ein Hebel auch dabei, die Schlösser der Kiste aufzubrechen. In diesen Beispielen kann der Schatzfinder mit Hebeln wesentlich größere Kräfte ausüben als er es ohne Hilfsmittel könnte. Seine Schatzkiste kann er auf zwei Arten mit einer Stange anheben:



zu ergänzen:

- Drehpunkte
- Angriffspunkte

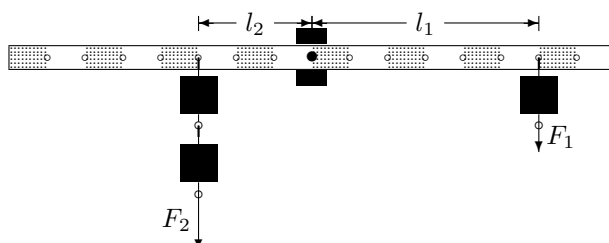
In jedem Fall gibt es eine Stelle D, an der die Stange fest, aber drehbar aufliegt, eine Stelle K, an der das Männchen Kraft ausübt, und eine Stelle L, auf der die Schatzkiste lastet. Je nachdem, ob die beiden Angriffspunkte K und L vom Drehpunkt D aus gesehen auf derselben Seite der Stange liegen oder auf verschiedenen, nennt man den Hebel **einseitig** bzw. **zweiseitig**. Die

Entfernung zwischen K und D wird oft **Kraftarm**, die zwischen L und D **Lastarm** genannt. Der Oberbegriff für beides ist **Hebelarm**. Bewegt sich der Hebel nicht, so heben sich offensichtlich die Einflüsse aller am Hebel angreifenden Kräfte gerade auf: der Hebel befindet sich *im Gleichgewicht*. Im Folgenden untersuchen wir, wann dies der Fall ist.

Zweiseitiger Hebel

V 5.1 Gleichgewicht am zweiseitigen Hebel

Aufbau:



Durchführung: Wir hängen rechts und links vom Drehpunkt mit den Hebelarmen l_1 bzw. l_2 Gewichtsstücke mit Gewichten F_1 bzw. F_2 an. Wir suchen dabei Kombinationen von Hebelarmen und Gewichten, bei denen der Hebel im Gleichgewicht ist.

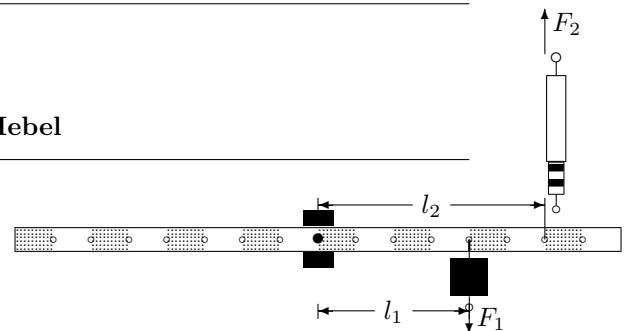
Beobachtung: /Auswertung: Gleichgewicht herrscht zum Beispiel bei folgenden Kombinationen:

l_1										
F_1										
l_2										
F_2										

Ergebnis:

Einseitiger Hebel

V 5.2 Gleichgewicht am einseitigen Hebel



Aufbau:

Durchführung: Wir stellen bei verschiedenen Werten für l_1 , F_1 und l_2 fest, wie stark man über den Kraftmesser am Hebel nach oben ziehen muss, um ihn im Gleichgewicht zu halten.

Beobachtung: /Auswertung: Gleichgewicht herrscht zum Beispiel bei folgenden Kombinationen:

l_1										
F_1										
l_2										
F_2										

Ergebnis:

5.1.2 Drehmoment

Die rechnerischen Ergebnisse unserer Untersuchungen am ein- und am zweiseitigen Hebel unterscheiden sich nicht. Die Unterschiede zwischen beiden Hebeltypen liegen ja auch nur in den Richtungen des Hebelarmes l_2 und der Kraft F_2 : Im Versuch zum zweiseitigen Hebel verlief der Hebelarm vom Drehpunkt aus nach links und die Kraft wirkte nach unten. Beim einseitigen Versuch wiesen der Hebelarm nach rechts und die Kraft nach oben.

Aber auch diese Unterschiede verschwinden, wenn

Zeigen Kraft und Hebelarm z.B. in dieselbe Richtung, so ergibt sich überhaupt keine Drehwirkung. Davon überzeugt man sich leicht, indem man an der Stange in Richtung der Stange zieht. Bei unseren bisherigen Versuchen wirkte die Kraft immer senkrecht zum Hebelarm.

Das gemeinsame Versuchsergebnis lässt sich nun noch einfacher formulieren, wenn man das darin vorkommende Produkt aus Kraft und Hebelarm als eine neue physikalische Größe benennt:

Das Produkt aus einer senkrecht am Hebel angreifenden Kraft F und ihrem Hebelarm l heißt **Drehmoment** M . Kurz:

$$M :=$$

Zugehörige Einheit:

$$1 \text{ N} \cdot 1 \text{ m} = 1 \text{ Nm} = 1 \text{ „Newtonmeter“}$$

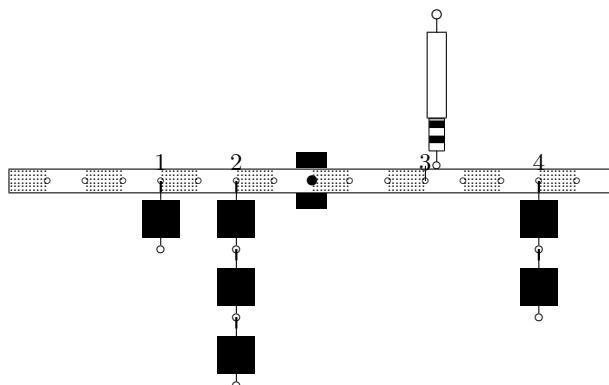
Man kann weiter festlegen, dass Drehmomente, die alleine eine Drehung z.B. gegen den Uhrzeigersinn verursachen, positiv sind, und Drehmomente für Drehungen in entgegengesetzter Richtung negativ. Eine andere Möglichkeit, den Drehsinn zu beschreiben, sind entsprechende Pfeilchen wie in folgenden Beispielen: $\curvearrowright$, $\curvearrowleft$, 20 Nm . Es liegt folgende Vermutung nahe, die für mehr als zwei Kräfte allerdings noch geprüft werden muss:

Hebelgesetz: Ein Hebel ist im Gleichgewicht, wenn die Summe aller im Uhrzeigersinn wirkenden Drehmomente gleich der Summe aller entgegengesetzt wirkenden Drehmomente ist.

Ebenso, wie ein Körper im Kräftegleichgewicht nicht beschleunigt wird, d.h. seine Geschwindigkeit nicht ändert, ändert ein Hebel im Drehmomenten-Gleichgewicht seine Drehgeschwindigkeit nicht.

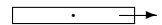
V 5.3 Hebel mit mehr als zwei Kräften

Aufbau:



man die Richtungen über einen Drehsinn beschreibt: In beiden letzten Versuchen waren l_2 und F_2 so gerichtet, dass F_2 alleine den Hebel *gegen* den Uhrzeigersinn gedreht hätte. In beiden Fällen waren jedoch l_1 und F_1 ebenso einflussreich auf eine Drehung *im* Uhrzeigersinn ausgerichtet.

Man beachte, dass die Richtung der Drehwirkung einer Kraft nicht nur von der Richtung der Kraft, sondern auch von der Richtung des zugehörigen Hebelarmes abhängt.



zu ergänzen:

- Formel

Durchführung: Wir belasten den Hebel an mehreren Stellen, z.B. wie abgebildet. Wir stellen dann fest, mit welcher Kraft an welcher Stelle man ihn im Gleichgewicht hält.

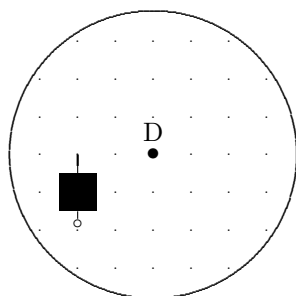
Messwerte: / Auswertung:

Nr.	1	2	3	4	5
Kraft					
Hebelarm					
Drehmoment					

Summe aller Drehmomente: $\curvearrowright$ $\curvearrowleft$

5.1.3 Beliebige Körper als Hebel

Hebel müssen nicht stangenförmig sein. Tatsächlich kann ein Hebel jede Form haben. Für die weiteren Untersuchungen ist es günstig, eine Scheibe als Hebel zu benutzen. Die Scheibe wird mittig auf eine Drehachse montiert.



V 5.4 Drehmomente an einer Scheibe

Aufbau: siehe Rand

Durchführung: Wir hängen an die Scheibe einen Gewichtsstein wie abgebildet und suchen alle Stellen der Scheibe, an denen ein weiterer Gewichtsstein angehängen werden kann um Gleichgewicht herzustellen.

Ergebnis: Es sind die im Bild markierten Stellen.

zu ergänzen:

- gefundene Stellen

Offensichtlich ändert sich das Drehmoment nicht, wenn man den Angriffspunkt einer Kraft in Richtung der Kraft verschiebt. Die Gerade, auf der der Angriffspunkt so verschoben werden kann, heißt **Wirklinie** der Kraft.

An obigem Versuch erkennen wir daher, dass man zur Berechnung eines Drehmoments *nicht* einfach den Abstand des Angriffspunktes von der Drehachse als Hebelarm verstehen darf.

Zur Berechnung von Drehmomenten muss die Kraft multipliziert werden mit dem Abstand ihrer Wirklinie von der Drehachse.

Denke oder zeichne Dir also einfach immer nur den Kraftpfeil nach beiden Seiten zu einer Geraden – der Wirklinie – verlängert und nimm deren Abstand zur Drehachse als Hebelarm. Dieser Abstand ist natürlich die Länge derjenigen Verbindungslinie zwischen Drehachse und Wirklinie, die senkrecht zu beiden steht.

Die Form des Hebels spielt überhaupt keine Rolle! Es spielt insbesondere auch keine Rolle, ob die Wirklinie und ihre senkrechte Verbindung zur Drehachse innerhalb des Hebelkörpers verlaufen oder nicht.

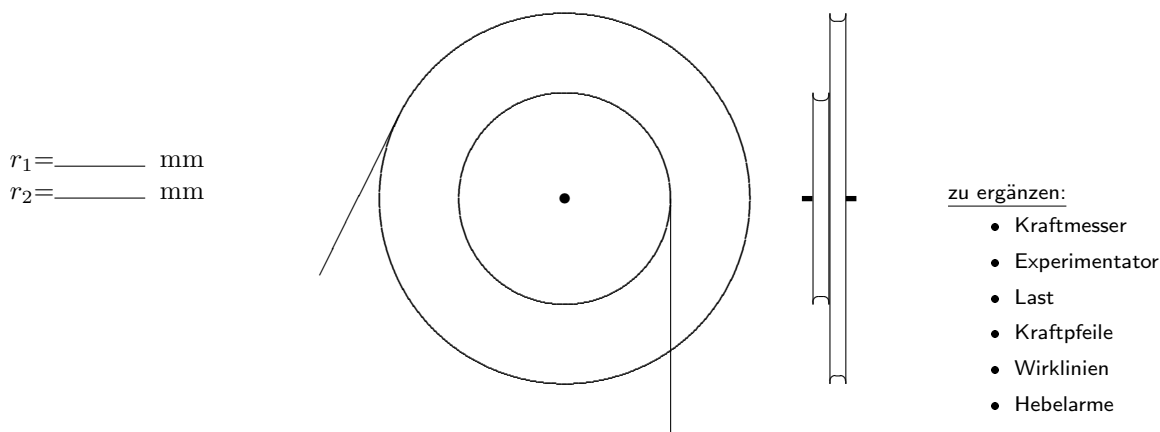
Alternative Berechnung eines Drehmoments: Multipliziere die Länge der Verbindungslinie zwischen Angriffspunkt und Drehachse mit der Komponente der Kraft, die senkrecht auf dieser Verbindungslinie steht.

5.1.4 Wellräder

... sind eine technisch bedeutsame Form von Hebeln. Ein Wellrad besteht grundsätzlich aus zwei oder mehr Rädern, die auf einer gemeinsamen Achse montiert und fest so miteinander verbunden sind, dass sie sich stets nur gemeinsam drehen können. Ein Beispiel für ein Wellrad sind z.B. am Hinterrad eines Fahrrades mit Kettenschaltung die unterschiedlich großen Ritzel (=Zahnräder) für die verschiedenen Gänge.

V 5.5 Kraftwandlung am Wellrad

Aufbau: /Durchführung: Zwei fest aufeinander montierte Scheiben sind in ihrer Mitte (•) drehbar gelagert. (rechts im Bild: Seitenansicht) Am Rand jeder Scheibe ist eine Schnur befestigt und um die Scheibe gewickelt. An einer Schnur hängt eine Last, an der anderen zieht der Experimentator über einen Kraftmesser.



Messwerte:

F_1									
F_2									

Ergebnis: 1) Die Hebelarme sind _____

2) Das Hebelgesetz hat hier die Form _____

Technische Betrachtung: Bei einem Wellrad überträgt die starre Verbindung der beiden Räder das Drehmoment, das man an einem Rad ausübt, auf das andere Rad. Sind die beiden Räder nicht – wie im vorigen Bild – direkt miteinander verbunden, sondern über eine Stange, so nennt man diese Stange eine **Welle**. (Daher also die Bezeichnung *Well-rad*!)

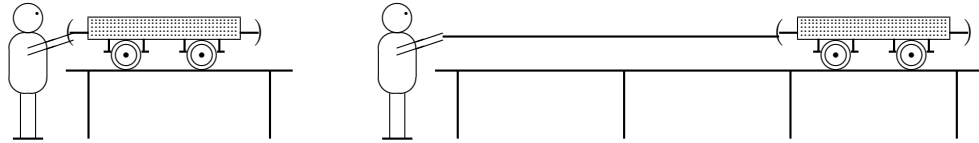
Bei manchen Aufgaben musst Du sorgfältig unterscheiden: Seile, Ketten und ineinandergreifende Zahnräder übertragen eine *Kraft* unverändert von einem Bauteil auf ein anderes, wobei sich das Drehmoment ändern kann. Wellräder und sonstige Hebel bestehen aus starr miteinander verbundenen Teilen, die ein *Drehmoment* unverändert übertragen, jedoch Kräfte dabei verändern können.

5.2 Flaschenzüge

5.2.1 Seil

Anstatt eine Zugkraft direkt an einem Körper angreifen zu lassen, kann man auch über ein Seil an dem Körper ziehen.

Beispiel: Ziehen an einem Waggon einer Modelleisenbahn:



zu ergänzen:

- Kraftpfeile

Aller Erfahrung nach braucht man für denselben Zweck mit Seil genausoviel Kraft wie ohne Seil. Das Seil überträgt unsere Kraft nur auf den Körper am anderen Ende. Dabei ändert sich auch nicht die Richtung der Kraft. Das Seil ermöglicht lediglich, einen anderen, z.B. bequemen Standpunkt einzunehmen. Von den drei Eigenschaften einer Kraft – nämlich Betrag, Richtung und Angriffspunkt – ändert ein Seil demnach

.....

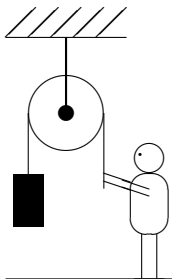
Wie immer gilt auch bei der Kraftübertragung durch ein Seil das Prinzip der Gleichheit von Kraft und Gegenkraft (*actio=reactio*): Übt Körper A auf Körper B eine Kraft aus, dann übt B auf A eine gleichgroße entgegengesetzte Kraft aus. Zieht z.B. ein Schüler als Körper A über das Seil als Körper B an einem Waggon als einem dritten Körper C, so sind eigentlich insgesamt 4 gleichgroße Kräfte im Spiel:

die Kraft, die der Schüler auf das Seil ausübt
die Kraft, die das Seil auf den Schüler ausübt

die Kraft, die das Seil auf den Waggon ausübt
die Kraft, die der Waggon auf das Seil ausübt

5.2.2 Seil und Rolle

Feste Rolle



Oft wäre es hilfreich, wenn man zumindest die Richtung einer aufzubringenden Kraft ändern könnte. Beispiel: Man will eine schwere Kiste hochheben. Es wäre dabei günstig, wenn man seine eigene Körpergewichtskraft zum Heben einsetzen könnte. Diese wirkt jedoch nach unten, und zum Heben muss eine Kraft nach oben an der Kiste angreifen. Hier hilft eine an der Decke befestigte Rolle mit darübergeführtem Seil, wie nebenstehend gezeichnet.

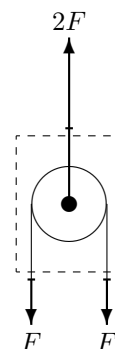
Erfahrung oder eine kurze Messung zeigen uns, dass man mit dieser Vorrichtung wieder genausoviel Kraft aufbringen muss wie ohne: Hat die Kiste z.B. ein Gewicht von 300 N, so muss man zum Halten dieser Kiste auch mit 300 N am anderen Seilende ziehen. Dennoch kann diese Vorrichtung das Heben von Lasten erheblich erleichtern, weil man eben die Zugkraft recht bequem mit Hilfe des eigenen Körpergewichtes aufbringen kann. Hauptgrund für die Bequemlichkeit ist wohl, dass die Wirbelsäule nicht belastet, ja sogar entlastet wird.

Kräfte an einer Rolle

In der folgenden Zeichnung ist nur die Rolle als Ausschnitt (gestrichelter Rahmen) aus dem vorigen Bild gezeichnet. Auf dieses Bildelement wirken von außen drei Kräfte ein: Die beiden Seilstücke ziehen mit gleichen Kräften an der Rolle nach unten, die Befestigung zieht an der Achse der Rolle nach oben.

Als Ganzes bewegt sich die Rolle nicht fort, insbesondere wird sie also auch nicht beschleunigt. Folglich muss die Summe aller an ihr angreifenden Kräfte Null sein. Die Befestigung muss daher vom Betrage her die doppelte Kraft eines Seiles ausüben. (genau genommen zuzüglich der Gewichtskraft der Rolle, die wir aber gewöhnlich vernachlässigen)

An einer Rolle sind die beiden Seilkräfte gleich groß.
An der Rollenachse wirkt die doppelte Kraft wie an jedem Seil.



Dass die beiden Seilstücke mit gleichen Kräften ziehen, gilt immer, wenn die Rolle sich nicht (– zumindest aber mit gleichbleibender Geschwindigkeit –) dreht. Da die Seile nämlich mit dem gleichen Rollenradius als Hebelarm an der Rolle angreifen, würden sie bei ungleichen Kräften auch vom Betrage her ungleiche Drehmomente bezüglich der Rollenachse ausüben. Die Befestigung greift mit Hebelarm Null an und übt folglich gar kein Drehmoment aus. Das gesamte Drehmoment wäre also bei ungleichen Seilkräften sicherlich ungleich Null, was zu einer beschleunigten Drehung der Rolle führen würde.

Lose Rolle

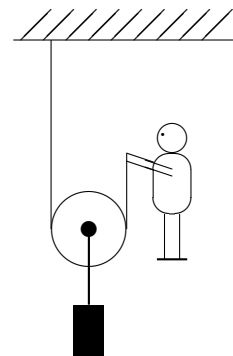
An der Rollenachse wirkt die doppelte Kraft wie an jedem Seil. Dies heißt natürlich umgekehrt auch, dass an jedem Seilende nur mit der halben Kraft gezogen werden muss wie an der Rollenachse. Was spricht also dagegen, die zu hebende Last an die Rollenachse zu hängen und dann nur noch mit der halben Kraft an jedem Seilende ziehen zu müssen?

Zunächst spricht dagegen, dass man ja an *zwei* Seilenden mit der halben Kraft ziehen muss, so dass doch wieder die gleiche Gesamtkraft erforderlich ist wie zuvor. Pfiffig wie wir sind, ziehen wir aber selbst nur an einem Ende und lassen die Decke an dem anderen Seilende ziehen, wie nebenstehend zu sehen. Über eine lose Rolle lässt sich also eine Last mit der Gewichtskraft F_L halten mit der Seil-Kraft $F_S = \frac{1}{2} F_L$. Man spart die halbe Kraft.

Einen Nachteil hat diese Vorrichtung mit loser Rolle aber auch: Um die Last z.B. um 1 cm zu heben, muss die Rolle um 1 cm gehoben werden. Dazu müssen aber die *beiden* Seilabschnitte links und rechts *jeweils* um 1 cm kürzer werden. Man muss also das selbst gehaltene Seilende 2 cm weit nach oben ziehen.

Um Dir das klarzumachen, denke Dir z.B. eine waagerechte Linie in Höhe der Hand des Männchens. Wird die Last um 1 cm gehoben, sind die beiden Seilabschnitte unterhalb dieser Linie jeweils um 1 cm kürzer. Aus dem Bereich unterhalb der Linie muss das Männchen daher insgesamt 2 cm Seil nach oben herausziehen.

Allgemeiner: Man muss beim Heben der Last um eine Strecke s_L das Seilende entlang einer Strecke $s_S = 2 \cdot s_L$ ziehen. Der Kraftersparnis mit dem Faktor $\frac{1}{2}$ steht also ein Weg-Mehraufwand mit dem Faktor 2 gegenüber.

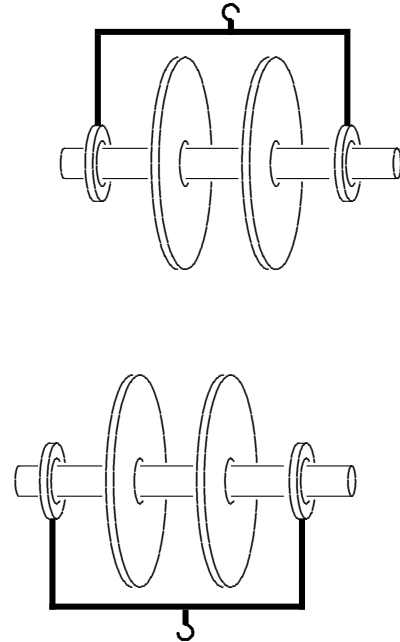
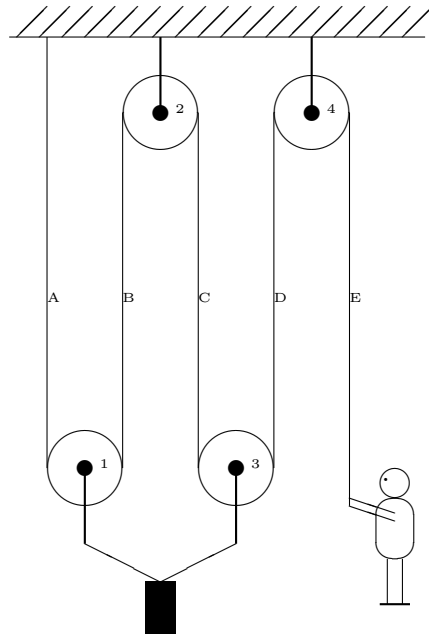


5.2.3 Flaschenzüge

Durch Kombination von festen und losen Rollen lässt sich noch mehr Kraft einsparen. Bei folgend gezeichneter Vorrichtung hängt die Last an den Achsen zweier loser Rollen. Links: schematischer Aufbau. In der Praxis setzt man die festen und die losen Rollen jeweils auf eine selbe Achse wie rechts gezeichnet. Einen so gebildeten Rollenverbund nennt man eine **Flasche**, den gesamten Aufbau aus zwei Flaschen und dem Seil einen **Flaschenzug**.

zu ergänzen:

- Seil
- Decke
- Last
- Männchen



?

Wo wirken welche Kräfte?

Gehen wir in Gedanken dem Seil entlang, von der Decke aus: Das Seilstück A ziehe mit einer Kraft F_S an der Decke. Wie auf Seite 54 dargelegt zieht Seilstück A dann auch mit F_S an Rolle 1. Nach dem Merksatz auf Seite 55 zieht dann auch Seilstück B mit F_S an Rolle 1. Wie auf Seite 54 dargelegt zieht Seilstück B dann auch mit F_S an Rolle 2. ...

⇒ alle Seilstücke ziehen mit F_S an den Rollen und am Männchen. Das Männchen zieht also ebenfalls mit F_S an seinem Seilende.

Nach dem Merksatz auf Seite 55 wirkt bei einer Seilkraft von F_S auf die Achse jeder Rolle eine Kraft vom Betrage $2F_S$.

Die Last, die ja an zwei Rollenachsen hängt, wird folglich gehalten von einer Gesamtkraft vom Betrage $2 \cdot 2F_S$. Also beträgt ihre Gewichtskraft $F_L = 4F_S$.

Zur Kontrolle: An der Decke ziehen Seilstück A mit $1F_S$ sowie die Rollen 2 und 4 mit jeweils $2F_S$. Der ganze Flaschenzug übt also insgesamt $5F_S$ auf die Decke aus. Dies sind genau die $5F_S$, mit denen die Last ($4F_S$) und das Männchen ($1F_S$) am Flaschenzug nach unten ziehen.

Natürlich wurden auch hier alle Eigengewichte im Flaschenzug vernachlässigt. Tatsächlich müsste man die Gewichte der Rollen 1 und 3 noch zu der Last hinzuaddieren, weil die Seilabschnitte A-D ja außer der Last auch diese beiden Rollen halten.

Um die Last zu halten, muss das Männchen also nur ein Viertel der Gewichtskraft der Last aufbringen. Allerdings gilt auch: Will das Männchen die Last um 1 cm anheben, muss es das Seil 4 cm weit ziehen. Es müssen ja zum Heben der Last die 4 Seilabschnitte A bis D um jeweils 1 cm verkürzt werden. Es müssen daher 4 cm Seil aus dem Flaschenzug herausgezogen werden.

Mechanische Energie

6.1 Arbeit

6.1.1 Kräfte und Wege

Wir haben im Laufe des Schuljahres einige Kraftwandler kennengelernt. Mit hydraulischen Pressen und Hebebühnen, mit Hebeln und mit Flaschenzügen ist es möglich, die für einen Zweck nötige Kraft dem Betrage nach zu verändern. Wir haben jedoch festgestellt, dass sich dabei auch die Länge des Weges verändert, entlang dessen die veränderte Kraft wirken muss.

Wir gehen im Folgenden davon aus, dass für ein bestimmtes Vorhaben ohne Kraftwandler eine Kraft F entlang eines Weges s wirken müsste. Es gehe z.B. darum, eine Last mit dem Gewicht F um eine gewisse Strecke s anzuheben. Dazu können Kraftwandler benutzt werden:

Flaschenzug

Die aufzubringende Kraft kann herabgesetzt werden auf einen Bruchteil $\frac{1}{r} \cdot F$, wobei r von der Bauart des Flaschenzuges abhängt. Bei allen von uns durchgerechneten Bauarten muss dann jedoch das Seil um $r \cdot s$ her-

ausgezogen werden. Grund: Wenn mehr Seilabschnitte die Last tragen, sind zwar geringere Seilkräfte nötig, aber es müssen beim Heben auch mehr Seilabschnitte verkürzt werden.

Hydraulische Hebebühne

Entscheidend ist hier das Verhältnis der Kolbenflächen von Hub- und Pumpkolben, $\frac{A_{Hub}}{A_{Pump}}$, im Folgenden r genannt. Durch Wahl der Kolbenflächen kann r beliebig vorgegeben werden. Die am Pumpkolben erforderliche Kraft beträgt $\frac{1}{r} \cdot F$, jedoch muss der Pumpkolben

mit dieser Kraft um die Strecke $r \cdot s$ bewegt werden. Grund: Liegt die Last auf einer größeren Hubkolbenfläche, genügt dort zwar ein geringerer Druck, aber es muss beim Heben auch ein größeres Volumen an Hydrauliköl unter den Hubkolben gedrückt werden.

6.1.2 Definition

Kraft F und Weg s werden durch Kraftwandler also immer so verändert, dass die eine Größe mit derselben Zahl r multipliziert wird, durch die die andere dividiert wird. Bildet man das Produkt aus Kraft und Weg, so kürzt sich diese Zahl heraus und man erhält unabhängig von r :

$$\frac{1}{r} \cdot F \cdot r \cdot s = F \cdot s$$

Für Physiker sind derart unabhängige Ausdrücke

sehr wichtig und nützlich. Einerseits kann man ja zur Lösung eines Problems nur dann *Gleichungen* aufstellen, wenn man Ausdrücke kennt, die eben unabhängig von den jeweiligen Umständen immer den *gleichen* Wert haben. Andererseits deutet die Unveränderlichkeit eines Ausdruckes darauf hin, dass er eine wesentliche Eigenschaft des Problems ausdrückt und für das Problem eine zentrale Bedeutung hat.

?

Welche Bedeutung hat der Ausdruck $Kraft \cdot Weg$?

Der Bedeutung eines Ausdruckes kommt man oft dadurch auf die Spur, dass man sich überlegt, was es bedeutet, wenn der Ausdruck einen größeren oder kleineren Wert hat.

Beim Hochheben eines Körpers hat $F \cdot s$ einen größeren Wert, wenn F größer ist, und auch, wenn s größer ist. Größeres F bedeutet praktisch, dass ein schwererer Körper gehoben wird. Größeres s bedeutet, dass der Körper höher gehoben wird. In beiden Fällen würde man schon ohne Physik-Kenntnisse sagen, dass beim Hochheben mehr **Arbeit** verrichtet wird. Man definiert demnach:

zu ergänzen:

- Terme

Wird entlang des Weges s die Kraft F in Richtung des Weges ausgeübt, so heißt das Produkt $F \cdot s$ die dabei verrichtete **mechanische Arbeit** W . Kurz:

$$W :=$$

Die Grundeinheit der Arbeit ist

$$1 \text{ Joule} = 1 \text{ J} :=$$

W wurde als Abkürzung des englischen Wortes *work* gewählt. Die Einheit Joule ist benannt zu Ehren des britischen Physikers James Prescott Joule (1818-1889)

Du hast die Einheit Newtonmeter bereits als Einheit für Drehmomente kennengelernt. Trotz übereinstimmender Bezeichnung der Einheiten und ähnlicher Definition als „Kraft mal Weg“ sind Arbeit und Drehmoment jedoch völlig verschiedene Größen. Ein Unterschied liegt bereits darin, dass beim Drehmoment Kraft und Weg *senkrecht* aufeinander stehen, bei der Arbeit jedoch *parallel* zueinander verlaufen.

In einigen Jahren kannst Du lernen, wie man Größen mit Berücksichtigung ihrer Richtungen miteinander

verrechnen kann. Die Multiplikation von Kraft und Weg zu Drehmoment wird dann als eine andere Multiplikation zu erkennen sein als die Multiplikation von Kraft und Weg zu Arbeit.

Aus Bequemlichkeit und weil Du ja bisher nur eine Art von Multiplikation kennst, schreiben wir in beiden Fällen dasselbe „mal“ und lesen die Einheit in beiden Fällen als „Newtonmeter“. Wir halten uns aber bereits jetzt daran, dass nur das Newtonmeter für Arbeit auch Joule genannt werden darf.

6.1.3 Anstrengung und Arbeit

In obiger Definition der Arbeit wurde eingeschränkt, dass die Kraft nicht nur *entlang* des Weges ausgeübt werden muss, sondern auch *in Richtung* des Weges wirken muss.

Ein Fall, bei dem dies nicht erfüllt ist, ist das Umhertragen einer Last auf ebenem Gelände. Es geht z.B. die Schülerin Carina mit einem 7 kg schweren Koffer in der Hand durch einen 50 m langen ebenen Gang. Sie muss dabei den ganzen Weg über eine Kraft von etwa 70 N auf den Koffer ausüben, damit er nicht der Schwerkraft folgend herabfällt. Diese Koffer-Halte-Kraft wirkt aber vertikal nach oben und damit senkrecht zum horizontal verlaufenden Weg. Die Kraft wirkt daher mit keiner Komponente in Richtung des Weges. Folglich ist dieses Tragen des Koffers nach obiger Definition keine physikalische Arbeit.

Obwohl das Tragen den Träger sicher anstrengt, ist die Definition von Arbeit physikalisch doch sinnvoll: Das Ergebnis des Tragens, nämlich den Ortswechsel des Koffers, könnte man ja ebenso ohne nennenswerte Anstrengung erreichen, indem man z.B. den Koffer auf einem Kuli schiebt. Der Ortswechsel alleine kann also nicht die Anstrengung verursachen. Die Anstrengung ist eigentlich nur Folge der Arbeitsweise unserer Muskeln.

Unvermeidbare Anstrengung tritt immer nur dann auf, wenn die Kraft – zumindest mit einer Komponente – in Richtung des Weges wirkt. Genau diese unvermeidbare Anstrengung deckt der physikalische Arbeitsbegriff ab. Falls Kraft und Weg schräg zueinander stehen, wird zur Berechnung der Arbeit nur die Kraftkomponente in Wegrichtung berücksichtigt.

6.1.4 Goldene Regel der Mechanik

Das Tragen des Koffers unterscheidet sich damit grundlegend vom Hochheben des Koffers. Bei letzterem kann keine Vorrichtung der Welt den nötigen Arbeitsaufwand

$F \cdot s$ verringern: Die erhöhte Lage des Koffers als Ergebnis des Hebens erfordert unausweichlich diese ganz bestimmte Arbeit, egal wie man es auch anstellen mag!

Goldene Regel der Mechanik: Der Einsatz eines (idealen) Kraftwandlers ändert die erforderliche mechanische Arbeit nicht.

ideal wurde hinzugefügt, weil ja z.B. ein realer Flaschenzug mit schlecht geölten Rollen und schweren, steifen Seilen sehr wohl die Arbeit erschweren, und damit erhöhen kann.

Natürlich kann man es *sich selbst* erleichtern, den Koffer hochzuheben, indem man einen Motor benutzt. Geeignet wäre z.B. ein Dieselmotor, der über Seile und Rollen den Koffer hebt.

Damit hat man aber die Arbeit nicht eingespart, sondern sie nur *dem Motor* überlassen: Seine Anstrengung lässt sich an seiner Erwärmung und seinem Diesel-

verbrauch erkennen, ganz entsprechend dem Schwitzen und dem gesteigerten Hunger eines arbeitenden Menschen. Auch der Motor müsste in jedem Fall eine Kraft F in Richtung und entlang eines Weges s ausüben, so dass $F \cdot s$ denselben Wert hätte wie beim Hochheben ohne Motor.

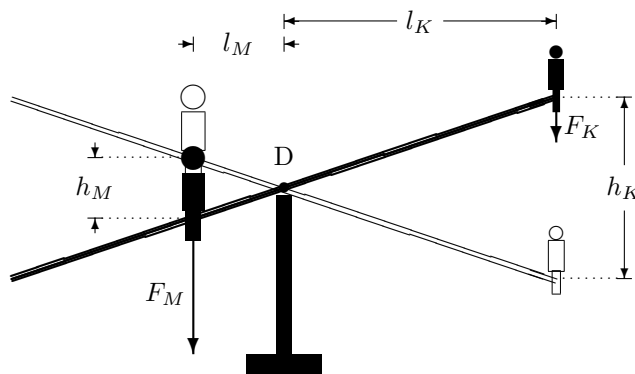
6.1.5 Kraft und Weg bei Hebeln

Neben Flaschenzügen und hydraulischen Systemen haben wir auch Hebel als Kraftwandler kennengelernt. An ihnen haben wir bisher zwar die Kräfte, jedoch noch nicht die zurückgelegten Wege betrachtet.

Auf einer Spielplatz-Wippe kann ein Kind, das nur ein Drittel des Gewichtes eines erwachsenen Mannes hat, dem Mann das Gleichgewicht halten. Dazu

muss nur sein Hebelarm dreimal so lang wie der des Mannes sein. Um aber den Mann um eine Höhe h_M anzuheben, muss dann das Kind auch um eine dreimal größere Strecke h_K absinken.

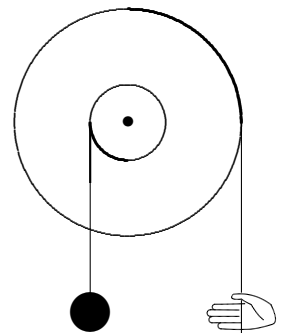
Mit $F_K = \frac{1}{3}F_M$ und $h_K = 3h_M$ ist also die auf der Seite des Kindes erbrachte Arbeit $F_K \cdot h_K = \frac{1}{3}F_M \cdot 3h_M = F_M \cdot h_M$ gleich der zum Heben des Mannes erforderlichen Arbeit.



Betrachten wir auch noch ein Wellrad, z.B. das nebenstehend abgebildete, bei dem der kleinere Radius 5 cm, der größere 15 cm betrage! Damit soll nun eine Last von 150 kg gehoben werden. Die Last übt eine Gewichtskraft von 1500 N auf das kleinere Rad aus. Sie erzeugt demnach ein linksdrehendes Drehmoment von $1500 \text{ N} \cdot 0,05 \text{ m} = 75 \text{ Nm}$.

Zum Halten der Last muss die Hand ein ebenso großes Drehmoment am größeren Rad erzeugen. Dazu ist nur die Kraft $75 \text{ Nm} / 0,15 \text{ m} = 500 \text{ N}$ erforderlich. Dies ist ein Drittel der Gewichtskraft der Last. Klar, denn der Hebelarm der Hand ist ja auch dreimal so lang wie der der Last.

Nun werde die Last z.B. mit einer viertel Drehung des Wellrades angehoben. Die Last selbst wird dabei um einen viertel Umfang des kleineren Rades gehoben. Dies sind $\frac{1}{4} \cdot 2\pi \cdot 5 \text{ cm} = 7,9 \text{ cm}$. Die Hand muss das Seil jedoch einen viertel Umfang des



größeren Kreises weit ziehen. Dies sind $\frac{1}{4} \cdot 2\pi \cdot 15 \text{ cm} = 23,6 \text{ cm}$, also dreimal so viel wie die Last gehoben wird. Klar, denn der Kreis mit dreifachem Radius hat ja auch dreifachen Umfang.

Also gehört auch beim Wellrad zu einer durch n geteilten Kraft ein n -facher Weg. Die Arbeit zum Heben der Last bleibt mit Einsatz des Wellrades so groß wie ohne. Die goldene Regel der Mechanik ist damit auch für diesen Kraftwandler bestätigt.

6.1.6 Formen mechanischer Arbeit

Hubarbeit

...ist wohl die einfachste Art von Arbeit: Ein Körper der Masse m werde um eine gewisse Höhe h angehoben. Es muss dazu entlang des Hubweges eine Kraft F_H entgegengesetzt zur Gewichtskraft des Körpers ausgeübt werden. Die Gewichtskraft beträgt $F_G = m \cdot g$, wobei g der Ortsfaktor ist. Der gesamte Hubvorgang hat dabei im einfachsten Fall drei Phasen:

1. Der noch unten ruhende Körper muss nach oben in Bewegung gesetzt, also beschleunigt werden. Dazu muss die ausgeübte Kraft F_H zumindest ein wenig größer als F_G sein. Nur dann weist die resultierende Kraft nach oben.
2. Sobald der Körper durch diese resultierende Kraft zumindest auf eine winzige Geschwindigkeit nach oben beschleunigt worden ist, ist eine weitere Beschleunigung nicht mehr erforderlich. Es genügt nun, dass F_H die Gewichtskraft F_G genau ausgleicht: $F_H = F_G$. Bei diesem Kräftegleichgewicht ist die resultierende Kraft Null und der Körper bewegt sich dann nach dem Trägheitssatz mit der zuvor erreichten Geschwindigkeit weiter nach oben.
3. Bevor der Körper die angestrebte Höhe erreicht, muss er abgebremst werden. Dazu wird F_H rechtzeitig zu-

mindest ein wenig unter F_G abgesenkt. Die resultierende Kraft wirkt dann nämlich nach unten, der Bewegungsrichtung entgegen.

Von den Abweichungen in der Start- und in der Stopp-Phase abgesehen beträgt also die zum Heben erforderliche Kraft entlang des ganzen Weges F_G . Es ist aus zwei Gründen zulässig, von diesen Abweichungen ganz abzusehen:

1. Die Abweichungen dürfen beliebig klein sein und beliebig kurz andauern, so dass sie beliebig kleine Rollen spielen. Man erreicht mit kleinen Abweichungen zwar nur kleine Beschleunigungen und daher in kurzen Beschleunigungsphasen auch nur geringe Geschwindigkeiten, so dass der Hubvorgang entsprechend lange dauert. Macht aber nichts: In der Theorie haben wir viel Zeit ...
2. Selbst wenn die Abweichungen größer sind und länger dauern, gleichen sich die durch sie verursachten Mehr- und Minderarbeiten immer genau aus: Was man unten beim Start mehr arbeiten muss, spart man oben wieder genau ein.

Wir können demnach ruhigen Gewissens die Arbeit so berechnen, als ob den ganzen Hub-Weg entlang die Kraft $F_H = F_G$ ausgeübt würde.

Um einen Körper der Masse m um die Höhe h hochzuheben ist die Hubarbeit W_h erforderlich mit

$$W_h =$$

zu ergänzen:

- Formelterm

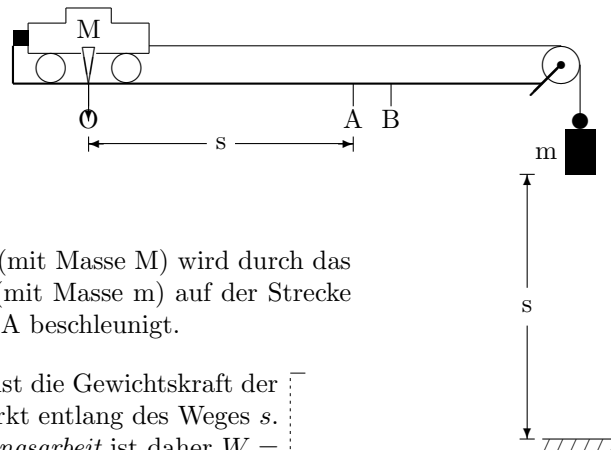
Beschleunigungsarbeit

...wird verrichtet, wenn die Geschwindigkeit eines Körpers erhöht oder vermindert wird. Wir gehen erst einmal davon aus, dass ein Körper bestimmter Masse *aus der Ruhe* auf eine Geschwindigkeit v beschleunigt werden soll. Es ist klar, dass die erforderliche Arbeit von v und der Masse des Körpers abhängen muss.

Mangels Rechenkünsten aus höheren Schuljahren können wir den Zusammenhang zwischen der beschleunigten Masse, der erreichten Geschwindigkeit und der Arbeit nicht rechnerisch herleiten. (Wir müssten dazu besser über den Zusammenhang zwischen der Beschleunigungsstrecke und der Geschwindigkeit Bescheid wissen.) Wir können jedoch in einem geeigneten Experiment die Natur selbst nach dem gesuchten Zusammenhang befragen.

V 6.1 Beschleunigungsarbeit

Aufbau:



Durchführung: Der Wagen (mit Masse M) wird durch das herabsinkende Gewichtsstück (mit Masse m) auf der Strecke (der Länge s) zwischen O und A beschleunigt.

- Die beschleunigende Kraft ist die Gewichtskraft der Masse m , also $m \cdot g$. Sie wirkt entlang des Weges s . Die geleistete *Beschleunigungsarbeit* ist daher $W = m \cdot g \cdot s$ und folglich von uns durch Vorgabe von m und s einstellbar. Da sich m durch Austausch des Gewichtsstückes leichter verändern lässt als s durch allerlei Verschiebungen am Versuchsaufbau, werden wir s unverändert lassen und die Arbeit W über m einstellen.
- Die beschleunigte *Masse* beträgt $M + m$, im folgenden $\tilde{m}$ genannt. Die Erde zieht ja mit $m \cdot g$ an dem „Verbund-Körper“ aus M und m . Auch das Gewichtsstück wird auf v beschleunigt!
- Nach der Beschleunigungsphase bewegt sich der Wagen kräftefrei und daher mit gleichbleibender Geschwindigkeit weiter. Um diese *Geschwindigkeit* zu ermitteln, stoppen wir die Zeit Δt , die der Wagen braucht, um die Strecke zwischen A und B zu durchfahren.

Mit diesem Versuch suchen wir den Zusammenhang zwischen *drei* Größen, nämlich $\tilde{m}$, v und W . *Zwei* davon, $\tilde{m}$ und W , geben wir indirekt über m und M vor, die dritte, nämlich v , ergibt sich aus der gemessenen Zeit Δt .

In der Auswertung ist es jedoch leichter, wenn zunächst nur *eine* vorzugebende Größe variiert. Wir werden daher bei der Auswahl von m und M für verschiedene Versuchsdurchgänge darauf achten, dass $\tilde{m}$ mehrmals den gleichen Wert hat. Bei einem gleichbleibenden Wert von $\tilde{m}$ können wir dann nämlich schon einmal untersuchen, welchen Einfluss W *alleine* auf v hat. Der zusätzliche Einfluss eines veränderten $\tilde{m}$ wird dann anschließend besser zu erkennen sein.

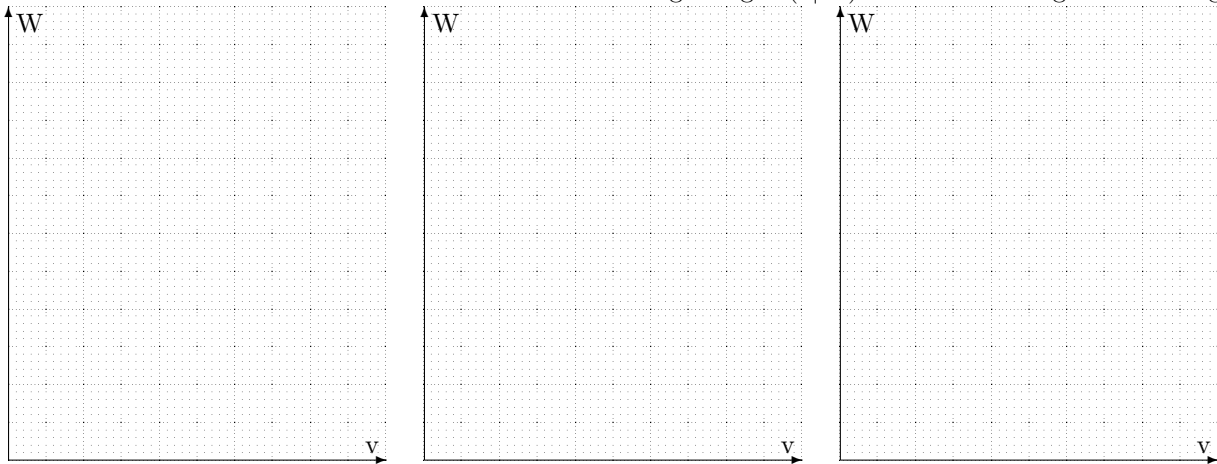
Messwerte: /Auswertung:

$|\overline{AB}| = \dots\dots\dots$

$s = \dots\dots\dots$

m													
M													
Δt													
W													
v													
$\tilde{m}$													

Um den Zusammenhang zwischen v und W besser zu erkennen, kann man jeweils für einen festen Wert von $\tilde{m}$ die zugehörigen $(v|W)$ -Paare in ein Diagramm eintragen:



Verdacht: Dem Schaubild nach könnte gelten $W \dots\dots\dots$
 Dies prüfen wir in der letzten Zeile obiger Tabelle. Der Vergleich mit der $\tilde{m}$ -Zeile ergibt das

Ergebnis: _____

Um einen Körper der Masse m aus der Ruhe auf die Geschwindigkeit v zu beschleunigen, ist die Beschleunigungsarbeit W_B erforderlich mit

$$W_B =$$

zu ergänzen:

- Formelterm

Diese Formel haben wir aus den Maßzahlen ermittelt, die wir in unserer Versuchsauswertung gewonnen haben. Physikalische Größenangaben haben jedoch auch Einheiten. Wenn die Formel für W_B richtig ist, (– und das ist sie –) dann muss auch auf beiden Seiten der Gleichung dieselbe Einheit stehen. Demnach muss gelten:

$$1 \text{ J} = 1 \text{ kg} \cdot \frac{1 \text{ m}^2}{1 \text{ s}^2}$$

Ersetzt man das Joule gemäß seiner Definition durch Newtonmeter und dividiert beide Seiten durch Meter, ergibt sich der Zusammenhang:

$$1 \text{ N} = 1 \text{ kg} \cdot \frac{1 \text{ m}}{1 \text{ s}^2}$$

Spontan magst Du jetzt denken: „Ich hab’s ja schon lange geahnt: Mit zuviel Mathematik kann man nur Unsinn produzieren.“ Bei genauerem Hinsehen bestätigt sich aber hier nur die Definition des Newton von Seite 24.

$$F = m \cdot a$$

Spannarbeit

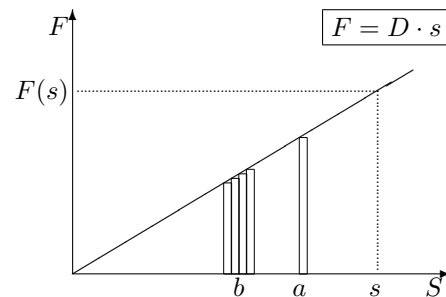
Eine weitere Form mechanischer Arbeit ist das Spannen einer Feder. (Allgemeiner: das Verformen eines elastischen Körpers) Zur Bestimmung der Arbeit W , die nötig ist, um eine Feder mit Federhärte D um eine Strecke s zu dehnen, kannst Du hier eine für Physiker typische Denkweise kennenlernen, einen Vorgeschmack der Integralrechnung.

Arbeit ist definiert als $F \cdot s$.

F ist dabei die Kraft, die entlang s wirkt. Hier wirkt aber keine gleichbleibende Kraft entlang der ganzen Dehnungsstrecke s , denn während des Dehnvorganges wächst die ausgeübte Kraft so wie hier dargestellt.

Denkt man sich die ganze Dehnungsstrecke in winzig kleine Teilstücke zerlegt, kann man aber immerhin für jedes Teilstück von einer etwa gleichbleibenden Kraft ausgehen. Entlang des bei a eingezeichneten Teilstückes ändert sich die Kraft ja relativ wenig: Stellt auf der Kraft-Achse jeder Millimeter z.B. ein Newton dar, dann wächst die Kraft von 18 N auf 18,6 N.

Der Unterschied wird offensichtlich beliebig klein, wenn man nur die Teilstücke kurz genug wählt. Die Arbeit entlang eines Teilstückes ist daher nur wenig größer als das Produkt aus der Teilstücklänge und der Kraft zu Beginn des Teilstückes. Die beiden Faktoren in diesem Produkt entsprechen im Diagramm der Breite und der Höhe des eingezeichneten Rechtecks. Das Produkt aus Breite und Höhe eines Rechteckes ist andererseits der Flächeninhalt (– bitte noch einfärben –) des Rechteckes.



Die gesamte Arbeit entlang aller Teilstücke entspricht der Summe aller entsprechenden Rechteckflächen, von denen an der Stelle b einige eingezeichnet sind. Man erkennt, dass diese Gesamtfläche aller Rechtecke gleich der Fläche unter dem Funktionsgraph ist, wenn man von den kleinen Dreieckchen zwischen Graph und Rechtecken absieht. Die Gesamtfläche dieser Dreieckchen wird sowieso beliebig klein, wenn man die Rechtecke ausreichend schmal wählt.

Die Arbeit W entspricht dem Flächeninhalt unter dem $F(s)$ -Graphen.

Für die Dehnung der anfänglich um 0 gedehnten (d.h. entspannten) Feder bis zur Dehnung s ergibt sich als Arbeit der „Flächeninhalt“ des Dreiecks unter dem Graphen zwischen Ursprung und dem Punkt s auf der

waagerechten Achse. Dieses Dreieck ist die Hälfte des gepunkteten Rechteckes mit Breite s und Höhe $F(s) = D \cdot s$. So ergibt sich $W = \frac{1}{2} \cdot s \cdot Ds = \frac{1}{2} Ds^2$

Um eine Feder mit Federhärte D aus der entspannten Lage um die Strecke s zu dehnen, ist die Arbeit W erforderlich mit

$$W =$$

zu ergänzen:

- Formelterm

6.2 Leistung

In den vorigen Abschnitten hat es keine Rolle gespielt, wie schnell eine Arbeit verrichtet wird. Häufig interessiert jedoch genau dies. Wer ein Auto kauft, will meist auch wissen, wie lange sein Motor braucht, um es z.B. von 0 auf $100 \frac{\text{km}}{\text{h}}$ zu beschleunigen. Für einen Kranfahrer spielt es eine Rolle, wie viel Hubarbeit sein Kran z.B. pro Minute verrichten kann. Allgemeiner gesagt interessiert häufig die *Leistungsfähigkeit* eines Gerätes. Ein leistungsfähigeres Gerät schafft dieselbe Arbeit in kürzerer Zeit, oder auch mehr Arbeit in derselben Zeit. Man definiert daher:

Entsprechend definierte Einheit:

$$1 \text{ Watt} = 1 \text{ W} :=$$

Unter der (mittleren) Leistung P in einer Zeitspanne t versteht man den Quotienten aus der in dieser Zeitspanne verrichteten Arbeit W und der Zeitspanne selbst. Kurz:

$$P :=$$

zu ergänzen:

- Terme

Verwechsle nicht die Einheit „Watt“ mit der Größe „Arbeit“, obwohl beide mit W abgekürzt werden! Ähnliches kennst Du ja bereits: m , s , A , V

6.3 Energie

6.3.1 Definition

Im vorigen Versuch zur Beschleunigungsarbeit hat das Gewichtsstück m beim Herabsinken Beschleunigungsarbeit verrichtet.

Die Fähigkeit eines Körpers oder Systems, Arbeit zu verrichten, heißt seine **Energie**. Sie wird zahlenmäßig angegeben als die Arbeit, die der Körper bzw. das System verrichten kann.
Formelsymbol: E oder W Einheit: Joule

6.3.2 Arten mechanischer Energie

Lageenergie

Arbeit zu verrichten war dem Gewichtsstück offensichtlich aufgrund seiner *erhöhten Lage* möglich. Da beim Herabsinken des Gewichtsstückes dessen Gewichtskraft $m \cdot g$ entlang der Strecke s gewirkt hat, betrug die verrichtete Arbeit $m \cdot g \cdot s$.

Geht man davon aus, dass das Stück nicht noch weiter als auf den Boden herabsinken kann, so war $m \cdot g \cdot s$ seine gesamte Energie: In der unteren Lage kann das Gewichtsstück keine weitere Arbeit mehr leisten.

Geht man davon aus, dass das Stück auch noch tiefer sinken kann (z.B. in eine Grube), so lässt sich immernoch zumindest sagen, dass es in der unteren Lage $m \cdot g \cdot s$ an Arbeit weniger verrichten kann als in der oberen. Es hat eben so viel Energie weniger.

Tatsächlich lassen sich immer nur solche Energiedifferenzen eindeutig feststellen. In welcher Lage man die Energie eines Körpers als Null ansieht, ist dieselbe Willkür, mit der man Nullpunkte für Zeitskalen, Positionsangaben oder Temperaturskalen festlegt.

Die Energie, die das Gewichtsstück in der oberen Lage mehr hat als in der unteren, ist genau gleich der Hubarbeit, mit der man es irgendwann zuvor aus der unteren in die obere Lage angehoben hat. Genausoviel Arbeit, wie man selbst beim Hochheben verrichtet, bekommt das Gewichtsstück dabei als Energie, so dass es später genau diese Menge an Arbeit verrichten kann. Energie kann demnach als *gespeicherte Arbeit* angesehen werden.

Bewegungsenergie

Die Fähigkeit, Arbeit zu verrichten, hat ein Körper auch aufgrund seiner *Geschwindigkeit*: Führt z.B. ein Wagen gegen eine Schraubenfeder, drückt er diese zusammen, wobei er selbst abgebremst wird. Entlang des Bremsweges übt er jedoch eine Kraft in Richtung des Weges auf die Feder aus: er verrichtet Arbeit, nämlich „Spannarbeit“.

Natürlich kann der bewegte Wagen auch über ein Seil und eine Rolle ein Gewichtsstück nach oben ziehen. Hierbei würde er bis zu seinem Stillstand Hub-

arbeit verrichten. Noch einfacher aufzubauen wäre ein Versuch bei dem der Wagen auf ansteigender Fahrbahn auf Kosten seiner Bewegungsenergie selbst an Höhe gewinnt.

Und wie viel Arbeit kann ein bewegter Körper verrichten? Du ahnst es sicher schon: Ein Körper kann aufgrund seiner Geschwindigkeit genausoviel Arbeit verrichten, wie es Beschleunigungsarbeit erfordert hat, ihn auf seine Geschwindigkeit zu bringen.

Spannenergie

Natürlich hat auch eine gespannte Feder die Fähigkeit, beim Entspannen Arbeit zu verrichten, also Energie. Auch hier beträgt die Energie ebensoviel, wie zum Spannen nötig war. Zum Nachweis dieses Sachverhalts könnte z.B. ein Versuch dienen, bei dem eine zusammengedrückte Feder einen anderen Körper hochkatapultiert.

Unsere bisherigen Energieformeln: Die Energie, die ...			
(↓)	aufgrund (↓) hat,	heißt (↓)	und beträgt (↓)
ein Körper der Masse m	seiner um h erhöhten Lage	Lageenergie oder potenzielle Energie	$E_{\text{pot}} =$
ein Körper der Masse m	seiner Geschwindigkeit v	Bewegungs- oder kinetische Energie	$E_{\text{kin}} =$
eine Feder der Härte D	ihrer elastischen Verformung um s	Spannenergie	$E_{\text{spann}} =$

6.3.3 Energieumwandlungen

Das Gewichtsstück, das beim Herabsinken einen Körper beschleunigt, verrichtet Beschleunigungsarbeit und verliert dabei Lageenergie. Der beschleunigte Körper erhält Bewegungsenergie.

Der bewegte Körper, der eine Feder verformt, verrichtet Spannarbeit und verliert dabei Bewegungsenergie. Die Feder erhält aufgrund ihrer Verformung Spannenergie.

Mit dieser Spannenergie könnte die Feder ein Gewichtsstück hoch katapultieren. Die Feder verlöre dabei Spannenergie, das Gewichtsstück gewänne Bewegungs- und Lageenergie.

Bei allen besprochenen Arbeitsvorgängen gibt es einen Körper, der Energie *verliert* und einen Körper, der Energie *gewinnt*. Manchmal ist es derselbe Körper, der Energie in einer Form verliert und in einer anderen Form gewinnt.

Beispiel: Ein hoch gehaltener Gummiball hat Lageenergie. Lässt man ihn los, verliert er an Höhe und damit an Lageenergie, gewinnt jedoch gleichzeitig an Geschwindigkeit und damit an Bewegungsenergie. Trifft

er auf hartem Boden auf, wird er verformt und abgebremst, verliert also Bewegungsenergie und gewinnt Spannenergie. Anschließend federt er wieder unter Verlust dieser Spannenergie in seine Kugelform zurück, wobei er sich unter Gewinn von Bewegungsenergie nach oben beschleunigt. Beim Hochfliegen geht diese Bewegungsenergie zugunsten von Lageenergie wieder verloren. Und so fort.

Bei vernachlässigbaren Luftreibungs- und anderen Störeffekten erreicht der Ball wieder etwa seine ursprüngliche Höhe. Er hat also nach dem Vorgang dieselbe Lageenergie wie vorher. Es liegt nahe, dass er die ganze Zeit über immer dieselbe Menge an Energie hatte, lediglich in wechselnden Formen. Alle Arbeiten, die stattgefunden haben, waren demnach nur Energieumwandlungen von einer Form in eine andere.

Auch im vorigen Fahrbahnversuch wurde nur die Lageenergie des Gewichtsstückes (fast) vollständig umgewandelt in Bewegungsenergie von Wagen und Gewichtsstück.

Energieerhaltungssatz der Mechanik: In einem abgeschlossenen System ohne Reibung bleibt die Summe aller mechanischen Energien alle Zeit gleich.

Wie nützlich es ist, bei verschiedensten Abläufen eine Größe zu kennen, die jedenfalls die ganze Zeit über *gleich* bleibt, wird sich in vielen Aufgaben zeigen. Wie sollte man nämlich ohne gleichbleibende Größen *Glei-*

chungen zur Lösung der Aufgabe aufstellen? Bei vielen Aufgaben kann man die entscheidende Gleichung aufstellen nach dem Muster:

$$\begin{array}{c} \text{Summe aller Energien} \\ \text{zu einem Zeitpunkt} \end{array} = \begin{array}{c} \text{Summe aller Energien zu} \\ \text{einem anderen Zeitpunkt} \end{array}$$

6.4 Reibungsarbeit

zu ergänzen:

- Formelterm

Wird ein Körper gegen eine Reibungskraft F_R gleitend oder rollend entlang einer Strecke s bewegt, so wird dabei die Reibungsarbeit W_R verrichtet mit

$$W_R =$$

Dies ergibt sich direkt aus der Definition von Arbeit. Im Gegensatz zur Hub- und Beschleunigungsarbeit führt Reibungsarbeit an einem Körper (System) jedoch nicht zu einer Zunahme der mechanischen Energie des Körpers (Systems). Der Reibungsarbeit entspricht keine mechanische Energieform.

Beispiel: Nachdem man eine Kiste ein Stück weit über den Boden gezogen hat, hat die Kiste weder Lageenergie aufgrund erhöhter Lage, noch Bewegungsenergie aufgrund ihrer Geschwindigkeit, noch Spannenergie aufgrund einer elastischen Verformung. Wir erkennen im Zustand der Kiste nach dem Reiben keine mechani-

sche Veränderung, die eine erhöhte Fähigkeit zum Verrichten von Arbeit – und d.h. keine erhöhte Energie – mit sich brächte.

Andererseits muss ein System, um eine Reibungsarbeit verrichten zu können, die Fähigkeit dazu, also Energie haben. (eine bloße Umformulierung!) Verrichtet das System aber unter Einsatz seiner mechanischen Energie eine Reibungsarbeit, so nimmt seine mechanische Energie um den Betrag ebendieser Reibungsarbeit ab, wie es bei jeder anderen Form von Arbeit auch wäre.

Zum Lösen vieler Aufgaben bedeutet dies, dass man als Ansatz benutzen kann:

$$\begin{array}{c} \text{Summe aller} \\ \text{mech. Energien} \\ \text{zu Beginn} \end{array} = \begin{array}{c} \text{dazwischen} \\ \text{verrichtete} \\ \text{Reibungsarbeiten} \end{array} + \begin{array}{c} \text{Summe aller} \\ \text{mech. Energien} \\ \text{am Schluss} \end{array}$$

6.5 Wirkungsgrad

Bei praktisch allen mechanischen Vorgängen tritt Reibung auf. Daher wird bei solchen Energieumwandlungen immer ein gewisser Anteil der anfänglich vorhandenen mechanischen Energie für Reibungsarbeit verwendet. Dieser Anteil ist am Ende des Vorganges nicht mehr in Form mechanischer Energie nutzbar. Zur Beschreibung des weiterhin nutzbaren Anteils verwendet man folgenden Begriff:

Der Wirkungsgrad η einer Energieumwandlung ist das Verhältnis aus der nutzbar bleibenden Energie zur insgesamt aufgewandten Energie:

$$\eta := \frac{W_{Nutz}}{W_{Aufwand}}$$

η ist der griechische Buchstabe, der einem $\tilde{a}$ entspricht, und *äta* gelesen wird.

Beispiel: Ein Körper hat zunächst wegen einer erhöhten Lage gegenüber dem Boden eine Lageenergie von 500 J und fällt dann auf den Boden herab. Beim Fall wird er durch den Luftwiderstand (Reibung an/in der Luft) gebremst, so dass er beim Auftreffen eine Bewegungsenergie von nur 400 J hat. Die Luftreibung hat also 100 J „gekostet“. Der Wirkungsgrad bei der Umwandlung der anfänglichen Lageenergie in Bewegungsenergie beträgt

$$\eta = \frac{400 \text{ J}}{500 \text{ J}} = \frac{4}{5} = 0,8 = 80\%$$

6.6 Weg-Zeit-Gesetz bei Beschleunigung

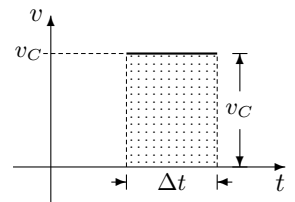
Die Formel $W = \frac{1}{2}mv^2$ für die Beschleunigungsarbeit kann auf *experimentellem* Wege gefunden werden, etwa so wie ab Seite 61 beschrieben. Mit diesem Abschnitt will ich Dir aber auch eine *rechnerische* Herleitung nicht vorenthalten. Wir wissen bereits:

1. Zum Beschleunigen eines Körpers der Masse m muss eine Kraft F entlang einer Beschleunigungsstrecke s auf den Körper einwirken. Nach der Definition von Arbeit wird dabei natürlich die Arbeit $W = F \cdot s$ (Gleichung I) verrichtet. Diese Formel erlaubt die Berechnung der Arbeit über die Beschleunigungsstrecke, nicht jedoch über die Endgeschwindigkeit.
2. Die Geschwindigkeit kommt aber folgendermaßen ins Spiel: Die Kraft F bewirkt die Beschleunigung

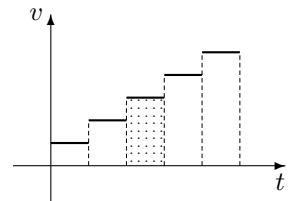
$a = \frac{F}{m}$. Dauert der Beschleunigungsvorgang die Zeitspanne Δt lang, so ergibt sich eine Geschwindigkeitsänderung $\Delta v = a \cdot \Delta t$. Sofern der Körper in Ruhe startet, erreicht er also bis zur Zeit t die Geschwindigkeit $v = a \cdot t$. Drückt man a durch F aus, so ergibt sich $v = \frac{F}{m} \cdot t$, und nach F aufgelöst: $F = \frac{v \cdot m}{t}$ (Gleichung II)

Die Gleichungen I und II lassen sich nun kombinieren, indem man F in Gleichung I durch den laut Gleichung II zu F gleichwertigen Ausdruck $v \cdot m/t$ ersetzt: $W = \frac{v \cdot m}{t} \cdot s$, was zwar auch $= v \cdot m \cdot \frac{s}{t}$ ist, wobei aber $\frac{s}{t} \neq v$. Die Formel $v = \frac{s}{t}$ definiert ja nur die *mittlere* Geschwindigkeit während der Zeitspanne t , aber v ist hier die *momentane* Geschwindigkeit am End-Zeitpunkt der Beschleunigung.

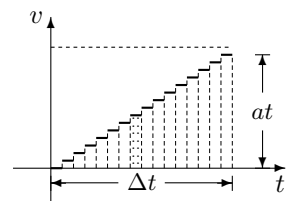
Um von hier aus weiter zu kommen, benötigen wir also den Zusammenhang zwischen Weg s und Zeit t bei beschleunigter Bewegung. Hier hilft derselbe „Trick“ wie bei der Berechnung der Spannarbeit: Bei konstanter Geschwindigkeit v_C ist $\Delta s = v_C \cdot \Delta t$, was im zugehörigen $s(t)$ -Diagramm dem Flächeninhalt („Höhe“ $v_C \times$ „Breite“ Δt) des Rechtecks unter dem Graphen entspricht.



Nimmt man an, dass sich die Geschwindigkeit stufenweise erhöht, ergibt sich die Gesamtstrecke als Summe der einzelnen Teilstrecken, welche jeweils durch den Flächeninhalt des Rechteckchens unter einer Stufe beschrieben werden können. (Einer davon ist im Bildchen nebenan mit Pünktchen hervorgehoben.) Die Gesamtstrecke entspricht dann immernoch der Gesamtfläche unter dem Graphen.



Werden die Stufen immer feiner, ändert sich an dieser Entsprechung nichts. Bei gleichmäßiger Stufung gleicht sich der Graph aber einer Ursprungsgeraden an, die eine gleichmäßig beschleunigte Bewegung gemäß $v = a \cdot t$ darstellt. Die aus immer schmäleren Rechteckchen zusammengesetzte Gesamtfläche gleicht sich der Dreiecksfläche unter dieser Geraden an.



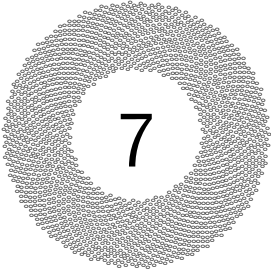
Sieht man das Dreieck als halbes Rechteck, erkennt man, dass sein Flächeninhalt dem Produkt $\frac{1}{2} \cdot \Delta t \cdot at$ als Gesamtstrecke s entspricht. Für einen Start zum Zeitnullpunkt schreiben wir für Δt einfach t .

Bei einer gleichmäßig mit a beschleunigten Bewegung aus der Ruhe heraus beträgt der in der Zeit t zurückgelegte Weg

$$s = \frac{1}{2}at^2$$

Damit ergibt sich nun auch wie erwartet:

$$W = v \cdot m \cdot \frac{\frac{1}{2}at^2}{t} = \frac{1}{2} \cdot v \cdot m \cdot at = \frac{1}{2} \cdot v \cdot m \cdot v = \frac{1}{2}mv^2$$



Licht und Schatten

7.1 Erfahrungen mit Licht

Es fällt schwer zu sagen, was *Licht* ist. Einfacher ist es, erst einmal zu erklären, was *kein Licht* bedeutet: Ohne Licht ist es dunkel, ohne Licht kann man auch mit den besten Augen nichts sehen.

- In einem dunklen Raum machen wir deshalb *Licht an*, um etwas zu sehen. Wir zünden z.B. eine Kerze an. Das so angemachte Licht wird offensichtlich von der brennenden Kerze erzeugt.

- Wir sehen mit unseren Augen Dinge erst bei Licht. Dies legt nahe, dass zum Sehen das Licht auch zu unseren Augen gelangen muss. Wir gehen also davon aus, dass Licht von Lichtquellen (wie der Kerze) ausgesandt wird und von unseren Augen empfangen und wahrgenommen wird.

- Eine Kerzenflamme kann von jeder Seite her gesehen werden. Die Flamme sendet also Licht in alle Richtungen aus.

- In diesem allseitigen Lichtschein der Flamme werden auch andere Gegenstände im Raum sichtbar. Also geht dann wohl auch von diesen Gegenständen Licht aus.

Wir unterscheiden: Lichtquellen, die von sich aus leuchten, und beleuchtete Körper. Beleuchtete Körper senden erst dann Licht aus, wenn sie selbst von Licht getroffen werden.

Viele Redewendungen aus dem Alltag verraten, dass die meisten Menschen eigentlich eine andere Vorstellung vom Sehen haben. Man spricht z.B. vom *Augenlicht*, so als ob das Auge irgendwelche Seh-Strahlen aussenden und damit die Umgebung abtasten würde. Ein Blick wird irgendwohin geworfen, ist stechend oder durchdringt ein Dunkel; man sieht, soweit das Auge reicht. Nach allen diesen Formulierungen geht etwas *vom Auge aus*, obwohl doch immer nur Licht *in das Auge hinein* fällt.

7.2 Lichtausbreitung

Wir haben die Vorstellung, dass Licht von Lichtquellen ausgesendet wird. Anders ausgedrückt: Licht breitet sich von Lichtquellen her aus. Wie hat man sich diese Ausbreitung genauer vorzustellen? Wie bei einer Menschenmenge in einem gerade geöffneten Kaufhaus zu Beginn des Sommerschlussverkaufs? Wie bei einer Grippe-Epidemie oder einem Gerücht? Wie bei Split-

tern einer Explosion?

Das Wichtigste über die Ausbreitung von Licht erfahren wir, wenn wir das Licht einer Lichtquelle daran hindern, sich überallhin auszubreiten: Dem Licht eines kleinen Lämpchens stellen wir auf einer Seite eine Wand mit Loch in den Weg. Die Wand lässt nur durch das kleine runde Loch Licht hindurch.

Tatsächlich muss es hier keine ganze Wand sein. Es genügt ein Schirm, der vom Lämpchen aus gesehen einen gewissen Raumbereich abdeckt. Einen solchen Schirm mit Loch nennt man eine **Lochblende**.



zu ergänzen:

- ausgeleuchteter Bereich

An welche Stellen im Raum hinter der Wand kommt nun noch Licht hin?

?

Und dazu zunächst: Wie stellen wir fest, wo Licht hinkommt? Wenn wir z.B. unsere Hand an eine Stelle halten, erkennen wir, ob die Hand vom Licht des Lämpchens angestrahlt leuchtet oder nicht. So könnten wir den gesamten Raum hinter der Wand Stelle für Stelle auf Lichteinfall hin untersuchen.

Einfacher geht es, indem wir Rauch oder Staub hinter die Wand blasen. Nur dort, wo Licht hinkommt, werden winzige Rauch- bzw. Staubpartikel angestrahlt und damit sichtbar. Beachte: Wir können das Licht, das sich an unseren Augen *vorbei* ausbreitet, nicht sehen! Licht ist zwar das Einzige, was unsere Augen wahrnehmen können, aber nur, wenn das Licht *in* das Auge fällt. Erst die beleuchteten Staubpartikel senden Licht auch *in* unsere Augen.

Wir erkennen: Der ausgeleuchtete Bereich hinter der

Wand ist begrenzt durch einen Kegel mit dem Lämpchen als Spitze. Man nennt einen so ausgeleuchteten Kegel einen **Lichtkegel**. Zeichne den Lichtkegel oben ein!

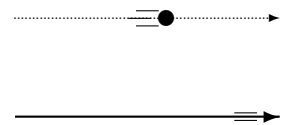
Verengt man das Loch in der Wand, wird auch der Lichtkegel schmaler. Je schmaler man den Kegel macht, desto weniger fällt auf, dass er hinter dem Loch immer breiter wird. Er erinnert mehr und mehr an einen geometrischen Strahl (=gerichtete Halbgerade) vom Lämpchen durch das Loch. Ein Lichtkegel kann bei dieser Vorstellung von **Lichtstrahlen** als ein Bündel solcher Strahlen angesehen werden. Wir nennen einen Lichtkegel daher meist ein **Lichtbündel**.

Und was wissen wir nun über die Ausbreitung von Licht? *Beobachten* können wir im Versuch Folgendes:

Licht gelangt nur zu Stellen, zu denen es von einer Lichtquelle aus eine geradlinige Verbindung ohne Hindernisse gibt.

Um uns diese Beobachtung plausibel zu machen, können wir uns vorstellen, dass Licht etwas ist, das sich geradlinig von Lichtquellen aus fortbewegt. Demnach beschreiben wir die *Wege* des Lichtes (•) mit Lichtstrahlen.

Wir können uns auch vorstellen, dass das Licht selbst strahlenförmig aus einer Lichtquelle heraustritt. Demnach beschreiben wir das *Licht selbst* (→) mit Lichtstrahlen.



In jedem Fall haben wir uns mit den Lichtstrahlen etwas über unsere gesicherte Beobachtung hinaus *dazugedacht*. Eine Beschreibung von Naturerscheinungen mit Hilfe solcher gedachter Elemente nennt man ein physikalisches **Modell**.

Modelle können die Beschreibung einer Naturerscheinung vereinfachen. Sie können aber darüber hinaus helfen, Voraussagen über künftige Abläufe zu machen. Ein Modell ist so gut, wie die auf ihm beruhenden Voraussagen zutreffen. Es wird sich zeigen, dass ein Modell mit Lichtstrahlen viele optische Phänomene gut beschreiben und voraussagen kann.

Lichtstrahlen sind gedachte Objekte. Man kann keine einzelnen Lichtstrahlen erzeugen, sondern immer nur Lichtbündel. Man kann zwar mit Lochblenden recht wenig auseinanderlaufende, strahlähnliche Lichtbündel erzeugen, aber in diesen Bündeln ist dann nur noch sehr wenig Licht enthalten. Schließlich hat man ja das meiste Licht, das aus der Quelle kommt, ausgeblendet.

Zur Untersuchung von Lichtstrahlen und zu anderen Zwecken hat man sich jedoch Lichtquellen konstruiert, die sehr schmale und dennoch lichtstarke Lichtbündel erzeugen: **Laser**. Da deren Lichtbündel fast gar nicht auseinanderlaufen, spricht man von Laserstrahlen.

Nach allem bisher Gesagten wird der ganze Raum um uns herum kreuz und quer von Licht durchströmt. Denn: Jeder sichtbare Gegenstand sendet in jede Richtung Lichtstrahlen aus. Wir nehmen jedoch nur die Lichtstrahlen wahr, die in unsere Augen fallen. Auf dem Weg zu uns müssen diese Strahlen jedoch ständig von anderen Lichtstrahlen durchkreuzt worden sein. Offensichtlich gilt:

Lichtstrahlen stören sich gegenseitig nicht, durchdringen einander ohne Wechselwirkung.

Die Tatsache, dass sich Licht so unbeirrt geradlinig ausbreitet, lässt sich nutzen. In Situationen, wo es darauf ankommt, dass etwas geradlinig verläuft, kann man dies mit Lichtstrahlen sicherstellen oder überprüfen. Lange gerade Tunnel können z.B. gebaut werden, indem man einem fest aufgestellten Laserstrahl einen Weg durch das Gebirge gräbt.

Wenn ein Körper den Blick auf einen anderen verdeckt, heißt dies, dass es keinen Lichtstrahl vom hinteren Körper zu unserem Auge gibt. Der vordere Körper

durchtrennt also jede geradlinige Verbindung zwischen hinterem Körper und unserem Auge. Die beiden Körper und das Auge liegen folglich auf einer Geraden.

Um zu prüfen, ob eine Reihe von dünnen Zaunpfosten einigermaßen gerade ist, genügt es daher schon, an der Reihe entlang zu schauen. Findet man keinen Standpunkt, von dem aus gesehen der vorderste Pfosten alle anderen Pfosten verdeckt, ist die Reihe nicht gerade.

7.3 Lichtgeschwindigkeit

Wenn wir sagen, dass Licht sich ausbreitet, denken wir dabei an eine Bewegung, die auch Zeit benötigt. Im Alltag erfahren wir jedoch nicht, dass Licht zum Ausbreiten Zeit benötigt. Oder hast Du schon einmal den Eindruck gehabt, nachdem Du eine Lampe links vor Dir eingeschaltet hast, diese Lampe mit dem linken Auge zuerst zu sehen? Sicher nicht. Wenn Licht dennoch Zeit zur Ausbreitung benötigt, können wir bereits sagen, dass sich Licht sehr schnell ausbreitet. Als Physiker wollen wir es natürlich genauer wissen: Wie schnell breitet sich Licht aus?

Um es vorweg zu sagen: Es ist recht aufwändig, die Lichtgeschwindigkeit zu bestimmen. Man muss dazu jedenfalls sehr kleine Zeitspannen und/oder sehr große Längen bestimmen. In den letzten Jahrhunderten schlugen deshalb einige Versuche zur Bestimmung der Lichtgeschwindigkeit fehl. Lange war deshalb nicht klar, ob Licht überhaupt Zeit benötigt, um sich auszubreiten.

1638 führten Galilei und ein Gehilfe einen Versuch aus, in dem bestimmt werden sollte, wie lange es dauert, bis Licht eine Strecke hin und zurück zurücklegt.

Sie stellten sich dazu nachts einige Kilometer voneinander entfernt auf. Jeder trug eine Laterne bei sich.

Laternen konnte man damals nicht an- und ausknipsen wie heute elektrische Glühlampen. Statt dessen kann man eine Laterne aber mit einem Schirm verdecken und wieder öffnen. Zu Beginn des Versuches hielten Galilei und sein Gehilfe ihre Laternen verdeckt.

Die Versuchsidee war nun folgende: Galilei öffnete seine Laterne. Deren Licht breitete sich bis zum Gehilfen aus. Sobald der Gehilfe dieses Licht sah, öffnete auch er seine Laterne. Deren Licht wiederum erreichte kurz darauf Galilei. Aus der Zeit für dieses Hin und Her und dem Abstand der beiden Männer wollte Galilei die Lichtgeschwindigkeit bestimmen. Galilei konnte die Lichtgeschwindigkeit so nicht bestimmen. Dies hat zwei Gründe:

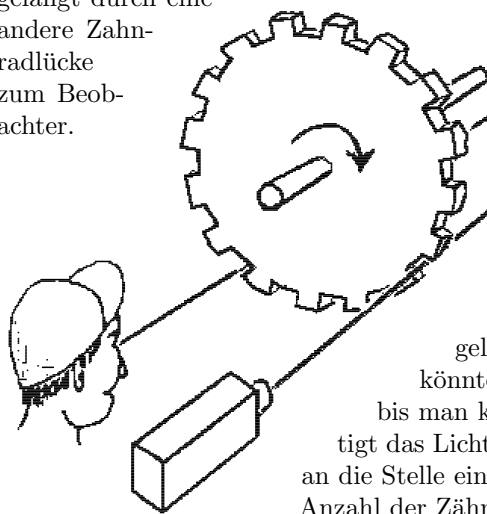
- 1) Licht ist so schnell, dass die Zeit, die es für ein paar Kilometer braucht, viel zu kurz ist, als dass man sie im 17. Jahrhundert hätte direkt messen können.
- 2) Die beiden Männer haben gewisse Reaktionszeiten: Bis jemand, der ein Licht sieht, eine Laterne geöffnet hat, ist viel mehr Zeit vergangen, als das Licht für seine Wege in diesem Versuch braucht. Reaktionszeiten eines Menschen schwanken außerdem unvorhersehbar viel zu stark, um Messungen der hier erforderlichen Präzision auszuführen.

Grundsätzlich ist Galileis Versuchsidee jedoch brauchbar. Man kann anstelle des Gehilfen einen Spiegel auf dem Mond verwenden, der ein einfallendes Laserlicht (ohne Reaktionszeit) zur Erde zurückwirft. Bei der Entfernung Erde-Mond von 380 000 km benötigt ein Laserlicht-Signal bis zu seiner Rückkehr 2,56 s. Daraus errechnet sich die Lichtgeschwindigkeit c zu

$$c = \frac{\text{Lichtweg}}{\text{Laufzeit}} = \frac{2 \cdot 380\,000 \text{ km}}{2,56 \text{ s}} \approx 300\,000 \frac{\text{km}}{\text{s}}$$

Die Lichtgeschwindigkeit c beträgt etwa $300\,000 \frac{\text{km}}{\text{s}}$.

1849 gelang es dem französischen Physiker Armand Hippolyte Louis Fizeau, die Lichtgeschwindigkeit rein irdisch zu messen. Seinen Versuchsaufbau zeigt diese Abbildung. Steht das Zahnrad in passender Stellung still, verläuft ein Lichtstrahl durch eine Zahnradlücke zu einem Spiegel, wird von diesem zurückgeworfen und gelangt durch eine andere Zahnradlücke zum Beobachter.



Dreht sich das Zahnrad, wird der Lichtstrahl abwechselnd von einem Zahn ausgeblendet und durch eine Lücke passieren gelassen. Hinter dem Zahnrad laufen daher Lichtblitze (Lichtstrahl-Abschnitte) zum Spiegel und zurück. Dreht sich das Zahnrad langsam, findet jeder Lichtblitz bei seiner Rückkehr das Zahnrad in fast der gleichen Stellung vor wie bei seinem Start. In der kurzen Zeit, in der der Lichtblitz zum Spiegel und zurück läuft, hat sich das langsame Rad kaum weitergedreht. Dreht sich das Zahnrad jedoch schnell genug, ist bis zur Rückkehr eines Blitzes bereits ein Zahn dorthin gelangt, wo der Blitz sonst durch eine Lücke zum Beobachter gelangen könnte. Man erhöht demnach die Drehgeschwindigkeit des Rades solange, bis man keine Lichtblitze mehr durch das Zahnrad sieht. Dann nämlich benötigt das Licht für Hin- und Rückweg genauso viel Zeit, wie es dauert, bis ein Zahn an die Stelle einer Lücke getreten ist. Diese Zeit lässt sich aus der Drehzahl und der Anzahl der Zähne berechnen.

Wissenschaftliche Zahldarstellung

In Aufgaben zur Lichtgeschwindigkeit treten häufig sehr große oder sehr kleine Maßzahlen für zurückgelegte Strecken bzw. benötigte Zeiten in den Standardeinheiten auf. Beispiele: Licht legt in einer Stunde 1 080 000 000 000 m zurück und benötigt für eine Strecke von 7,5 m nur 0,000 000 025 s.

Um mit solchen Zahlen generell bequemer umgehen und insbesondere sie auch bei begrenzter Stellenzahl eines Taschenrechners überhaupt ein- und ausgeben zu können, beschreibt man solche Zahlen als Produkt aus einer „handlichen“ (Komma-)Zahl und einer Zehnerpotenz. Dabei kann man „mal 10 hoch“ durch ein kleines oder großes E abkürzen. Für die Ein- und Ausgaben von Computern und Taschenrechnern wird noch statt eines Kommas ein Punkt gesetzt.

Es gilt also: $1\,080\,000\,000\,000 = 1,08 \cdot 1\,000\,000\,000\,000 = 1,08 \cdot 10^{12} = 1.08E12$ mit „Mantisse“ 1.08 und „Exponent“ 12. Es ist üblich, wie hier Mantissen zwischen 1 und 10 zu wählen. Natürlich sind aber auch z.B. 10.8E11 oder 108E10 zur Angabe derselben Zahl möglich.

Am Taschenrechner wird das E über eine Taste eingegeben, die je nach Geräte-modell z.B. **[EXP]**, **[E]** oder **[·10ⁿ]** heißen kann. In der Ausgabe taucht aus Platzgründen evtl. gar kein E auf, sondern es wird der Exponent von der Mantisse abgesetzt und/oder hochgestellt ausgegeben. Am Rand siehst Du ein Beispiel, welche Tasten-Eingabe welcher Ausgabe entspricht. (Nicht mit 1,08¹² verwechseln!)

Nun musst Du nur noch wissen, dass es auch negative Exponenten gibt. Im Einklang mit den Rechenregeln für Potenzen wird kurzerhand $a^{-n} := \frac{1}{a^n}$ definiert. Es ist also z.B. $10^{-4} = \frac{1}{10^4} = \frac{1}{10\,000} = 0,0001$ und $0,000\,000\,025 = 25 \cdot 0,000\,000\,001 = 25 \cdot 10^{-9} = 25E-9 = 2.5E-8$

Zur Eingabe des Vorzeichens des Exponenten kannst Du die Taste **[+/-]** zum Wechseln des Vorzeichens verwenden, wie es auf dem Seitenrand einige Varianten zur Eingabe derselben Zahl zeigen.

Falls Dir diese Erklärungen mit Zehnerpotenzen nicht gefallen, kannst Du die Sache auch so sehen: Mit dem Zusatz ...E+? hinter einer Zahl gibt man an, dass das Komma in der Zahl noch um ? Stellen nach rechts verschoben werden soll, mit E-? um ? Stellen nach links.

$$c \approx 3E8 \frac{\text{m}}{\text{s}}$$

1 · 0 8 E 1 2

↓
1.08 12

2 5 E ± 9

2 5 E 9 ±

2 · 5 E ± 8

7.4 Schatten

In den folgenden Versuchen erzeugen wir Schatten auf einer **Mattscheibe**. Dies ist eine flache Platte mit heller, matter Oberfläche. Auf einer solchen Oberfläche sieht man besonders deutlich, wo Licht auf sie fällt. Eine Mattscheibe nennen wir auch einen Schirm.

V 7.1 einfacher Schattenwurf

Aufbau:

Wir stellen
ein Lämpchen,
einen Gegenstand
und eine Mattscheibe
hintereinander auf.



zu ergänzen:

- Lichtstrahlen
- hell/dunkel auf der Mattscheibe

Durchführung: Wir dunkeln den Raum ab und schalten das Lämpchen ein.

Beobachtung: Ein Teil der Mattscheibe wird erleuchtet. Es bleibt nur eine Fläche dunkel, deren Form dem Umriss des Gegenstandes entspricht.

Erklärung: Von der Lichtquelle gehen Lichtstrahlen aus. Den Gegenstand können Lichtstrahlen jedoch nicht durchdringen. Auf der Mattscheibe treffen deshalb nur die Lichtstrahlen ein, die am Gegenstand vorbei gegangen sind.

Ein Schatten ist demnach einfach ein unbeleuchteter Bereich, der von beleuchteten Bereichen begrenzt ist. Ähnlich wie ein Loch als umrandetes Nichts erklärt werden kann, ist ein Schatten ein „umleuchtetes Dunkel“.

In unserer Sprache behandeln wir Schatten ebenso wie Löcher dennoch als etwas eigenständig existierendes. Wir schneiden *Löcher in* ein Blatt Papier, obwohl wir eigentlich ein Stück *Papier herausschneiden*. Genauso wird ein *Schatten an* die Wand geworfen, obwohl eigentlich nur *Licht von* der Wand abgehalten wird.

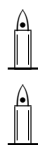
Eine Blende wirft eigentlich auch nur Schatten.

Nicht jeder Gegenstand eignet sich gleich gut, um Schatten zu erzeugen. Verschiedene Gegenstände lassen Licht verschieden gut durch. Je nachdem, ob ein Gegenstand kein Licht, einen Teil des einfallenden Lichtes oder praktisch das ganze Licht durchlässt, nennen wir den Gegenstand **undurchsichtig**, **durchscheinend** bzw. **durchsichtig**.

V 7.2 Schatten bei zwei Lämpchen

Aufbau: wie zuvor, jedoch mit zwei Lämpchen.

Durchführung: wie zuvor



zu ergänzen:

- Lichtstrahlen
- Helligkeiten auf der Mattscheibe

Beobachtung: Auf dem Schirm bleibt ein Teil ganz dunkel. Im daran angrenzenden Teil wird der Schirm nur schwach erleuchtet.

Erklärung: Den ganz dunklen Bereich trifft überhaupt kein Licht. Die halbdunklen Bereiche trifft nur das Licht jeweils eines Lämpchens. Die ganz erleuchteten Bereiche trifft Licht beider Lämpchen.

Man nennt den ganz dunklen Bereich den **Kernschatten**, und den halbdunklen nennt man **Halbschatten**.

Bisher haben wir Strahlengänge immer so gezeichnet, als ob das Lämpchen nur ein einziger leuchtender Punkt (im geometrischen Sinne) sei. Tatsächlich enthält seine Glühwendel aber unendlich viele dicht liegende leuchtende Punkte. Das Licht kommt daher nicht von einem einzigen Punkt, sondern aus einem ausgedehnten Bereich.

V 7.3 Schatten bei ausgedehnter Lichtquelle

Aufbau: wie zuvor, jedoch mit großflächiger Lichtquelle

Durchführung: wie zuvor



Beobachtung: Es gibt nun einen fließenden Übergang von einem ganz dunklen Bereich zum voll erleuchteten Teil des Schirmes.

Erklärung: Den ganz dunklen Teil trifft überhaupt kein Licht, den ganz hellen Bereich trifft Licht von allen leuchtenden Punkten der Lichtquelle. An jede Stelle des Übergangsbereiches trifft nur das Licht aus einem bestimmten Bereich der Lichtquelle. Je nach Lage der Stelle ist dieser Bereich mehr oder weniger groß.

zu ergänzen:

- Lichtstrahlen
- Helligkeiten auf der Mattscheibe

Den Übergangsbereich vom Kernschatten zum Hellen nennt man den **Übergangsschatten**.

7.5 Mond

Du weißt wahrscheinlich, dass die Erde, der Mond und die Sonne sehr große, in etwa kugelförmige Körper sind, die sich im Weltraum bewegen. Ein paar Dinge, die man darüber noch wissen sollte, findest Du hier zusammengefasst:

- Die **Sonne** ist sehr groß, sehr heiß und sie leuchtet.
- Es gibt sehr viele weitere solcher Sonnen im Weltraum. Da diese Sonnen jedoch von uns sehr viel weiter entfernt sind als unsere Sonne, sehen wir sie viel kleiner und weniger hell. Wir nennen sie **Sterne**. Tagsüber ist der Himmel so hell, dass man die Sterne gar nicht erkennen kann. (vergleiche: helles Pünktchen auf hellem Hintergrund) „Sonne“ und „Stern“ sind für Astronomen gleichbedeutend.
- Um unsere Sonne herum kreisen in unterschiedlichen Entfernungen die **Planeten**. Sie sind viel kleiner als die Sonne und leuchten (so gut wie) nicht. Vom sonnennächsten zum fernsten genannt heißen die Planeten: Merkur, Venus, Erde, Mars, Jupiter, Saturn, Uranus, Neptun, Pluto. (Pluto ist allerdings so klein, dass er offiziell heute nicht mehr als Planet bezeichnet wird.) Die Zeit, die die Erde benötigt, um einmal die Sonne zu umkreisen, ist ein Jahr.
- Planeten können ihrerseits von **Monden** umkreist werden. Diese sind noch einmal kleiner und leuchten ebenfalls nicht selbst. Die Erde wird von einem einzigen

Mond umkreist. Ein Umlauf des Mondes um die Erde dauert etwa 1 Monat, genauer 27,3 Tage.

- Jeder Himmelskörper dreht sich auch um sich selbst. Die Zeit, in der sich die Erde einmal um sich selbst dreht, ist ein Tag.

Der Mond benötigt für eine Umdrehung genauso lange wie für eine Erdumrundung, nämlich einen Monat. Du kannst dies daran erkennen, dass der Mond uns Erdlingen immer dieselbe Seite zeigt. Dies ist kein Zufall, sondern Folge physikalischer Gesetzmäßigkeiten, nach denen die Erddrehung, die Mondrotation und der Mondumlauf sich immer mehr aneinander anpassen.

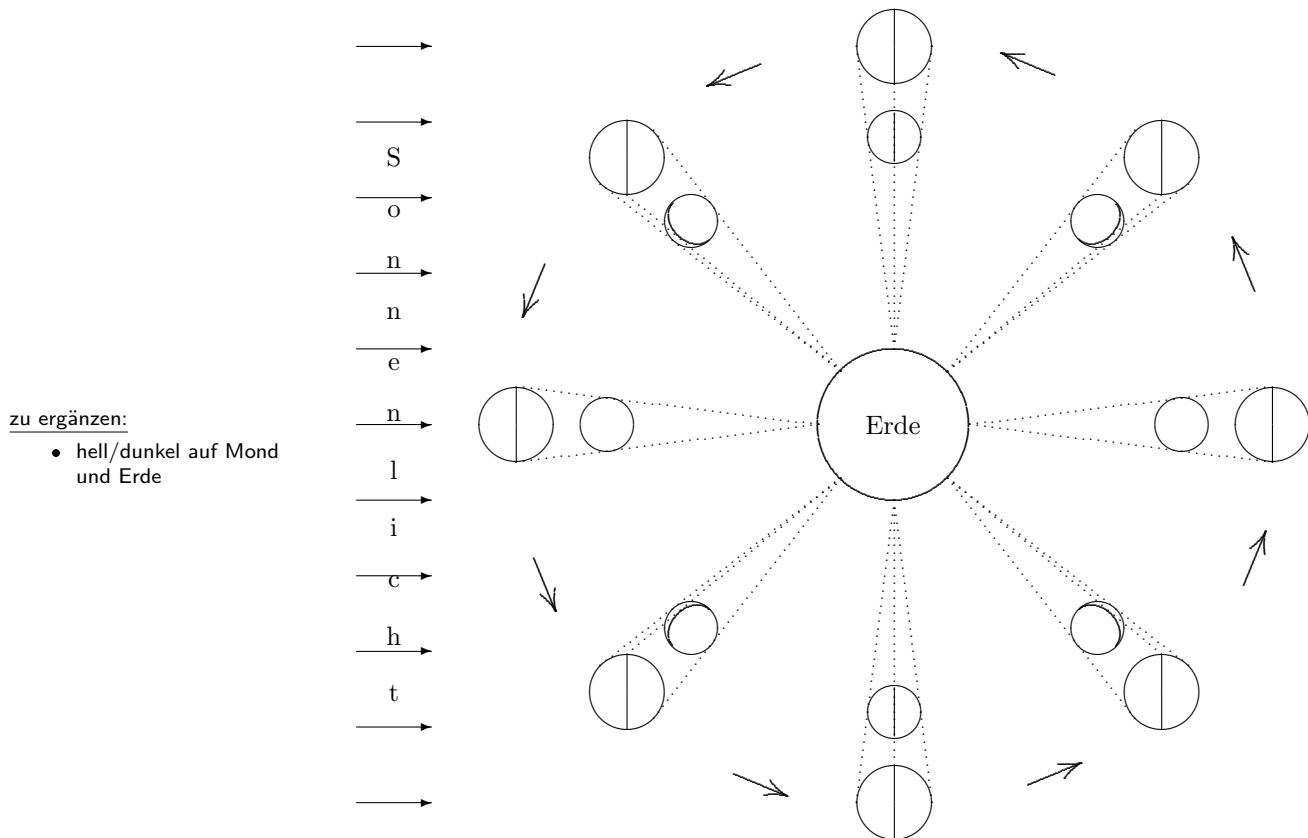
O bwohl unser Mond nicht selbst leuchtet, sehen wir ihn. Er wird nämlich von unserer Sonne beleuchtet. Und ebenso, wie die Erdkugel immer nur auf einer Hälfte beschienen wird, – auf dieser Hälfte ist dann *Tag* – so wird auch der Mond immer nur auf der Seite beleuchtet, die der Sonne zugewandt ist. Je nachdem, wie Erde und Mond gerade stehen, sehen wir von der Erde aus nur einen bestimmten Teil dieser beleuchteten Seite des Mondes. So kommt es zu den **Mondphasen**:

Neumond: Die unbeleuchtete Seite ist der Erde zugewandt. Wir sehen fast gar nichts vom Mond.

zunehmender Mond: Die beleuchtete Seite wird uns von Tag zu Tag mehr zugewandt, wir sehen mehr und mehr von der beleuchteten Seite. Eine nach rechts gebogene Sichel „)“ wird immer breiter.

Vollmond: Die beleuchtete Seite ist uns zugewandt. Wir sehen diese ganze Seite, eben einen „vollen“ Mond.

abnehmender Mond: Die beleuchtete Seite ist uns von Tag zu Tag mehr abgewandt. Der volle Mond schrumpft immer weiter zu einer in der Mitte nach links gebogenen Sichel.



Das Bild oben zeigt einen Blick auf die Nordhalbkugel der Erde. Die Sonne stehe

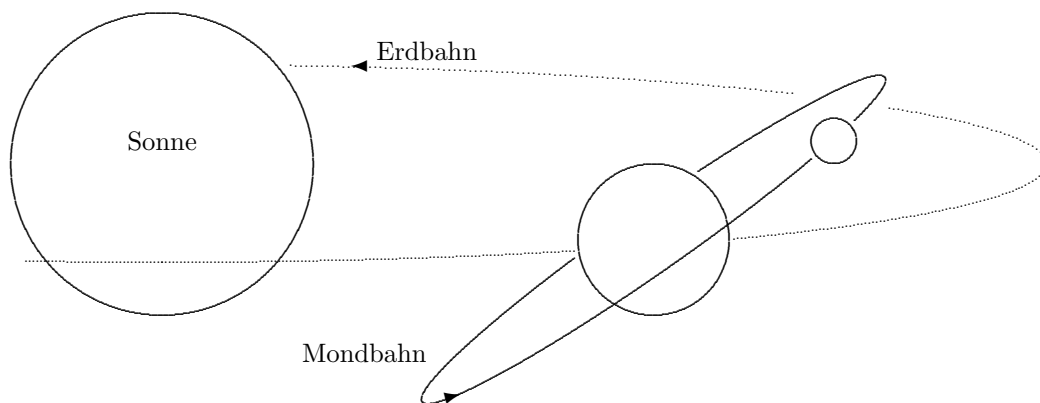
links von diesem Bild. Sie ist so weit entfernt, dass ihre Lichtstrahlen im relativ kleinen Bereich von Erde und Mond praktisch parallel zueinander verlaufen.

Eingezeichnet ist außen der Mond zu verschiedenen Zeiten eines Umlaufes. Zwischen Mond und Erde ist jeweils eingezeichnet, wie ein Beobachter auf der Erde den Mond in dieser Stellung etwa sehen würde. Die gepunkteten Linien sind die Sichtlinien vom Beobachter zu den Rändern des Mondes bzw. zur Hell-Dunkel-Grenze.

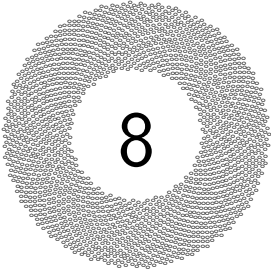
Das Bild ist allerdings noch nicht ganz fertig. Färbe doch alle angestrahlt hellen Flächen mit gelb, die dunklen mit grau!

Du fragst Dich vielleicht bei der Vollmond-Stellung ganz rechts, wie die linke Mondseite beschienen werden kann, wo doch die Erde das Sonnenlicht abhalten müsste. Wenn in dieser Stellung Sonne, Erde und Mond genauso auf einer Geraden liegen, wie es die Zeichnung andeutet, liegt der Mond tatsächlich im Schatten der Erde. Der Mond ist dann nur noch (aufgrund Dir noch nicht bekannter Effekte) ganz schwach erleuchtet. Dieses Ereignis nennt man deshalb eine **Mondfinsternis**.

Die Bahn des Mondes um die Erde verläuft jedoch nicht genau in der Papierebene, sondern ein wenig (5°) dazu geneigt. Der Mond könnte z.B. in der Vollmond-Stellung gerade etwas oberhalb der Zeichenebene stehen. Dann würde ihn das Sonnenlicht treffen, das über den Nordpol hinweg an der Erde vorbei läuft. Das folgende Bild stellt diese Neigung der Mondbahn stark übertrieben dar. Es ist nur für die Durchmesser von Erde und Mond maßstabsgetreu. Im gleichen Maßstab müsste ohne Perspektive der Mond etwa 60 cm weit von der Erde entfernt gezeichnet sein, die Sonne als Kreis mit 2,20 m Durchmesser 235 m entfernt.



Betrachten wir die linke Seite des Bildes zu den Mondphasen: Wenn der Mond in der Neumond-Stellung links gerade in der Papierebene steht, liegen Sonne, Mond und Erde ebenfalls auf einer Geraden, diesmal mit dem Mond zwischen Sonne und Erde. Nun wirft der Mond einen Schatten auf die Erde! Wo der Schatten hinfällt, wird es dunkel. Die Menschen in diesem Bereich sehen, wie der Mond vor der Sonne vorbeizieht und dabei die Sonne verdeckt. Dieses nur wenige Stunden andauernde Ereignis ist eine **Sonnenfinsternis**.



Lochkamera

8.1 Aufbau

Unser bisheriges Wissen genügt bereits, um die einfachst gebaute Kamera zu verstehen: die Lochkamera. Wie der Name bereits sagt, ist das Wichtigste an dieser Art Kamera ein Loch.

V 8.1 Strahlengang durch ein Loch

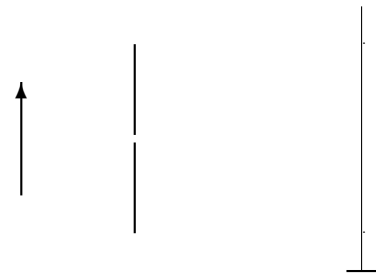
Aufbau: Ein leuchtender Gegenstand (hier als aufrecht stehender Pfeil gezeichnet) steht vor einer engen Lochblende. Auf der anderen Seite der Lochblende steht eine Mattscheibe.

Durchführung: Wir dunkeln den Raum ab.

Beobachtung: Auf der Mattscheibe ist ein Bild des Gegenstandes zu erkennen. Das Bild ist jedoch

zu ergänzen:

- Lichtstrahlen
- Bild auf der Mattscheibe



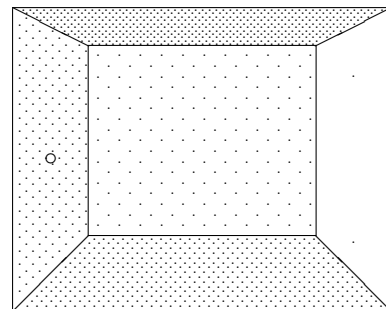
Erklärung: Auf die Mattscheibe kann Licht aus dem Raum vor der Blende praktisch nur durch das Loch fallen. Das Licht, das auf eine bestimmte Stelle des Schirmes fällt, kann daher nur aus Richtung Loch kommen, und damit nur von der Stelle des Gegenstandes, die in dieser Richtung liegt.

Ist diese Stelle des Gegenstandes z.B. rot, sendet sie rotes Licht aus und die entsprechende Stelle der Mattscheibe wird nur rot erleuchtet. So nimmt jede Stelle der Mattscheibe die Farbe einer bestimmten Stelle des Gegenstandes an und ein Bild entsteht.

Für die Bildentstehung ist es wichtig, dass zu jedem Punkt der Mattscheibe möglichst *nur* das Licht *eines* einzigen Punktes des Gegenstandes gelangt. Deshalb mussten wir beim letzten Versuch auch den Raum abdunkeln: Andernfalles wäre auch Licht anderer Herkunft auf die Mattscheibe gelangt und hätte das Bild überstrahlt.

Statt den Raum abzudunkeln kann man auch den Raum zwischen Lochblende und Mattscheibe in ein undurchsichtiges Gehäuse packen. Auch dies verhindert, dass Licht auf anderen Wegen als durch das Loch auf den Schirm gelangt.

Eine Lochkamera ist dementsprechend einfach ein abgeschlossener Kasten mit Loch als Blende auf der einen Seite und Mattscheibe oder Fotofilm auf der anderen Seite. Folgende Zeichnung zeigt einen Blick in den Kasten hinein. Ergänze darin das Bild der Kerze sowie die abbildenden Lichtstrahlen von einigen markanten Punkten der Kerze aus!



Bei der Mattscheibe gibt es ein kleines Problem: Das Bild entsteht *innen* auf einer Seite des Kartons. Wenn Du das Bild auch von *außen* sehen willst, brauchst Du eine durchscheinende Mattscheibe. Ersetze dazu einfach die „Bild-Kartonseite“ durch ein Stück Butterbrotpapier o.ä.!

zu ergänzen:

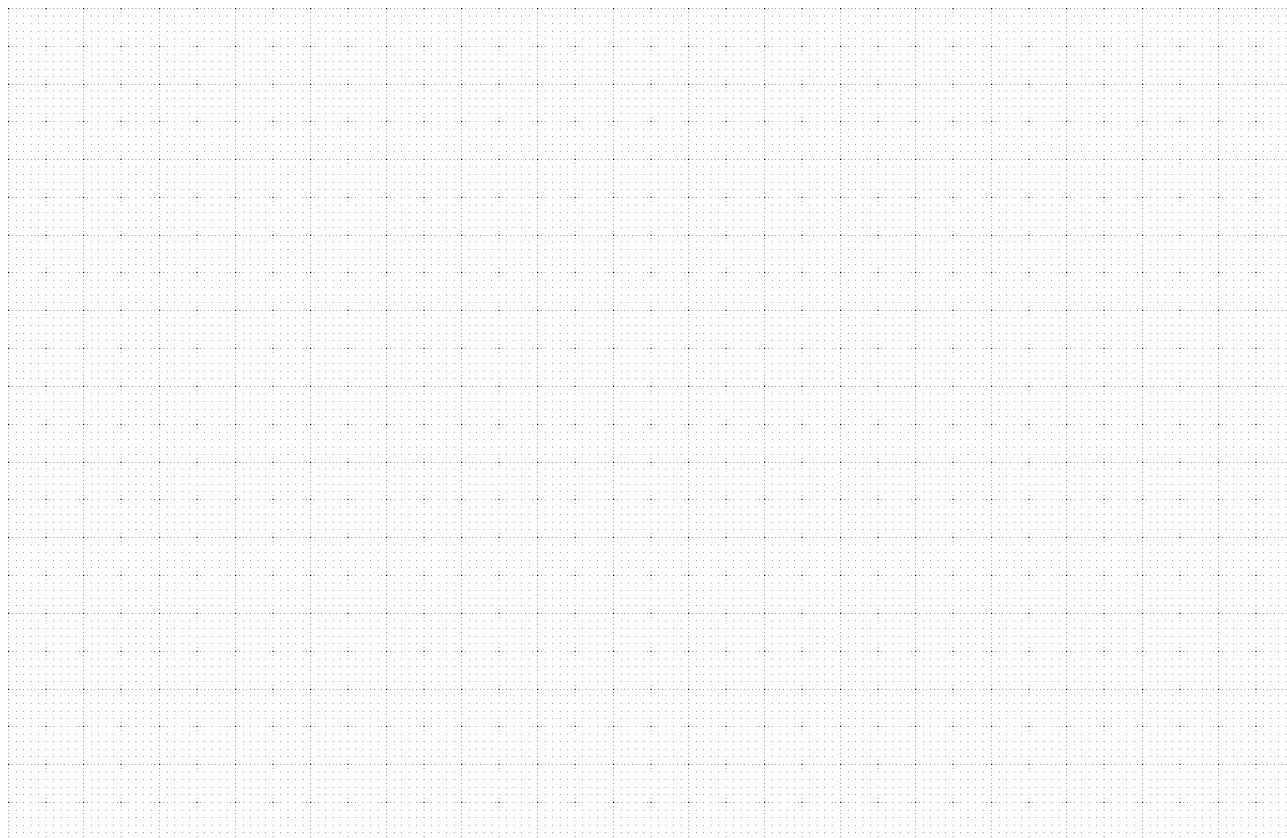
- Lichtstrahlen
- Bild auf der Mattscheibe

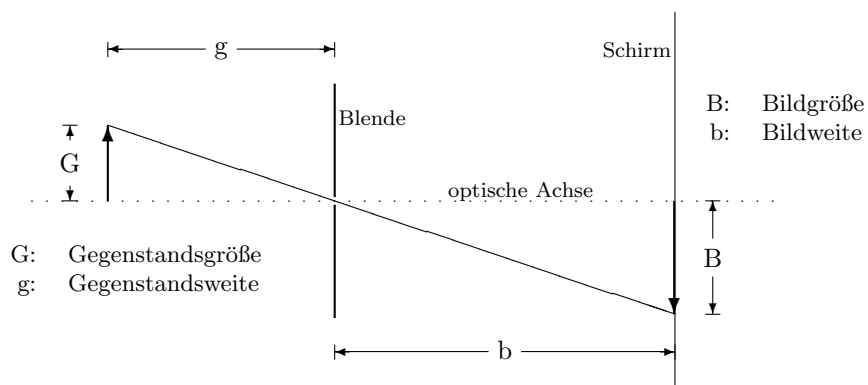
8.2 Abbildungsgleichung

Je nachdem, wie weit Gegenstand, Lochblende und Schirm voneinander entfernt stehen, ergeben sich verschieden große Bilder. Um die Zusammenhänge zwischen diesen Entfernungen und den Größen von Gegenstand und Bild zu ermitteln, benutzen wir die Bezeichnungen, die im folgenden Bild erklärt sind.

Um zeichnerische Auswertungen zu erleichtern, zeichnen wir mitten durch die Lochblende die sogenannte optische Achse des Systems und betrachten nur Gegenstände, die auf der optischen Achse stehen. Als Gegenstand zeichnen wir wieder einen Pfeil.

Indem man einen Gegenstand auswählt, legt man die Gegenstandsgröße G fest. Indem man den Gegenstand, die Lochblende und den Schirm aufstellt, legt man die Weiten g und b fest. Bei jedem Werte-Satz aus G , g und b stellt sich eine gewisse Bildgröße B ein.





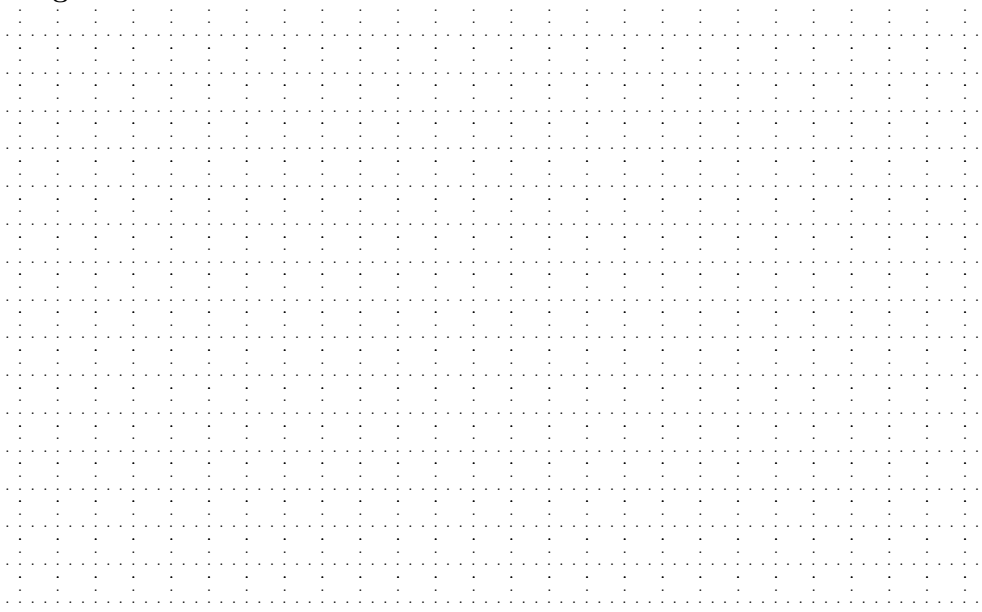
V 8.2 Bildgrößen

Aufbau: wie oben skizziert

Durchführung: Wir stellen systematisch Werte-Sätze aus G , g und b her und stellen dazu jeweils B fest. (Versuch kann auch zeichnerisch simuliert werden.)

Messwerte:

Ergebnis:



8.3 Bildqualität

Auf Fotofilm lässt sich das Bild in einer Lochkamera dauerhaft festhalten. Das eintreffende Licht löst nämlich chemische Reaktionen auf dem Filmmaterial aus. Zusammen mit den chemischen Reaktionen bei der Filmentwicklung führt dies dazu, dass jede Stelle des Films die Farbe und Helligkeit (– bei Schwarz-Weiß-Fotografie nur die Helligkeit –) des dort aufgetroffenen Lichtes annimmt.

So wird auf den Film mit Licht ein Bild geschrieben. Nach den griechischen Wörtern *phos*=Licht und *graphein*=schreiben heißen solche Bilder Fotografien. Das Wort Kamera stammt übrigens aus dem Lateinischen: *camera obscura* heißt *abgedunkelte Kammer*, und genau das ist eine Lochkamera ja!

Bilder, die man mit einer Lochkamera gemacht hat, sind allerdings meist nicht so gut wie Bilder aus einem „richtigen“ Fotoapparat: Sie sind recht unscharf oder lichtschwach. Wie kann man ein Lochkamera-Bild lichtstärker, heller machen? Offensichtlich muss dazu mehr Licht auf die Mattscheibe gelangen. Drei Möglichkeiten bieten sich dazu an:

- Man sorgt dafür, dass der Gegenstand mehr Licht aussendet, indem man ihn intensiver beleuchtet.
- Man lässt mehr Licht auf den Fotofilm einwirken, indem man den Film länger dem Licht aussetzt. Allerdings darf sich der Gegenstand in dieser längeren Zeitspanne nicht bewegen, weil sich ja sonst auch das Bild während der Aufnahme verändert.
- Man macht das Loch der Kamera größer: Durch ein größeres Loch kommt natürlich auch mehr Licht.

V 8.3 Änderung der Blendenöffnung

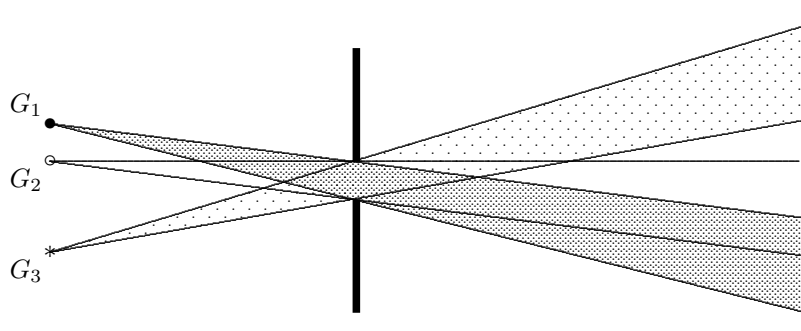
Aufbau: Leuchtender Gegenstand, Lochblende mit verstellbar großem Loch und Schirm auf einer optischen Bank.

Durchführung: Wir verändern den Durchmesser des Lochs.

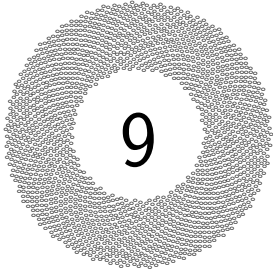
Beobachtung: Bei größerem Loch wird das Bild heller, allerdings auch verschwommener, unschärfer.

Erklärung: Wir haben bisher so getan, als ob nur einzelne Lichtstrahlen eines Gegenstandspunktes durch das Loch träten. Tatsächlich treten durch das Loch aber immer Lichtbündel, wenn auch schmale.

Zu jedem Punkt des Gegenstandes gibt es daher nicht nur einen Bildpunkt, sondern einen ausgedehnten Fleck auf dem Bild. Die Flecken benachbarter Gegenstandspunkte überlappen sich, (siehe G_1 und G_2) verschwimmen ineinander, das Bild wird unscharf. Je größer das Loch, desto größer die Flecken, desto größer die Überlappung, desto unschärfer das Bild.



(nicht nur) Bei einer Lochkamera sind daher Helligkeit und Schärfe einander entgegengesetzte Ziele.

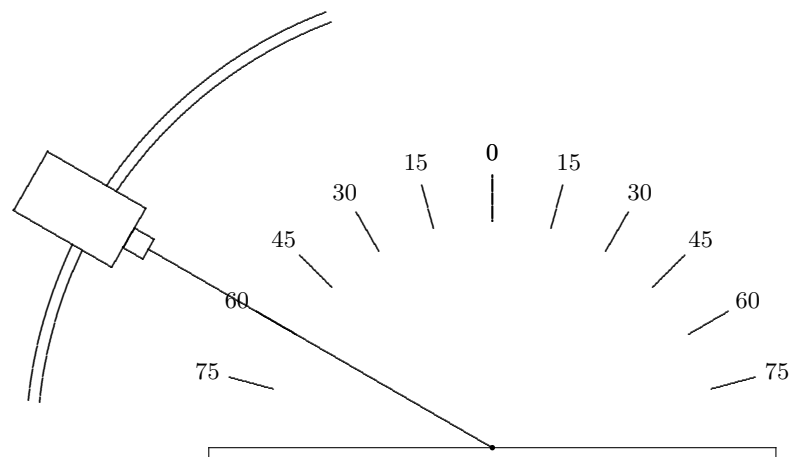


Spiegelbilder

9.1 Reflexion an einem ebenen Spiegel

V 9.1 Reflexion eines Lichtstrahls

Aufbau: Ein Lichtstrahler ist schwenkbar so montiert, dass sein Lichtstrahl unter einstellbarem Winkel auf eine feste Stelle eines Spiegels trifft. Der Lichtstrahl streift an einer Tafel mit Winkeleinteilung vorbei, so dass sein Verlauf sichtbar wird.



zu ergänzen:

- Strahlenverläufe

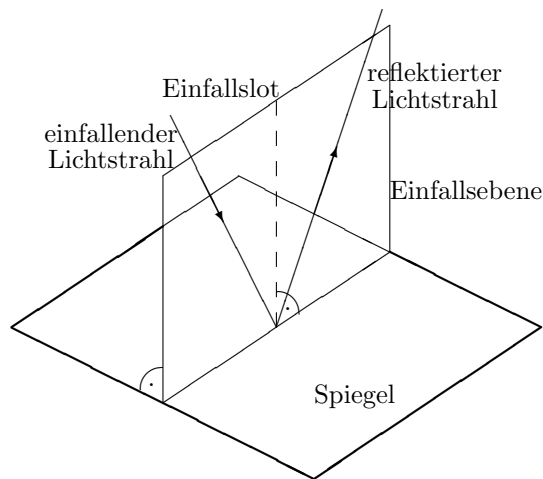
Durchführung: Für verschiedene Winkel des einfallenden Strahls lesen wir ab, unter welchem Winkel der gespiegelte Strahl ausfällt.

Beobachtung: siehe Zeichnung (Zeichne die beobachteten Strahlenverläufe ein. Verwende bei jeder Einstellung eine andere Farbe.)

Um das Ergebnis dieses Versuches zu formulieren, führen wir noch ein paar Bezeichnungen ein. Zeichne in folgendes Bild noch ein:

Einfallswinkel α_E = Winkel zw. einfallendem Strahl und Einfallslot

Reflexionswinkel α_R = Winkel zw. reflektiertem Strahl und Einfallslot



An einer spiegelnden Oberfläche wird ein einfallender Lichtstrahl in ganz bestimmter Richtung zurückgeworfen. Vom lateinischen Verb *re-flectere*= zurückbiegen her nennt man diesen Vorgang

Reflexion.

Reflexionsgesetz:

Einfallender Strahl, Einfallslot und reflektierter Strahl verlaufen in einer Ebene, die senkrecht auf der Spiegelebene steht. Einfalls- und Reflexionswinkel sind gleich.

Reflexion beobachten wir nur an sehr glatten Oberflächen. Beispiele:

9.2 Streuung

Eine raue Oberfläche kann unter dem Mikroskop wie eine zerklüftete Gebirgslandschaft aussehen. Ein noch so schmales Lichtbündel trifft dann auf ganz unterschiedlich gerichtete Flächenstückchen auf. Es werden daher in alle Richtungen Teile des einfallenden Lichtes zurückgeworfen. Dieses im Ganzen ungerichtete Zurückwerfen nennt man **Streuung** oder **diffuse Reflexion**. (lat.: *diffundere* = auseinandergießen)

9.3 Entstehung von Spiegelbildern

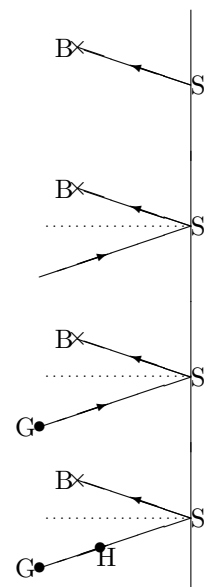
Wenn ein Beobachter B auf eine Stelle S eines Spiegels blickt, empfängt sein Auge von dort alle Lichtstrahlen, die von dort in Richtung auf das Auge verlaufen.

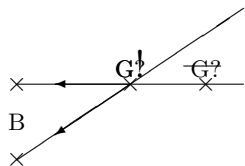
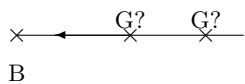
Vom Spiegel gehen nur reflektierte Strahlen aus. Die empfangenen Lichtstrahlen müssen also zuvor auf den Spiegel eingefallen sein, und zwar aus einer Richtung, die sich aus dem Reflexionsgesetz ergibt.

Verfolgt man diese Richtung zurück, trifft man wohl früher oder später auf einen Gegenstandspunkt G. Der Beobachter empfängt demnach aus S das Licht, das ursprünglich von G ausgegangen ist. So ist klar, dass der Beobachter überhaupt etwas von Dingen vor dem Spiegel im Spiegel sieht.

Aus dieser Überlegung geht noch nicht hervor, in welcher Entfernung der Beobachter den Gegenstandspunkt G sieht. Der von B empfangene Lichtstrahl könnte ebensogut wie von G auch z.B. von H ausgegangen sein.

Befinden sich in G und H tatsächlich Gegenstände, verdeckt aus der Sicht von B der Gegenstand in H denjenigen in G.





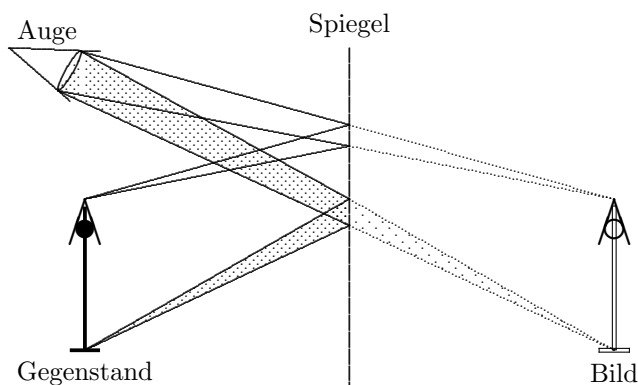
Generell kann ein Beobachter B die Entfernung zu einem Gegenstandspunkt G nicht erkennen, wenn er nur *einen* Lichtstrahl von G auffängt. Er weiß mit einem Lichtstrahl nur, auf welcher Geraden G liegen muss.

Fängt der Beobachter einen zweiten Strahl von G auf, kennt er eine zweite Gerade, auf der G liegen muss. Da G auf beiden Geraden liegen muss, muss G folglich der Schnittpunkt der beiden Geraden sein.

Am genauesten kann das Gehirn des Beobachters den Schnittpunkt bestimmen, wenn die aufgefangenen Lichtstrahlen möglichst verschiedene Richtungen haben, wenn also die auffangenden Stellen möglichst weit auseinander liegen. Dies ist bei den meisten Lebewesen mit einem Auge der Grund dafür, dass in einigen cm Entfernung von diesem Auge noch ein zweites Auge angebracht ist.

In den folgenden Zeichnungen sehen wir jedoch von diesen biologischen Details ab. Wir zeichnen Beobachter als ein einziges Auge, das dafür allerdings ausgedehnt genug ist, um an verschiedenen Stellen Lichtstrahlen zu empfangen.

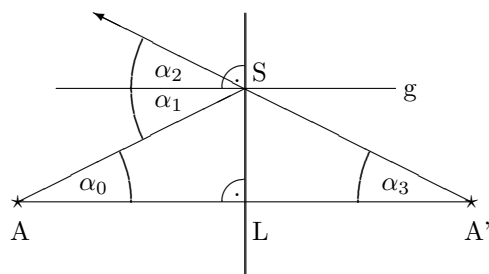
Das folgende Bild zeigt (links vom Spiegel) die Lichtbündel, die von der Spitze bzw. dem Fuß eines Gegenstandes ausgehen und nach Reflexion am Spiegel in das Auge eines Beobachters gelangen. Alle anderen Lichtstrahlen, die von denselben Gegenstandspunkten ausgehen, gelangen nicht in das Auge und sind daher für den Beobachter unerheblich.



Wenn das Gehirn unseres Beobachters nun wie gewohnt ermittelt, woher die eingetroffenen Lichtstrahlen wohl kommen, wird es glauben, dass die Lichtstrahlen von dem rechts eingezeichneten Pfeil kämen. Man erhält die genaue Lage z.B. der scheinbaren Pfeilspitze, indem man die entsprechenden Lichtstrahlen vom Auge aus rückwärts über den Spiegel hinaus verlängert.

Wie wir gleich beweisen werden, schneiden sich alle diese Verlängerungen in einem einzigen Punkt, eben dem Bildpunkt der Pfeilspitze. Die Lichtstrahlen der Pfeilspitze gelangen also genauso in das Auge des Beobachters, als ob sie von diesem einen Schnittpunkt der rückwärtigen Verlängerungen kämen.

Für jeden Punkt des Gegenstandes gibt es einen solchen Bildpunkt. Aus obiger Zeichnung und unserer Erfahrung mit Spiegelbildern vermuten wir, dass Bild und Gegenstand symmetrisch zur Spiegelebene liegen. Mit ein wenig Dreieckslehre lässt sich dies sogar beweisen:



Von Punkt A gehe ein Lichtstrahl in willkürlicher Richtung $\alpha_0 (\neq 0^\circ)$ aus, bezogen auf die Lotgerade AL von A auf die Spiegelebene. Im Punkt S des Spiegels wird er reflektiert. Die Gerade g ist als Einfallslot eingezeichnet. Die rückwärtige Verlängerung des reflektierten Strahles schneide die Gerade AL in A'.

$$\text{Nun gilt: } \left\{ \begin{array}{ll} \alpha_1 = \alpha_0 & (\text{Wechselwinkel an Parallelen}) \\ \alpha_2 = \alpha_1 & (\text{Reflexionsgesetz}) \\ \alpha_3 = \alpha_2 & (\text{Stufenwinkel an Parallelen}) \end{array} \right\} \Rightarrow \alpha_3 = \alpha_0$$

Die Dreiecke ALS und A'LS sind nach dem Kongruenzsatz SWW kongruent. (Übereinstimmungen siehe Rand)

Aus der Kongruenz folgt: $|\overline{AL}| = |\overline{A'L}|$

Dies bedeutet, dass A' und A achsensymmetrisch zur Geraden LS liegen, und zwar *unabhängig* von α_0 . Diese Unabhängigkeit bedeutet:

$$\sphericalangle ALS = \sphericalangle A'LS \quad (= 90^\circ)$$

$$\alpha_0 = \alpha_3 \quad (\text{siehe oben})$$

$$|\overline{LS}| = |\overline{LS}| \quad (\text{gemeinsame Seite})$$

Für jeden reflektierten Lichtstrahl aus dem Gegenstandspunkt A verläuft die rückwärtige Verlängerung durch den Bildpunkt A'.

Beachte, dass es daher für die Lage des Bildes überhaupt keine Rolle spielt, welche reflektierten Strahlen der Beobachter auffängt. Das Bild ist folglich das gleiche, egal, wo der Beobachter steht. Jeder Beobachter sieht das Bild an gleicher Stelle und kann es nicht ohne Weiteres vom Original-Gegenstand unterscheiden.

Aufgrund der Symmetrie zwischen Gegenstand und Bild ist offensichtlich, dass Bild und Gegenstand stets gleich groß sind.

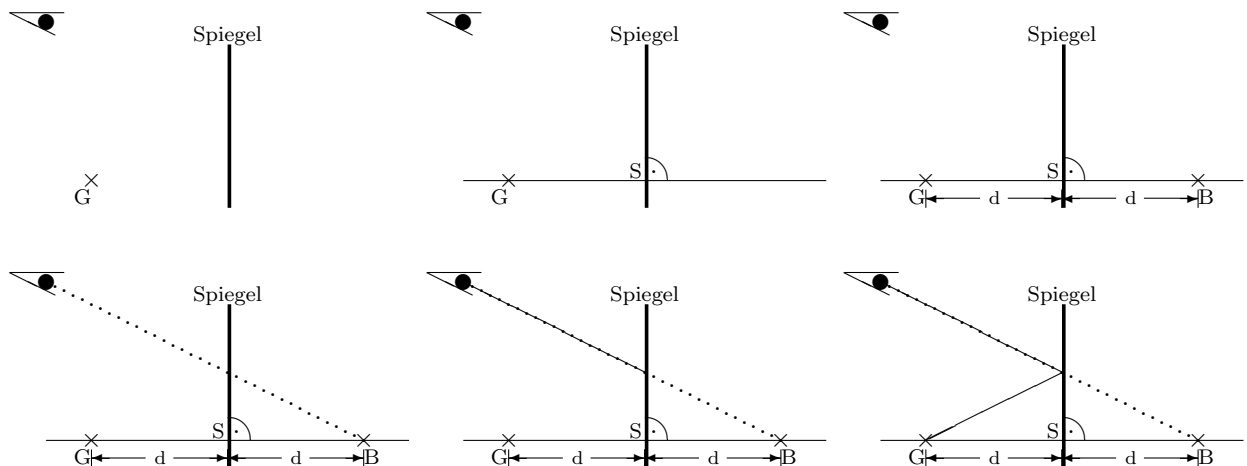
9.4 Konstruktion von Spiegelbildern

Aufgrund der Symmetrie lassen sich Spiegelbilder leicht konstruieren:

- Wähle markante Gegenstandspunkte, von deren Bildpunkten aus das Bild gut ergänzt werden kann. Ermittle zu jedem gewählten Gegenstandspunkt G den Bildpunkt B wie folgt:
- Konstruiere durch den Gegenstandspunkt die Gerade, die senkrecht zur Spiegelebene verläuft. (Lotgerade) Ihr Schnittpunkt mit der Spiegelebene heiße S.
- Trage auf dieser Geraden den Abstand d zwischen S und G von S aus in die andere Richtung ab und erhalte B.

Nach dem Bild lässt sich dann auch der Strahlenverlauf vom Gegenstand zu einem Beobachter leicht konstruieren: Zeichne zunächst dünn, geradlinig und durch den Spiegel hindurch die Strahlen, die vom Bild zum Beobachter verlaufen würden. Ab dem Spiegel stellen diese Linien die tatsächlichen Lichtstrahlen nach der Reflexion dar und können dicker gezeichnet werden. Ergänze die noch nicht reflektierten Strahlen, indem Du die Schnittpunkte zwischen den bisherigen Linien und dem Spiegel mit den zugehörigen Gegenstandspunkten verbindest.

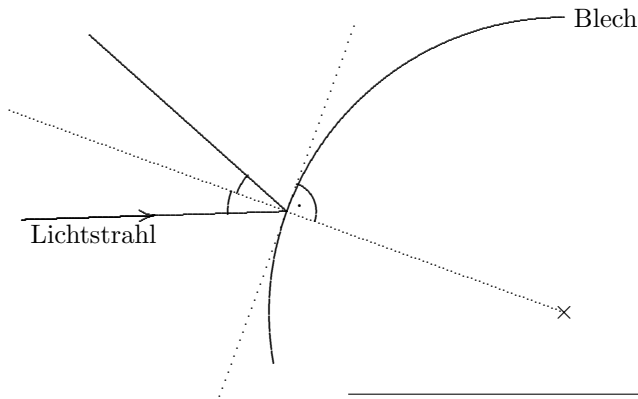
Die folgende Bildsequenz zeigt die Stationen einer solchen Konstruktion für einen Gegenstandspunkt und einen Strahl zum Beobachter:



9.5 Reflexion an gekrümmten Flächen

V 9.2 Reflexion an einem gebogenen Blech

Aufbau:



Durchführung: Ein Lichtstrahl wird auf einen kreisförmig gebogenen Blechstreifen gerichtet. (an der Tafel vorbeistreichend oder durch Rauch, damit man seinen Verlauf verfolgen kann)

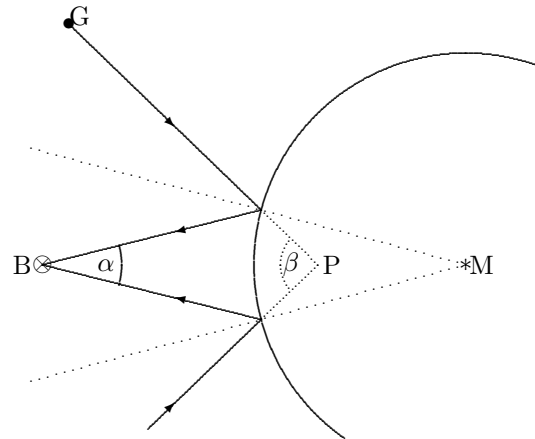
Beobachtung: Der Strahl verläuft wie gezeichnet.

Erklärung: Durch den Auftreffpunkt des Strahls verläuft das Blech wie die dünn gepunktete Tangente. Der Strahl wird demnach so reflektiert, als ob die dünn gepunktete Linie das spiegelnde Blech darstellte. Das Einfallslot (dicker gepunktet) ist also der verlängerte Radius durch den Auftreffpunkt. Aus dem Reflexionsgesetz ergibt sich damit der tatsächlich beobachtete Verlauf des Strahles.

Spiegelbilder in gekrümmten Flächen sind verzerrt, Größen und ihre Verhältnisse stimmen im Allgemeinen nicht mit denen am Originalgegenstand überein.

Warum z.B. sieht man sich in einer Kugel verkleinert?

Um dies zu verstehen müssen wir nicht zu einem Gegenstand das ganze Bild konstruieren. Es genügt, mit folgender Zeichnung zurückzuverfolgen, welche Lichtstrahlen ein Beobachter ($\otimes$) aus einem gewissen Bereich der Kugeloberfläche auffängt. (Obwohl das Licht natürlich *zum* Beobachter verläuft, muss man eine solche Zeichnung *vom* Beobachter aus konstruieren.) Auf einem Sichtkegel mit Öffnungswinkel α fängt der Beobachter im Punkt B also Lichtstrahlen von der Kugeloberfläche auf, die auf einem Kegel mit Öffnungswinkel β zur Kugel gekommen sind. Offensichtlich ist $\alpha < \beta$. Das aus dem Winkel β eintreffende Licht wird bei der Reflexion an der Kugel also auf den Winkel α verengt. Denselben Raumbereich, den ein Beobachter von P aus in einem Blickwinkel β sehen würde, sieht der Beobachter in B auf den Blickwinkel α verengt, also verkleinert.



9.6 Vergleich: Spiegel- und Lochkamerabilder

Wir haben bei der Lochkamera und bei Spiegeln Bilder kennengelernt. Zwischen den Bildern einer Lochkamera und Spiegelbildern gibt es jedoch einen wesentlichen Unterschied.

Das Bild bei einer Lochkamera besteht aus Punkten, auf die Licht jeweils nur eines Gegenstandspunktes *fällt*. Im Bildpunkt für z.B. einen roten Gegenstandspunkt *gibt es* daher überhaupt Lichtstrahlen und zwar nur rote. Eine Fotoplatte, die man an den Ort des Bildes hält, verfärbt sich an jedem Punkt entsprechend und kann so das Bild aufnehmen. Eine Mattscheibe am Ort des Bildes streut das Licht in alle Richtungen, so dass das Bild von allen Richtungen her zu sehen ist.

Das Bild bei einem Spiegel besteht aus Punkten, von denen Licht einer Farbe *herzukommen scheint*. Am Ort

des Bildes braucht es überhaupt keine Lichtstrahlen zu geben. Eine Fotoplatte oder eine Mattscheibe am Ort des Bildes fängt daher auch kein Bild auf.

Bilder, die aus tatsächlich lichtdurchfluteten Punkten bestehen und die man daher auf einer Fotoplatte aufnehmen kann, nennt man **reell**.

Bilder, die aus nur scheinbaren Ausgangspunkten von Licht bestehen und die man daher nicht auf einer Fotoplatte aufnehmen kann, nennt man **virtuell**. Also:

Ein Lochkamera-Bild ist

Ein Spiegelbild ist

Brechung

10.1 Strahlengang durch Grenzflächen

10.1.1 Experimentelle Befunde

Früher hast Du bereits gelernt, wie an einer Glasscheibe Spiegelbilder entstehen: Sie entstehen dadurch, dass Lichtstrahlen von der Glasscheibe reflektiert werden. Wir wissen aber auch, dass nicht das ganze Licht, das auf eine Glasscheibe trifft, reflektiert werden kann.

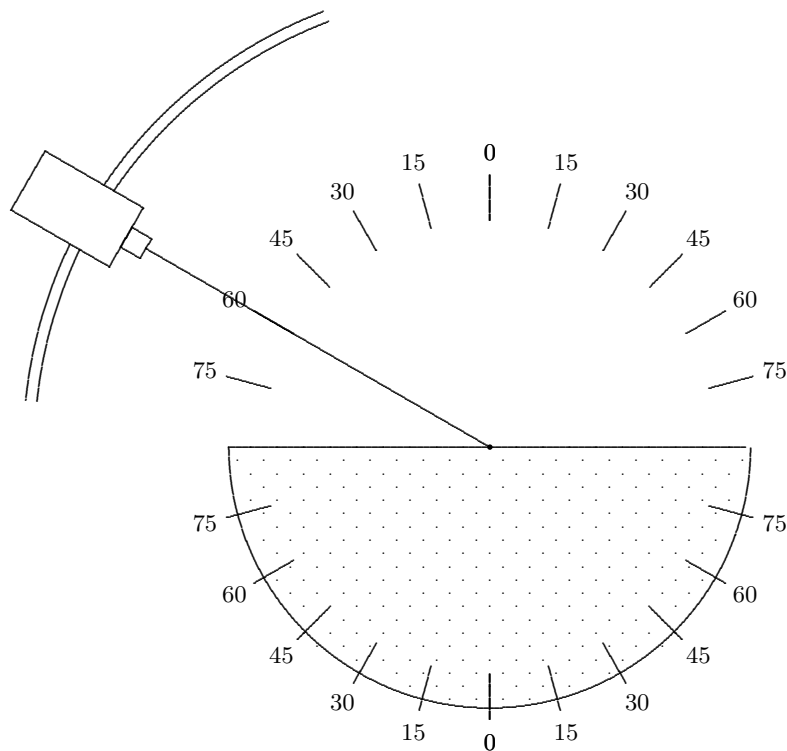
Du fragst Dich, wieso nicht? Nun: Wie könntest Du durch eine Glasscheibe hindurch Dinge sehen, wenn

kein Licht von diesen Dingen durch die Scheibe ginge?

Wir müssen also annehmen, dass sich ein Lichtstrahl beim Auftreffen auf eine Glasscheibe o.ä. aufspaltet: Ein Teil wird reflektiert, und ein Teil dringt in die Scheibe ein und läuft hindurch. In diesem Kapitel geht es um den eindringenden Teil.

V 10.1 Übergang eines Lichtstrahls von Luft in Glas

Aufbau: Ein Lichtstrahler ist schwenkbar so montiert, dass sein Lichtstrahl unter einstellbarem Winkel auf eine feste Stelle einer Glasoberfläche trifft. Der Lichtstrahl streift an einer Tafel mit Winkeleinteilung vorbei, so dass sein Verlauf sichtbar wird.



zu ergänzen:

- reflektierter Strahl
- gebrochener Strahl
- Winkel

Durchführung: Wir schwenken den Strahler, so dass der Lichtstrahl aus verschiedenen Richtungen auf die Grenzfläche zwischen Luft und Glas fällt.

Beobachtung: Außer dem einfallenden Lichtstrahl erkennt man einen reflektierten Strahl und einen Strahl, der in das Glas eindringt. Der reflektierte Strahl verläuft gemäß dem bereits bekannten Reflexionsgesetz. ($\alpha_R = \alpha_L$, Index L für Luft)

Wenn der einfallende Lichtstrahl senkrecht auf die Grenzfläche trifft ($\alpha_L = 0^\circ$), verläuft der eindringende Teil in gleicher Richtung weiter.

Bei schrägem Eintritt in das Glas ändert der Strahl seine Richtung. Wir beschreiben die Richtung des eingedrungenen Strahles mit α_G . (siehe die von Dir ergänzte und passend beschriftete Zeichnung, Index G für Glas)

Messwerte:

$\alpha_L [^\circ]$							
$\alpha_G [^\circ]$							

Auswertung:

Wir finden zunächst keine einfache rechnerische Beschreibung des genauen Zusammenhanges zwischen α_L und α_G . Wir stellen lediglich fest:

- Nur bei senkrechtem Einfall ist $\alpha_G = \alpha_L (= 0^\circ)$, d.h. der Strahl verläuft ohne Richtungsänderung
- In allen anderen Fällen ist $\alpha_G < \alpha_L$, d.h. der Strahl knickt beim Eintritt zum Einfallslot hin ab.
- Vergrößert man α_L , nimmt auch α_G zu, allerdings nicht so stark.

Das Abknicken eines Lichtstrahles an einer Grenzfläche zwischen zwei Materialien nennt man **Brechung**. Man sagt: der Lichtstrahl wird **gebrochen**. Wie der reflektierte Strahl liegt auch der gebrochene Strahl in der Einfallsebene.

V 10.2 Brechung beim Übergang in andere Materialien

Aufbau: wie zuvor, jedoch mit anderen Materialien (Wasser in einer kleinen Wanne, Eis, Diamant¹, andere Glassorte)

Durchführung: wie zuvor

Beobachtung: wie zuvor, jedoch mit anderen Brechungswinkeln bei gleichen Einfallswinkeln

Ergebnis: Der Brechungswinkel hängt nicht nur vom Einfallswinkel ab, sondern auch vom verwendeten Material.

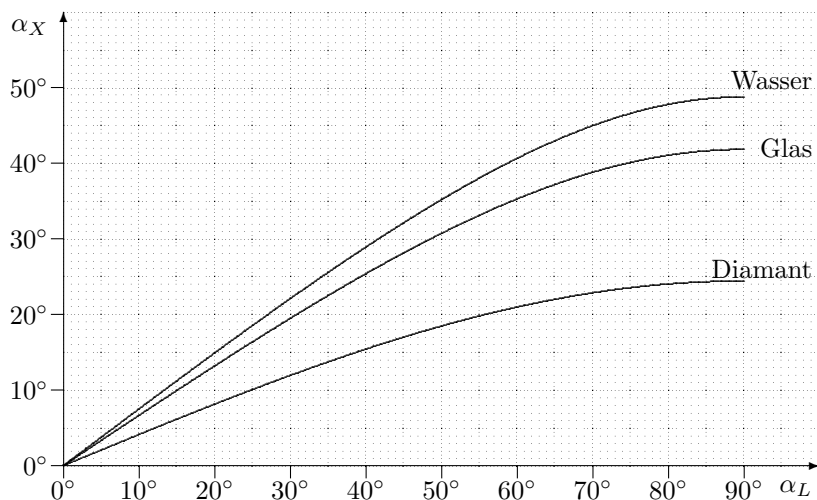
Ein Ziel wissenschaftlicher Forschung ist es, den Verlauf zukünftiger Ereignisse möglichst genau vorhersagen zu können. Für eine Licht-Brechung müsste man dazu bei bekanntem Einfallswinkel und bekanntem Material den Brechungswinkel voraussagen können.

Im Idealfall findet man eine Gleichung, mit der sich die vorauszusagende Größe aus den gegebenen Größen errechnen lässt. Für die Abhängigkeit des Brechungswinkels vom Einfallswinkel können wir eine solche Gleichung aber (noch) nicht angeben.

Mangels Gleichung können wir lediglich z.B. für Glas genügend viele ($\alpha_L | \alpha_G$)-Paare messen und tabellieren oder sonst irgendwie aufschreiben. Zweckmäßig trägt man die Paare in ein Koordinatensystem ein.

¹lag gerade in der Physik-Sammlung rum

Mathematisch: Die Menge aller Messwertpaare ist eine Funktion von der Menge aller untersuchten Einfallswinkel in die Menge der möglichen Ausfallswinkel. Die eingetragenen einzelnen Punkte bilden den Funktionsgraphen.



Trägt man genügend viele Punkte ein, fällt auf, dass die Punkte entlang einer recht einfach geformten Linie liegen. Man geht nun davon aus, dass jeder Punkt dieser Linie ein zu erwartendes Messwertpaar darstellt. Statt *Linie* sagt man in diesem Zusammenhang *Kurve*.

Mathematisch: Man erweitert die Definitionsmenge von der Menge der tatsächlich untersuchten Einfallswinkel auf die Menge aller möglichen Einfallswinkel. Funktionsgraph ist die aufgefallene Linie.

Nun kann man für alle möglichen Einfallswinkel den zu erwartenden Brechungswinkel ablesen, obwohl man nur für endlich viele Einfallswinkel den Brechungswinkel gemessen hat.

10.1.2 Brechungsgesetz

Wie kommt es zu Brechung?

?

Die Vorgänge an der brechenden Grenzfläche können wir mit unserem Strahlenmodell des Lichtes nicht vollständig verstehen. Das Phänomen Brechung liegt nämlich darin begründet, dass Licht auch eine wellenartige Natur hat, die dem Strahlenmodell völlig entgeht.

Um diese Erklärungslücke in unserem bisherigen Modell einigermaßen zu schließen, nimm zunächst folgenden Vergleich als Lückenbüßer: Ein Auto fährt von einer Straße schräg nach rechts in einen Acker hinein. Das rechte Vorderrad gelangt dabei zuerst in den unwegsamen Acker und kommt daher schon früh nicht mehr so schnell voran, während das linke Rad noch leicht weiterrollen kann. Das Auto zieht folglich während des Eintritts in den Acker nach rechts. (zum „Einfallslot“ hin.)

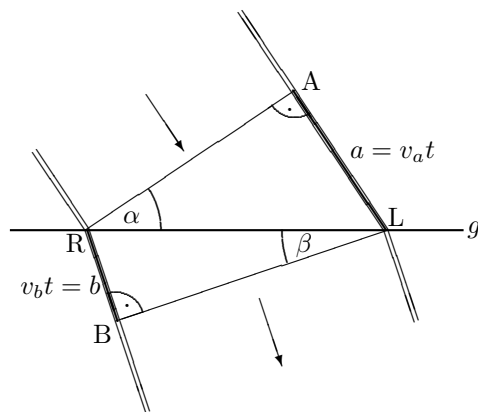
Zutreffend an diesem Vergleich ist, dass sich Licht in verschiedenen Materialien tatsächlich unterschiedlich schnell ausbreitet. Am schnellsten ist Licht im völlig leeren Raum (=Vakuum), in Luft nur ein wenig langsamer. Nicht zutreffend ist, dass ein Lichtstrahl ein rechtes und ein linkes Vorderrad hat.

Beim Vergleich zweier Materialien nennt man das Material, in dem Licht stärker gebremst wird, das **optisch dichtere Medium**, das andere das **optisch dünnere Medium**. Bei einer Brechung findet man also den größeren Winkel zwischen Strahl und Einfallslot im optisch dünneren Medium.

Die Lichtgeschwindigkeit im Vakuum beträgt:

Zur Beschreibung der optischen Dichte eines Medium dient der **Brechungsindex**, definiert als $v_{\text{Vakuum}}/v_{\text{Medium}}$.

Zur Herleitung des rechnerischen Zusammenhangs zwischen den Geschwindigkeiten des Lichtes in verschiedenen Medien und den Brechungswinkeln betrachte die nebenstehend dargestellte Situation: Eine Lichtwelle treffe von oben kommend auf die Grenzfläche g zwischen zwei Medien. Wie die Vorderachse obigen Autos bewege sich ein Wellenkamm zwischen den eingezeichneten „Reifenspuren“ in Pfeilrichtung auf die Grenzfläche zu. Dabei verläuft ein Wellenkamm stets senkrecht zu seiner Bewegungsrichtung. Beim ersten Kontakt mit der Grenzfläche verläuft der Kamm als $\overline{RA}$.



Während sich nun das obere Ende des Wellenkamms („das linke Vorderrad“) in einer gewissen Zeit t um die Strecke a von A nach L bewegt, kann sich in derselben Zeit das untere Ende („rechtes Vorderrad“) im optisch dichteren Medium nur um die kürzere Strecke b von R nach B bewegen. Danach verläuft der Kamm also als $\overline{BL}$ und bewegt sich – wie immer senkrecht zu sich selbst – wie es die „Reifenspuren“ und der Pfeil andeuten.

Der Rest ist nur noch ein wenig Mathematik: (Definition des Sinus in rechtwinkligen Dreiecken, Auflösen nach $|\overline{RL}|$ und Gleichsetzen, Ersetzen von a und b gemäß $v = s/t$, Kürzen)

$$\begin{aligned} \frac{a}{|\overline{RL}|} &= \sin(\alpha) \quad \wedge \quad \sin(\beta) = \frac{b}{|\overline{RL}|} \\ \frac{v_a \cdot t}{\sin(\alpha)} &= \frac{a}{\sin(\alpha)} = |\overline{RL}| = \frac{b}{\sin(\beta)} = \frac{v_b \cdot t}{\sin(\beta)} \\ \frac{v_a}{\sin(\alpha)} &= \frac{v_b}{\sin(\beta)} \end{aligned}$$

Du musst Dich nun nur noch davon überzeugen, dass α und β aus der vorigen Zeichnung mit dem Einfallswinkel bzw. Brechungswinkel zwischen Strahl und Einfallslot übereinstimmen.

Hat das Licht in zwei Medien die Geschwindigkeiten v_a bzw. v_b , so gilt für die Winkel α und β , die ein Lichtstrahl beim Übertritt von einem ins andere Medium mit dem Einfallslot bildet, die Beziehung:

$$\frac{\sin(\alpha)}{\sin(\beta)} = \frac{v_a}{v_b}$$

10.1.3 Umgekehrter Lichtweg

Wir kennen nun den Verlauf eines Lichtstrahles, der schräg in eine Glasscheibe eindringt: Er wird beim Eintritt ins Glas zum Einfallslot hin gebrochen.

?

Ändert der Strahl auch beim Austritt aus dem Glas seine Richtung?

Um diese Frage experimentell zu klären, können wir den vorigen Versuchsaufbau fast unverändert wieder benutzen: Es muss wieder ein Lichtstrahl auf die Grenzfläche zwischen Luft und Glas fallen, dieses Mal jedoch nicht von der Seite der Luft her, sondern vom Glase her.

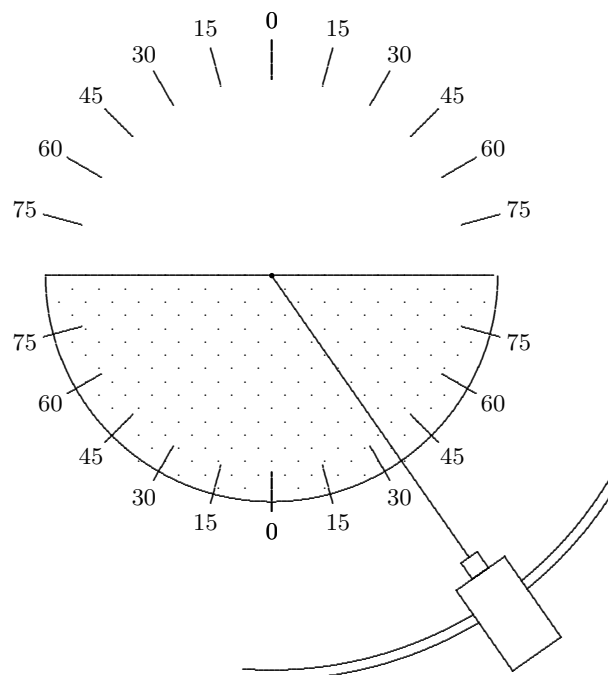
Der Strahl muss natürlich erst einmal in das Glas hinein, bevor er austreten kann. Damit nicht schon beim Eintritt eine Brechung stattfindet, die uns verwirren könnte, lassen wir den Strahl stets senkrecht in das Glas eintreten. Hier zeigt sich der Vorteil des halbzylindrischen Glaskörpers: Auf die halbkreisförmig gebogene Seite des Glaskörpers trifft der Lichtstrahl bei jeder Stellung des Strahlers senkrecht auf.

V 10.3 Übergang eines Lichtstrahls von Glas in Luft

Aufbau: wie zuvor, jedoch: Strahler von unten auf die gebogene Seite des Glaskörpers gerichtet

zu ergänzen:

- reflektierter Strahl
- gebrochener Strahl



Durchführung: Wir schwenken den Strahler so, dass der Lichtstrahl nacheinander unter den Winkeln α_G aus dem ersten Brechungsversuch auf die Grenzfläche trifft.

Beobachtung: Der einfallende Strahl wird zum Teil an der Grenzfläche reflektiert, und zum Teil passiert er die Grenzfläche. Bei schrägem Einfall findet Brechung vom Einfallslot weg statt.

Messwerte:

$\alpha_G [^\circ]$							
$\alpha_L [^\circ]$							

Ergebnis:

Jeder Weg eines gebrochenen Lichtstrahles kann von Licht auch in entgegengesetzter Richtung durchlaufen werden. Das Diagramm auf Seite 87 kann also auch zur Bestimmung von α_L von einem gegebenen α_G, w, D ausgehend umgekehrt gelesen werden.

Wir erinnern uns: Auch bei der Reflexion kann jeder Lichtweg ebenso in umgekehrter Richtung durchlaufen werden. Es gilt ganz allgemein:

Lichtwege sind umkehrbar.

10.2 Totalreflexion

10.2.1 Theorie

Für jeden Lichtstrahl, der von Luft auf Glas trifft, können wir seinen weiteren Verlauf im Glas konstruieren. Wir wissen weiterhin, dass jeder so konstruierte Lichtweg auch in umgekehrter Richtung – vom Glas zur Luft – durchlaufen werden kann. Mit jedem dieser Lichtwege kennen wir demnach den weiteren Verlauf eines Strahls, der von der Glas-Seite her den Weg in umgekehrter Richtung zur Grenzfläche hin zurücklegt.

Aus den Kurven im Diagramm auf Seite 87 entnehmen wir jedoch, dass kein Lichtstrahl nach dem Übergang Luft→Glas im Glas unter einem Winkel größer 42° verlaufen kann. Nichts und niemand verbietet uns aber, von der Glas-Seite her einen Lichtstrahl unter einem größeren Einfallswinkel auf die Grenzfläche zu richten.

Einen weiteren Verlauf dieses Strahles in der Luft können wir nicht durch Umkehrung eines uns bereits bekannten Lichtweges angeben. Kein Lichtweg führt aus der Luft unter so großem Winkel ins Glas hinein.

Da alle Lichtwege umkehrbar sind, bedeutet dies jedoch gleichzeitig: Kein Lichtweg mit so großem Winkel im Glas führt aus dem Glas heraus in die Luft hinein.

Nach dieser Überlegung dürfte ein Strahl mit so großem Einfallswinkel im Glas überhaupt nicht aus dem Glas heraustreten. Anders gesagt: es dürfte keinen Übertritt zwischen den Medien Glas und Luft geben und damit natürlich auch keine Brechung.

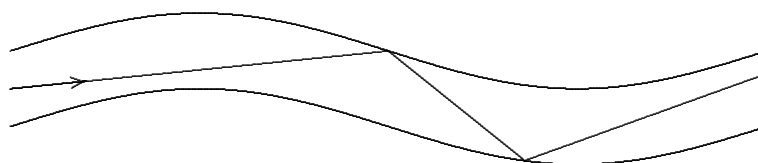
Schwenkt man im Versuch auf Seite 89 den Strahler auf größere Winkel als 42° , so verschwindet tatsächlich der gebrochene Strahl. Es bleibt nur der reflektierte Strahl zu sehen. Die eben angestellte Überlegung ist damit bestätigt.

Da von dem einfallenden Strahl kein Teil gebrochen, sondern alles reflektiert wird, nennt man diese Erscheinung **Totalreflexion**. (lat.: totus=ganz) Den kleinsten Einfallswinkel, bei dem Totalreflexion auftritt, nennt man **Grenzwinkel der Totalreflexion**.

10.2.2 Anwendungen

Totalreflektierende Flächen sind optimale Spiegel. Schließlich reflektieren sie nicht nur einen Teil des einfallenden Lichtes, sondern 100%ig das ganze Licht. Es gibt deshalb einige Anwendungen totalreflektierender Flächen.

Glasfaserkabel sind die heute vielleicht wichtigste Anwendung der Totalreflexion. In ein Ende einer Glasfaser wird ein Licht eingestrahlt. Soweit die Faser nicht geknickt ist, trifft der Lichtstrahl immer nur unter Winkeln auf die Faserwand, die größer als der Grenzwinkel der Totalreflexion sind. Der Lichtstrahl wird daher an der Faserwand immer nach innen totalreflektiert und bleibt somit in der Faser. Das Licht lässt sich folglich mit der Glasfaser auf beliebigen Wegen von einem Ort zu einem anderen leiten.



Damit lassen sich nicht nur dekorative Lampen bauen: Man kann Informationen aller Art in Form von Lichtsignalen durch Glasfasern übertragen — viel besser, als man es mit elektrischen Signalen durch Kupferleitungen kann. („besser“ heißt: mehr Information in kürzerer Zeit über längere Entfernungen bei weniger Störungen)

Totalreflexion ist auch im Spiel, wenn sich an heißen Tagen auf einer Teerstraße weit entfernte Din-

ge spiegeln, so als ob die Straße ein großer Spiegel oder ein ruhiger See wäre. Die Straße ist jedoch trocken und rau, so dass sie sicher nicht als Spiegel wirken kann.

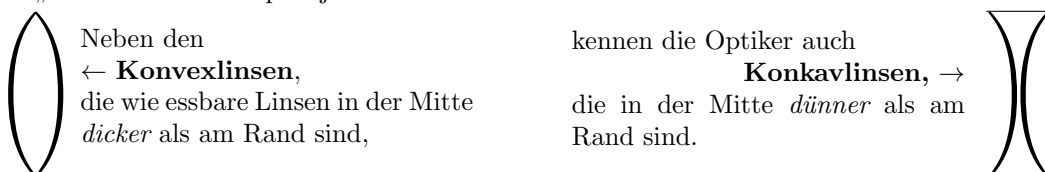
Bei starker Sonneneinstrahlung und Windstille bildet sich jedoch über der beschienenen Straße eine dünne heiße Luftschicht. Diese heiße Luft hat nun aber eine geringere optische Dichte als die darüberliegende kühlere Luft. Lichtstrahlen, die flach genug auf die Grenzfläche treffen, werden daher an ihr totalreflektiert.

Abbildung durch Linsen

11.1 Strahlengang durch Linsen

11.1.1 Linsenformen

Die Linsen, um die es hier und z.B. bei Fotoapparaten geht, sind Glaskörper. Eine mögliche Form dieser Glaslinsen ist die Form der essbaren Linsen. Die Bezeichnung „Linse“ ist in der Optik jedoch nicht auf diese Form beschränkt:



Dabei müssen Linsen noch nicht einmal symmetrisch sein. Sie können durchaus auf beiden Seiten verschieden gekrümmt sein: Zum Beispiel links konvex und rechts konkav, links plan (d.h. gar nicht gekrümmt) und rechts konvex, links schwach gekrümmt und rechts stark, ...

Für die optischen Eigenschaften im Grundsätzlichen ist es auch nicht nötig, dass Linsen aus Glas bestehen. Beispielsweise zeigt auch Wasser mit gekrümmter Oberfläche dieselben Eigenschaften.

11.1.2 Strahlengänge bei Konvexlinsen

Parallel einfallende Strahlen

V 11.1 Achs-parallele Lichtstrahlen auf Konvexlinse

Aufbau: An einer Hafttafel werden eine Lampe und eine Konvexlinse angebracht wie gezeichnet. Die Lampe erzeugt einige zueinander parallele Lichtstrahlen. (genauer: sehr enge Lichtbündel)



zu ergänzen:

- weitere Strahlverläufe

Beobachtung:

?

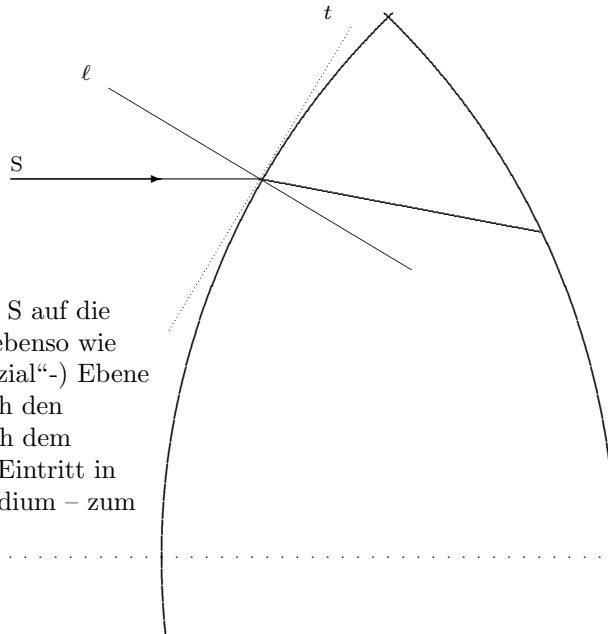
Warum ändern die Strahlen ihre Richtung zur optischen Achse hin?

Wenn Licht in die Linse *eintritt*, und wenn es wieder *austritt*, tritt es jedenfalls von einem Medium in ein anderes über, nämlich von Luft in Glas bzw. von Glas in Luft. Dabei tritt im Allgemeinen **Brechung** auf. Bisher haben wir Brechung an *ebenen* Grenzflächen kennengelernt. Offensichtlich findet sie aber auch an den *gekrümmten* Flächen einer Linse statt. Dazu die Groß-Aufnahme im folgenden Bild.

zu ergänzen:

- weiterer Strahlenverlauf

In dem Punkte, in dem der Lichtstrahl S auf die Linse trifft, verläuft deren Oberfläche ebenso wie die gepunktet eingezeichnete („Tangenzial“-) Ebene t . Die Senkrechte zu dieser Ebene durch den Auftreffpunkt ist das Einfallslot ℓ . Nach dem Brechungsgesetz wird der Strahl beim Eintritt in die Linse – als das optisch dichtere Medium – zum Einfallslot hin gebrochen.



?

Warum schneiden sie sich alle im selben Punkt?

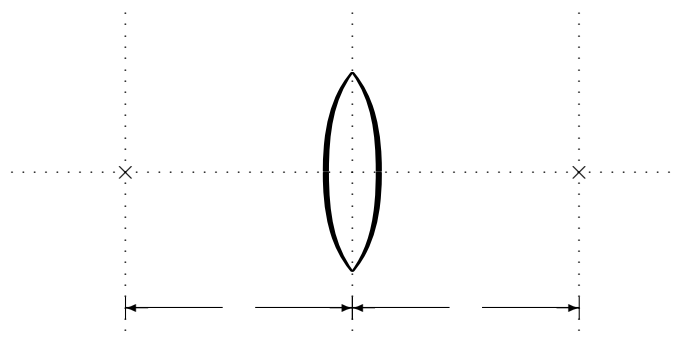
Eigentlich schneiden sich die gebrochenen Strahlen gar nicht alle *exakt* im selben Punkt. Um die Schnitte rechnerisch zu untersuchen, benötigt man neben dem Brechungsgesetz als Formel auch allerlei Rechenkünste über Dreiecke und Kreise. Auch wenn wir sie hier nicht durchführen können, bestätigen diese Rechnungen unsere Beobachtungen:

Lichtstrahlen, die *parallel* zur optischen Achse und *nahe* der optischen Achse einfallen, werden so gebrochen, dass sie die optische Achse hinter der Linse praktisch im selben Punkt schneiden.

(Die achsfernen Strahlen, für die dies nicht so gut erfüllt ist, blendet man in optischen Versuchen meist aus, indem man Lochblenden vor die betroffenen Linsen stellt.)

Den Schnittpunkt nennt man den **Brennpunkt** oder **Fokus** der Linse. Seine Entfernung von der Linsenmitte heißt **Brennweite**, gewöhnlich abgekürzt mit f . Die Ebene senkrecht zur optischen Achse durch einen Brennpunkt heißt **Brennebene**.

Eine Linse kann natürlich ebenso von rechts wie von links bestrahlt werden. Die Situation ist dann symmetrisch zu der bisher besprochenen. Es gibt demnach auf jeder Seite einen Brennpunkt, $F_{1,2}$. Die Brennweiten sind beiderseits gleich. (Sogar dann, wenn die Linse links und rechts verschieden stark gekrümmt ist.) Folgendes – noch zu beschriftendes – Bild zeigt die eingeführten Begriffe:



Achs-parallel einfallende Lichtstrahlen verlaufen hinter der Linse durch den dortigen Brennpunkt.
Kurz: Parallelstrahlen werden zu Brennstrahlen.

Was hat der Brennpunkt mit „brennen“ zu tun?

?

Wenn alle Strahlen, die vor der Linse parallel nebeneinander verlaufen, hinter der Linse durch denselben Punkt verlaufen, dann verlaufen dort die Lichtstrahlen viel dichter, und das Licht ist im Bereich dieses Punktes viel intensiver als außerhalb. Dasselbe gilt für Wärmestrahlung, die wir ja bereits als eine unsichtbare Art von Licht kennengelernt haben. Bündelt man also z.B. genügend viel Sonnenlicht, kann die gebündelte Sonnenstrahlung im Brennpunkt ein Papierstückchen entzünden!

Welchen Verlauf nehmen Lichtstrahlen, die zwar untereinander, aber nicht zur optischen Achse parallel sind?

?

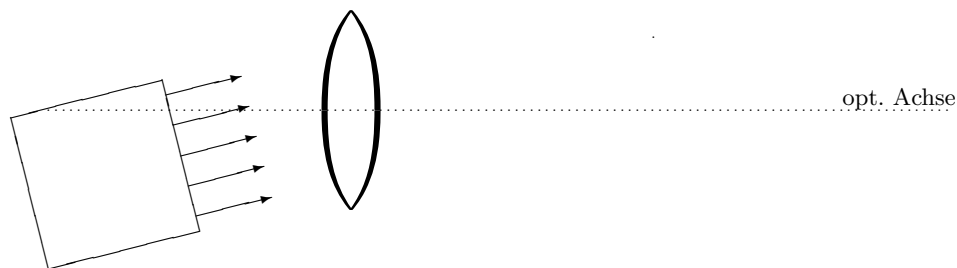
V 11.2 Zueinander parallele Lichtstrahlen auf Konvexlinse

Aufbau: Wie zuvor.

Durchführung: Die Lampe wird so verschoben, dass ihr Licht aus anderen Richtungen auf die Linse fällt.

zu ergänzen:

- weitere Strahlverläufe
- Brennebene



Beobachtung:

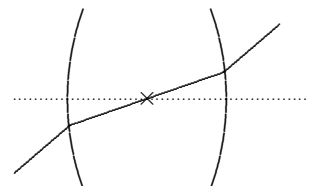
Insgesamt ergeben die beiden letzten Versuche:

Zueinander parallel einfallende Lichtstrahlen schneiden sich alle in einem Punkt der Brennebene.

Mittelpunktstrahlen

Die beiden vorigen Versuche zeigen außerdem, dass Lichtstrahlen, die durch den Mittelpunkt der Linse verlaufen, keine Richtungsänderung erfahren. Dies ist bereits aus Symmetrie-Überlegungen so zu erwarten. Ein Mittelpunktstrahl kann höchstens innerhalb der Linse ein wenig seitlich versetzt werden. Davon abgesehen gilt also:

Strahlen durch den Mittelpunkt einer Linse durchqueren die Linse ohne Ablenkung.



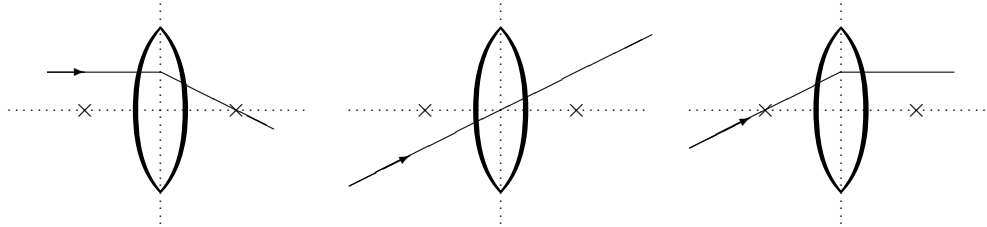
Brennstrahlen

Wir betrachten nun Lichtstrahlen, die von einem Brennpunkt her auf eine Konvexlinse treffen. Diese Strahlen durchlaufen also rückwärts denselben Weg, den ein Parallelstrahl von der anderen Seite der Linse her nach seiner Brechung nehmen würde. Da wir wissen, dass die Lichtwege bei der Brechung (wie auch sonst immer) umkehrbar sind, können wir voraussagen:

Von einem Brennpunkt her einfallende Strahlen verlaufen hinter der Linse parallel zur optischen Achse.
Kurz: Brennstrahlen werden zu Parallelstrahlen.

Vereinfacht gezeichnete Strahlengänge

So wie wir bei den Mittelpunktstrahlen den seitlichen Versatz eines Strahles während des Durchquerens der Linse außer Acht gelassen haben, werden wir zukünftig immer den Lichtweg innerhalb der Linse außer Acht lassen. Solange wir mit relativ dünnen Linsen arbeiten, kann das kein großer Fehler sein. Anstatt bei einem Lichtstrahl *beide* stattfindenden Brechungen zu zeichnen, werden wir unsere Konstruktionen so zeichnen, als ob nur *eine* Brechung an der Mittelebene der Linse stattfände. Die drei besonderen Strahlarten zeichnen wir also wie folgt:



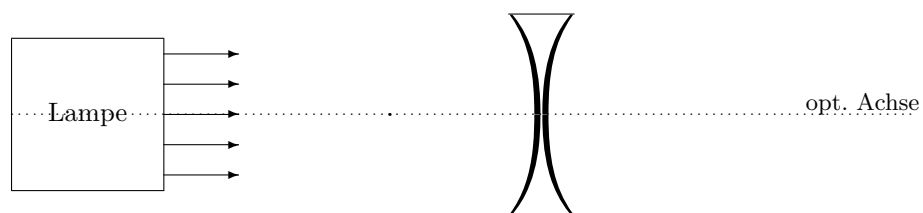
11.1.3 Strahlengänge bei Konkavlinsen

V 11.3 Parallele Lichtstrahlen auf Konkavlinse

Aufbau: An einer Hafttafel werden eine Lampe und eine Konkavlinse angebracht wie gezeichnet. Die Lampe erzeugt einige zueinander parallele Lichtstrahlen. (zumindest sehr enge Lichtbündel)

zu ergänzen:

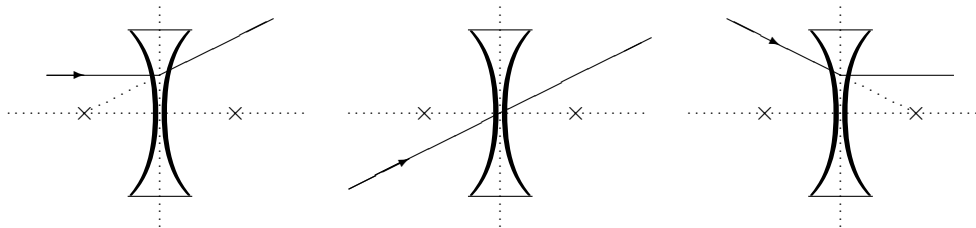
- weitere Strahlverläufe



Beobachtung:

Bei Konkavlinsen schneiden sich die achsparallel eingefallenen gebrochenen Strahlen zwar gar nicht, wohl aber ihre rückwärtigen Verlängerungen. Sie schneiden sich sogar alle im selben Punkt, der auch hier *Brennpunkt* genannt wird. (Obwohl hier sicher nichts besonders gut in Brand gesetzt werden kann.) Zur Unterscheidung gibt man Brennweiten von Konkavlinsen jedoch negativ an.

Auch sonst gibt es weitere Gemeinsamkeiten mit Konvexlinsen: Versteht man hier unter einem Brennstrahl einen Strahl, dessen Verlängerung durch den gegenüberliegenden Brennpunkt verläuft, so gelten sogar alle Merkgeln über die besonderen Strahlen unverändert auch hier:



Der Unterschied zwischen Konvex- und Konkavlinsen ist anschaulich formuliert folgender: Treffen parallele Strahlen auf die Linse, so laufen diese hinter einer Konvexlinse zunächst *zusammen*, hinter einer Konkavlinse jedoch sofort *auseinander*. Man teilt Linsen nach diesem Kriterium ein in

Sammellinsen und Zerstreulinsen

11.2 Bildentstehung bei Konvexlinsen

Zu Beginn dieses Kapitels hast Du gelesen, dass man Fotoapparate mit Linsen baut. Wir haben nun genügend Grundlagen-Wissen gesammelt, um zu verstehen, wie bei einer Linse das Bild eines Gegenstandes entsteht.

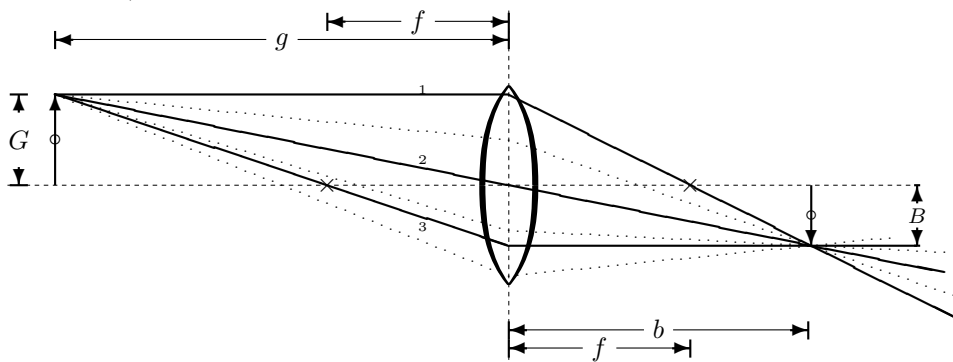
Wie verlaufen die Lichtstrahlen, die von einem Gegenstand auf eine Konvexlinse fallen, hinter der Linse weiter?

?

11.2.1 Konstruktion von Bildern

Wir betrachten dazu zunächst einen einzigen Gegenstandspunkt. Wie üblich zeichnen wir dazu als Gegenstand einen aufrecht auf der optischen Achse stehenden Pfeil ($\uparrow$) und konzentrieren uns auf das Licht von dessen Spitze.

Von dieser Spitze aus konstruieren wir die drei besonderen Strahlen, deren weiteren Verlauf wir kennen: Parallel-₍₁₎, Mittelpunkt-₍₂₎ und Brennstrahl₍₃₎. (durchgezogen gezeichnet)



Es fällt auf, dass diese drei Strahlen hinter der Linse einen gemeinsamen Punkt durchlaufen. Es lässt sich experimentell zeigen, dass sogar *alle* (achsnahe) Strahlen, die von der Pfeilspitze ausgehen, durch diesen Punkt verlaufen. (Einige davon sind gepunktet dargestellt.)

Für jeden anderen Punkt des Gegenstandes gibt es ebenfalls einen solchen *Bildpunkt*, den alle vom Gegenstandspunkt durch die Linse verlaufenden Lichtstrahlen treffen. Sofern der Gegenstandspunkt nicht auf der optischen Achse liegt, lässt sich sein Bildpunkt dabei genauso konstruieren, wie oben gezeigt. In obiger Zeichnung markieren zwei Kringelchen ($\circ$) ein weiteres Beispiel-Paar aus Gegenstands- und Bildpunkt.

Jeder Bildpunkt wird nur vom Licht *eines einzigen* Gegenstandspunktes durchflutet. Eine Fotoplatte im Bereich der Bildpunkte wird sich daher an jeder Stelle gemäß der Farbe *eines* Gegenstandspunktes verfärben. Insgesamt entsteht so ein (kopfstehendes) Bild ($\downarrow$) des ganzen Gegenstandes. Im Allgemeinen kann das Bild dabei verkleinert, vergrößert oder originalgroß sein.

Abbildungsmaßstab

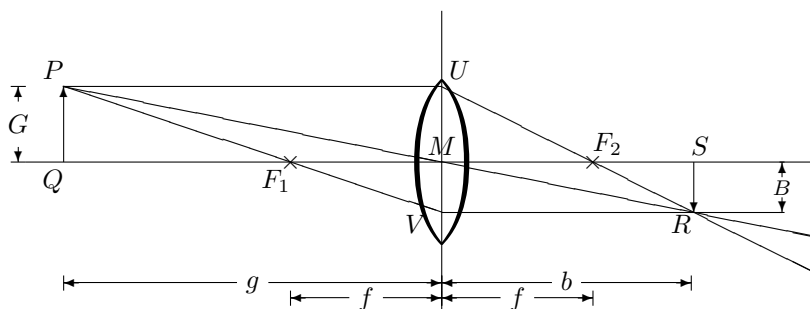
$$A := \frac{B}{G}$$

Zur Beschreibung der geometrischen Verhältnisse bei einer solchen Abbildung benutzt man die gleichen Größen, die Du bei der Lochkamera bereits kennengelernt hast: Gegenstandsgröße G und -weite g , Bildgröße B und -weite b . Nützlich ist auch der Begriff **Abbildungsmaßstab**. Dies ist einfach der Faktor, um den das Bild gegenüber dem Gegenstand vergrößert ist.

Zur Konstruktion eines ganzen Bildes geht man folgendermaßen vor: Man konstruiert die Bildpunkte markanter Gegenstandspunkte und ergänzt den Rest. Zur Konstruktion eines Bildpunktes genügt es dabei, zwei der drei besonderen Strahlen zu zeichnen, weil mit deren Schnittpunkt der Bildpunkt bereits bekannt ist. Es ist auch nicht erforderlich, dass diese Strahlen durch die gezeichnete Linse verlaufen: Man orientiert sich immer nur an der Mittelebene. Die Linse selbst kann in der Zeichnung sogar ganz weggelassen werden.

11.2.2 Rechnerische Zusammenhänge

Die nebenstehende Zeichnung zeigt (mal wieder) einen Gegenstand $\overline{QP}$ auf der optischen Achse QS einer Linse im Punkt M mit Brennpunkten $F_{1,2}$, das Bild $\overline{SR}$ und die drei besonderen Strahlen von P nach R durch die Stellen U , M bzw. V der Mittelebene.



Du erkennst darin mehrere Paare ähnlicher Dreiecke, aus denen Du die Gleichheit entsprechender Seitenverhältnisse herauslesen kannst:

$$\left. \begin{array}{l} \triangle MQP \sim \triangle MSR \Rightarrow \frac{G}{B} = \frac{g}{b} \\ \triangle QPF_1 \sim \triangle MVF_1 \Rightarrow \frac{G}{B} = \frac{g-f}{f} \\ \triangle SRF_2 \sim \triangle MUF_2 \Rightarrow \frac{G}{B} = \frac{f}{b-f} \end{array} \right\} \Rightarrow \frac{g}{b} = \frac{g-f}{f}$$

Die gewonnene Gleichung lässt sich noch ein wenig schöner umformen:

$$\frac{g}{b} = \frac{g-f}{f} \Leftrightarrow \frac{g}{b} = \frac{g}{f} - \frac{f}{f} \Leftrightarrow \frac{g}{b} = \frac{g}{f} - 1 \Leftrightarrow \frac{1}{b} = \frac{1}{f} - \frac{1}{g}$$

Mit den üblichen Bezeichnungen gelten bei einer (dünnen) Linse

die Linsengleichung

$$\frac{1}{f} = \frac{1}{b} + \frac{1}{g}$$

und

die Abbildungsgleichung

$$\frac{b}{g} = A = \frac{B}{G}$$

Mit der Gleichung aus dem dritten und bisher nicht ausgenutzten Paar ähnlicher Dreiecke lässt sich die Linsengleichung ebenso gut herleiten. Dies bestätigt uns, dass auch der Strahl $P \rightarrow U \rightarrow F_2$ durch R verläuft, dass sich also alle drei besonderen Strahlen im selben Punkt schneiden.

11.2.3 Variation der Gegenstandsweite

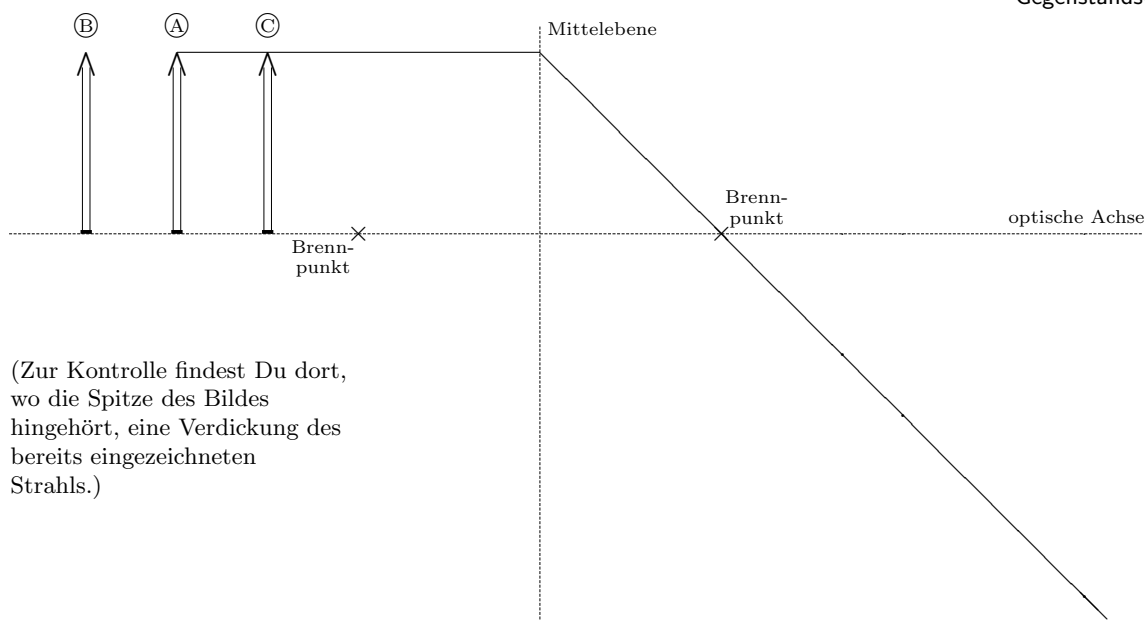
Abbildung in Originalgröße

Wir suchen zunächst eine Gegenstandsweite, bei der das Bild ebenso groß wie der Gegenstand ist. Von dieser Weite jeweils ausgehend werden wir dann untersuchen, wie sich das Bild beim Vergrößern und beim Verkleinern der Weite ändert.

Laut Abbildungsgleichung ist bei $B = G$ auch $b = g$. Laut Linsengleichung ist dann $\frac{1}{f} = \frac{2}{g}$ und damit $g = 2f$. Durch diese kleine Rechnung oder durch experimentelles Probieren erhält man also:

Befindet sich der Gegenstand in doppelter Brennweite, dann ist das Bild originalgroß und ebenfalls in doppelter Brennweite.

In folgender Zeichnung ist Pfeil ① in doppelter Brennweite aufgestellt. Färbe ihn in einer Farbe Deiner Wahl! Der Parallelstrahl von seiner Spitze aus ist bereits eingezeichnet. Ergänze die Zeichnung in derselben Farbe um den Mittelpunktstrahl, den Brennstrahl und das Bild des Pfeils!



(Zur Kontrolle findest Du dort, wo die Spitze des Bildes hingehört, eine Verdickung des bereits eingezeichneten Strahls.)

Unter der Internet- Adresse
softfrutti.de
(→ „Materialien“)
wurde die Datei
LINSE.GEO
zum Download bereit-
gestellt und kann mit dem
Programm *Euklid*
(→ www.dynageo.de)
benutzt werden.
Sie zeigt den hier
behandelten Strahlengang
bei veränderbarer
Gegenstandsweite.

Vergrößerung der Gegenstandsweite

Pfeil ③ weist gegenüber ① eine vergrößerte Gegenstandsweite auf. Ergänze die Zeichnung mit einer anderen Farbe um die drei besonderen Strahlen von B's Spitze aus und um B's Bild! Streiche dann in folgendem Satz das Nicht-Zutreffende:

Entfernt man außerhalb der Brennweite den Gegenstand von der Linse, dann bewegt sich das Bild [zur Linse hin|von der Linse weg] und wird [kleiner|größer].

Verringerung der Gegenstandsweite

Verringert man von ① ausgehend die Gegenstandsweite, so ergibt sich beispielsweise Pfeil ②. Zeichne mit einer weiteren Farbe auch für diesen Pfeil die Strahlen und das Bild! Man erkennt:

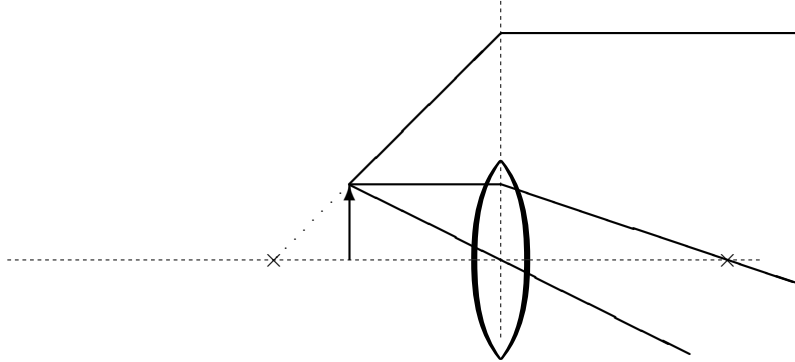
Rückt man außerhalb der Brennweite den Gegenstand näher an die Linse, dann bewegt sich das Bild [zur Linse hin|von der Linse weg] und wird [kleiner|größer].

Verringert man die Gegenstandsweite unter die Brennweite, so ändert sich der Strahlengang grundlegend: Die gebrochenen Strahlen schneiden sich überhaupt nicht! Auf der rechten Seite der Linse kann folglich auch kein Bild aufgefangen werden.

Es schneiden sich allerdings die rückwärtigen Verlängerungen aller gebrochenen Strahlen. Für einen Beobachter rechts von der Linse sieht es so aus, als ob alle Strahlen von diesem Schnittpunkt her kämen. Für den Beobachter scheint das Licht von dem großen Pfeil ganz links zu kommen, den Du noch einzeichnen musst. Der Beobachter sieht diesen scheinbaren großen Pfeil.

zu ergänzen:

- rückwärtige Verlängerungen
- virtuelles Bild

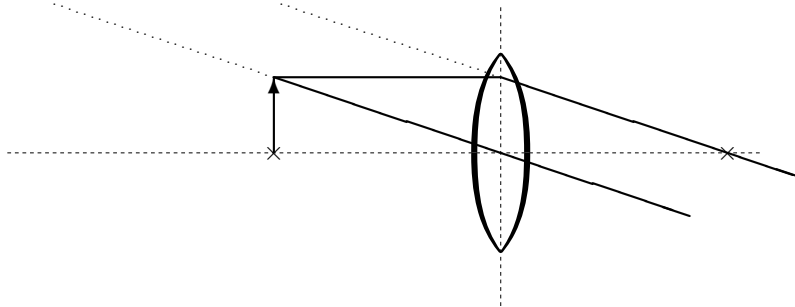


Zur Unterscheidung zwischen dieser Art Bilder und den bisher vorgekommenen benutzt man folgende Adjektive:

Bilder, die aus *tatsächlich lichtdurchfluteten Punkten* bestehen, heißen **reell**.
 Bilder, die aus *scheinbaren Ausgangspunkten von Licht* bestehen, heißen **virtuell**.

Im Grenzfall $g = f$ ist der Strahlengang wieder anders. Einen Brennstrahl vom Gegenstand zur Linse gibt es gar nicht. Die übrigen Strahlen werden zu parallelen Strahlen gebrochen, so dass sich weder die gebrochenen Strahlen selbst noch ihre rückwärtigen Verlängerungen schneiden. Es gibt deshalb hier weder ein reelles noch ein virtuelles Bild.

Dieser Strahlengang folgt übrigens schon aus den beiden Regeln, dass parallel zueinander einfallende Strahlen in der Brennebene gebündelt werden, und dass Lichtwege umkehrbar sind.

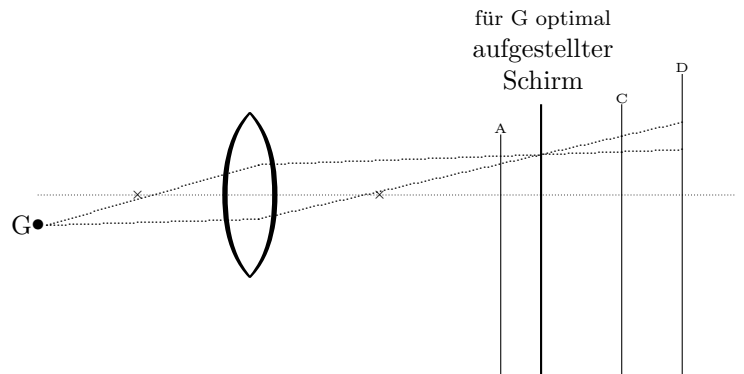


11.3 Geräte zur Bildaufnahme

11.3.1 Bildschärfe

Ein Bild ist ganz scharf, wenn jeden Punkt des Bildes nur das Licht eines einzigen Gegenstandspunktes trifft. Unschärfe entsteht, wenn zusätzlich Licht anderer, meist benachbarter Gegenstandspunkte auftrifft. Dies geschieht, wenn das Licht eines Gegenstandspunktes nicht nur einen einzigen Punkt des Bildes beleuchtet, sondern einen ausgedehnten Bereich. Bereiche, die zu benachbarten Gegenstandspunkten gehören, können sich dann überlappen, so dass das Bild verschwimmt. Dasselbe Problem hast Du bereits bei der Lochkamera kennengelernt.

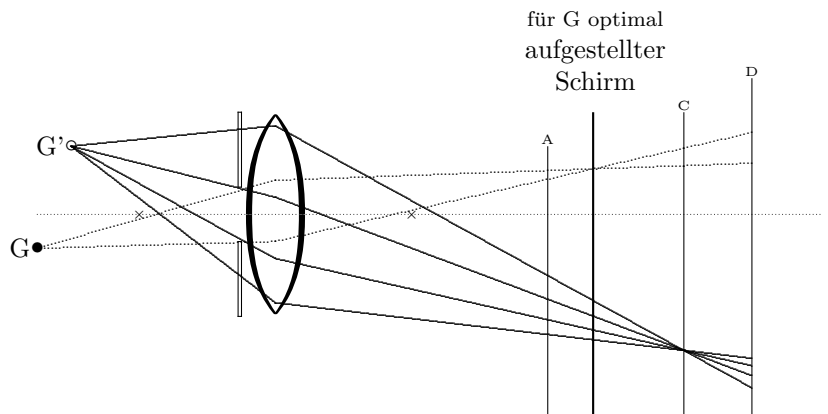
Hinter einer Sammellinse erhält man von einem Gegenstandspunkt G nur dann einen scharfen Bildpunkt auf einem Schirm, wenn der Schirm genau dort steht, wo sich die entsprechenden Lichtstrahlen schneiden, also genau in Bildweite hinter der Linse. Steht der Schirm etwas näher an oder weiter von der Linse, so beleuchtet das Licht des Gegenstandspunktes einen ausgedehnten Fleck.



Die Größe des Fleckes hängt davon ab, wie weit der Schirm von der optimalen Position verrückt ist. (siehe Schirme C und D)

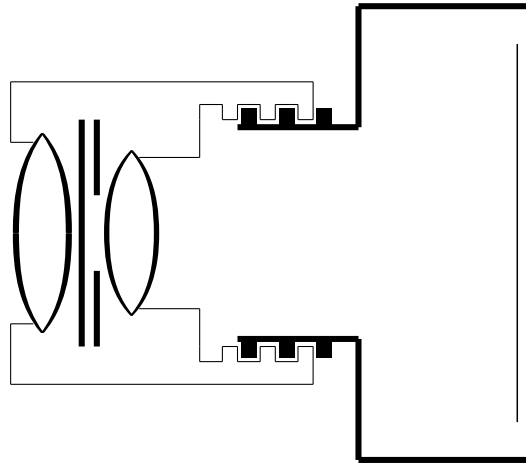
Schirm C stünde optimal für einen (näher als G gelegenen) Gegenstandspunkt G' wie in folgender Abbildung. Der für G optimale Schirm fängt von diesem G' jedoch einen ausgedehnten Fleck auf.

Die Größe dieses Fleckes hängt mit davon ab, wie weit der beleuchtende Lichtkegel geöffnet ist, was seinerseits von der Blendenöffnung abhängt. Die Abbildung rechts zeigt die Strahlengänge von G' aus, mit und ohne enge Blende. Auf dem für G optimal aufgestellten Schirm beleuchtet G' entsprechend verschieden große Flecken.



Obwohl eine bestimmte eingestellte Bildweite streng genommen nur für eine einzige Gegenstandsweite genau passt, bleibt also immerhin bei kleinerer Blendenöffnung für einen größeren Wertebereich von Gegenstandsweiten eine einigermaßen scharfe Abbildung erhalten. Da die Gegenstandsweite in den erzeugten Bildern üblicherweise als „Tiefe“ bezeichnet wird, nennt man die Größe dieses Bereiches mit scharfer Abbildung die **Tiefenschärfe**.

11.3.2 Fotoapparat



Bei geöffnetem **Verschluss** ist der **Film** dem Licht ausgesetzt, das durch das **Objektiv** aus der Außenwelt in die Kamera eindringt. Das Filmmaterial wird durch diese Belichtung chemisch so verändert, dass daraus später im Labor Bilder entwickelt werden können.

Die belichtende Lichtmenge muss natürlich einerseits groß genug sein, um ausreichende chemische Umwandlungen hervorzurufen, darf aber andererseits nicht so groß sein, dass sich alle belichteten Stellen ununterscheidbar vollständig umwandeln. (Für Schwarz-

Weiß formuliert: genug Licht, damit nicht alles schwarz bleibt, und nicht zu viel Licht, damit nicht alles weiß wird. Auf die Grau-Töne dazwischen kommt es an.)

Das Objektiv (=das Teil, das dem zu fotografierenden *Objekt* zugewandt ist) hat die Aufgabe, das einfallende Licht so zu bündeln, dass in der Filmebene ein Bild der Außenwelt entsteht. Es enthält dazu mehrere Linsen, die wir in ihrer Gesamtheit jedoch vereinfacht wie eine einzige Sammellinse betrachten können.

Es ist erforderlich, mehrere Linsen zu kombinieren, weil man mit einem Fotoapparat ja in einem *handlichen* Gehäuse *helle* und *scharfe* Bilder machen will. Man muss dazu eine relativ kurze Brennweite wie bei dickeren Linsen haben, damit die Bildweite in das Gehäuse passt. Um genügend Licht einzufangen muss man auch achsferne Lichtstrahlen auffangen und scharf bündeln, wie es nur dünne Linsen können. Durch ein geeignetes System mehrerer Linsen lassen sich diese Kriterien gleichzeitig erfüllen. Nachteilig an der Kombination von Linsen ist allerdings, dass die Lichtausbeute mit jeder weiteren Linse abnimmt, weil eine Linse immer auch ein wenig Licht „schluckt“.

Der Abstand zwischen Objektiv und Film lässt sich durch Drehung des Objektivs auf einem Schraubengewinde einstellen. Bei vorgegebener Brennweite des Objektivs gehören ja zu verschiedenen Gegenstandsweiten auch verschiedene Bildweiten. Mit dieser Einstellung entscheidet man daher, welche Gegenstandsweite mit ihrer zugehörigen Bildweite von der Linse genau bis zum Film reicht, welche Tiefen-Ebene des Motivs also scharf abgebildet wird.

Die **Blende** ist wie die verstellbaren Lochblenden aus der Physik-Sammlung aufgebaut: Blechplättchen sind verschiebbar so um den Mittelpunkt der Blende angebracht, dass sie dank eines geeigneten Mechanismus ein variabel großes, etwa kreisförmiges Loch in der Mitte offen lassen. Die Einstellung der Blende hat zwei Auswirkungen: Die Blendenöffnung beeinflusst die einfallende *Lichtmenge* und die *Tiefenschärfe*.

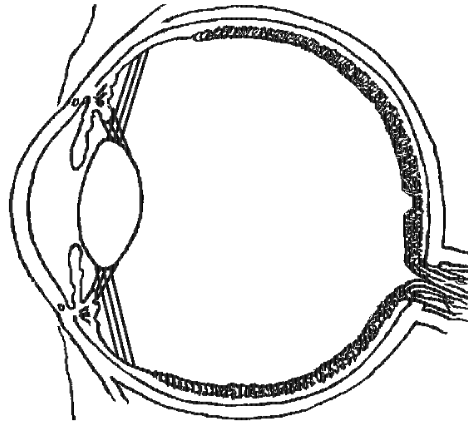
Neben dem Abstand Objektiv-Film und der Blendenöffnung ist die Belichtungszeit (=Dauer, wie lange

der Verschluss geöffnet wird) die dritte für die technische Bildqualität entscheidende Einstellung. Wie die Blende beeinflusst auch die Belichtungszeit die einfallende Lichtmenge und die Schärfe. Erhöht man nämlich die Belichtungszeit, kann sich das Motiv während der Belichtung um größere Strecken bewegen: das Bild verwischt.

Neben diesen optischen Gegebenheiten setzt auch noch die Chemie der Schärfe Grenzen: Da die lichtempfindlichen Stoffe des Filmes körnerförmig vorliegen und bei Lichteinfall immer nur ganze Körner chemisch umgesetzt werden, kann ein Bild nicht schärfer sein als es die Korngröße vorgibt.

In einer Digitalkamera befindet sich anstelle des Films ein Raster aus winzigen elektronischen Schaltungen, deren elektrische Eigenschaften sich je nach Lichteinfall ändern, und die somit jeweils ein Fleckchen Bild („Pixel“) auf elektrische Weise erfassen.

11.3.3 Menschliches Auge



Querschnitt durch rechtes Auge, waagerechte Schnittebene, von oben gesehen

zu ergänzen:

- Beschriftung

Die Gegenstände in der Außenwelt werden durch das „Linsensystem“ aus **Hornhaut**, Kammerwasser und **Linse** auf die **Netzhaut** abgebildet. Den Hauptbeitrag zur gesamten Lichtbrechung liefert dabei die Grenzfläche zwischen Luft und Hornhaut. Die Brechungen zwischen Hornhaut und Wasser, zwischen Wasser und Linse, sowie zwischen Linse und Glaskörper fallen geringer aus, weil die hier aneinandergrenzenden Stoffe recht ähnliche optische Dichten haben.

Die Netzhaut enthält lichtempfindliche Zellen, die dem Bild entsprechend gereizt werden und über den **Sehnerv** entsprechende Signale an das Gehirn weitergeben.

Von vorne betrachtet kann man sich den Mechanismus zur Veränderung der Linsenkrümmung schematisch wie nebenstehend gezeichnet vorstellen: Die Linse **L** ist über die Zonulafasern am ringförmigen Ziliarmuskel **Z** befestigt. Der Ziliarmuskel ist über ein elastisch gespanntes Gewebe **G** an der Augen-Innenwand befestigt. Da das Auge mit dem Glaskörper prall gefüllt ist, behält diese Wand ihre Form weitgehend bei und kann daher im Folgenden als fest angesehen werden.

Die Linse alleine wäre in entspanntem Zustand relativ dick, hätte also eine große Krümmung und in der gezeichneten Sicht einen kleinen Durchmesser. Aufgrund der Spannung des Gewebes **G** wird dieser Linsendurchmesser jedoch bei entspanntem Ziliarmuskel auseinandergezogen. Die Linse ist dann abgeflacht, ihre Brennweite groß und damit geeignet, um ferne Gegenstände abzubilden.

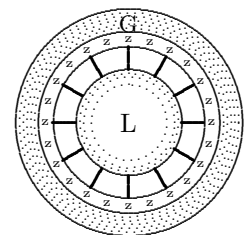
Zieht sich der Ziliarmuskel zusammen, werden außen das Gewebe **G** weiter gedehnt und innen die Zonulafasern entspannt, so dass die Linse sich wieder zu ihrem entspannten Grundzustand hin zusammenziehen kann. In diesem Zustand ist die Brennweite klein, so dass nahe Gegenstände scharf abgebildet werden.

Das Bild auf der Netzhaut kann allerdings nicht schärfer werden als durch den Abstand unabhängiger Sehzellen vorgegeben. Im Bereich des **gelben Fleckes** sieht man am schärfsten. Dorthin fällt das Licht von den Gegenständen, auf die man seinen Blick richtet.

Der **blinde Fleck** ist die Stelle, durch die der Sehnerv das Auge verlässt. Hier gibt es überhaupt keine Sehzellen! Mit den Bild-Informationen aus der Umgebung dieses Fleckes ergänzt das Gehirn das fehlende

Zur Regulierung der Helligkeit kann sich die **Regenbogenhaut**, auch **Iris** genannt, durch Muskelkraft verengen und erweitern. Bei großer Helligkeit ist die **Pupille** (=das Loch, das die Iris offen lässt) recht klein. Dadurch wird nicht nur die Netzhaut geschont, sondern auch das Bild schärfer und tiefenschärfer.

Natürlich müssen auch beim Auge für ein scharfes Bild die Brenn-, die Gegenstands- und die Bildweite aufeinander passen. Zu diesem Zweck kann die Linsenform mit Hilfe des ringförmigen Ziliarmuskels verändert werden. Bei stärkerer Krümmung nimmt die Brennweite der Linse ab.



Bildstück normalerweise sinnvoll, wenn es nicht sowie-so die fehlenden Bildteile dem anderen Auge entnehmen kann.

Du kannst Dich mit folgendem Versuch von der Existenz des blinden Fleckes überzeugen: Richte bei geschlossenem linken (rechten) Auge das rechte (linke) Auge auf das Kreuzchen (Kringelchen) in folgendem Bild, achte jedoch auf das Kringelchen, ohne den Blick vom Kreuzchen zu wenden! Verändere dann den Abstand zwischen dem Bild und Deinem Auge zwischen etwa 10 cm und 20 cm! Bei einer bestimmten Lage verschwindet das Kringelchen (Kreuzchen) in den blinden Fleck.

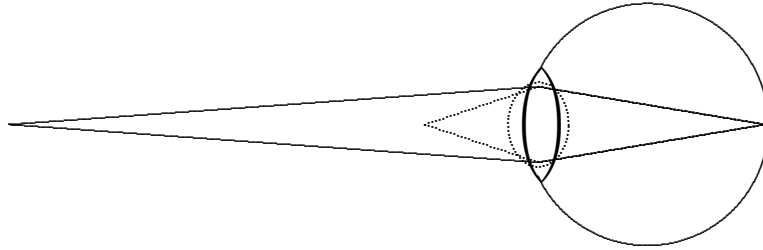
×

○

Normalsichtigkeit

Das Licht, das von einem weit entfernten Gegenstandspunkt in unser Auge fällt, läuft weniger stark auseinander als das Licht eines nahen Punktes. Es muss daher von der Linse auch nicht so stark gebrochen werden, um gebündelt auf die Netzhaut zu treffen. Zum Sehen weit entfernter Gegenstände darf die Linse nicht so stark gekrümmt sein.

Strahlengang und Linse für weit entfernten (durchgezogen) und nahen Gegenstand (punktiert):



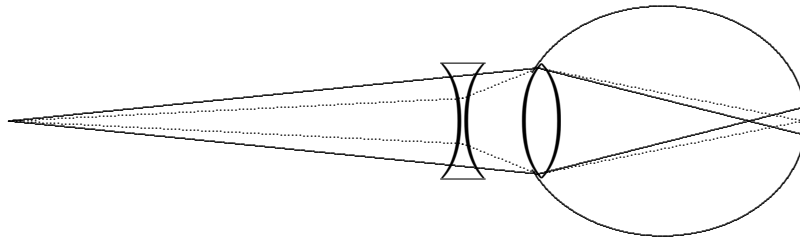
Kurzsichtigkeit

Bei manchen Menschen ist der Augapfel zu lang. Die Netzhaut ist von der Linse so weit entfernt, dass die Linse selbst bei geringstmöglicher Krümmung das Licht weit entfernter Gegenstände bereits vor der Netzhaut fokussiert. Bis zur Netzhaut läuft das Licht wieder auseinander und erzeugt daher ein unscharfes Bild: Der Mensch sieht weit entferntes unscharf. Positiv ausgedrückt: Er sieht kurz vor ihm befindliches besser. Er

wird deshalb **kurzsichtig** genannt.

Man kann ihm helfen, indem man ihm Zerstreuungslinsen vor die Augen setzt. Eine solche lässt die Strahlen vor dem Auge stärker auseinanderlaufen, so als ob sie von einem näher gelegenen Punkt kämen. Die eigentlich für die Augenlänge zu starke Augenlinse ist für diese stärker auseinanderlaufenden Strahlen nun gerade richtig.

Strahlengang mit (punktiert) und ohne (durchgezogen) Brille:

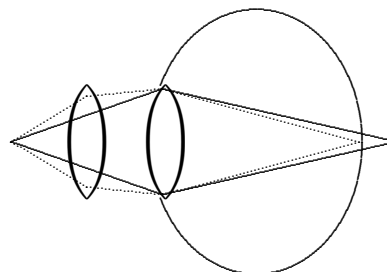


Weitsichtigkeit

Bei anderen Menschen ist der Augapfel zu kurz. Die Linse bricht selbst bei stärkster Krümmung das Licht naher Gegenstände nicht stark genug, um es bis zur Netzhaut zu fokussieren. Diese Menschen sehen daher Nahes unscharf, weit entferntes besser: Sie sind **weit-**

sichtig.

Man hilft ihnen mit Sammellinsen, die das zu stark auseinanderlaufende Strahlenbündel schon mal ein wenig sammeln, damit die eigentlich zu schwache Linse nur noch den Rest schaffen muss.



Brechkraft

In den letzten Abschnitten wurde statt von Brennweiten mehr davon gesprochen, wie stark eine Linse Licht brechen kann. Eine geringe Brennweite bedeutet entspricht einer hohen Brechkraft. Man definiert daher:

Die Brechkraft einer Linse ist der Kehrwert ihrer Brennweite.
Man misst die Brechkraft in der Einheit Dioptrie:

$$1 \text{ Dioptrie} := \frac{1}{1 \text{ Meter}} \quad \text{kurz:} \quad 1 \text{ dpt} := \frac{1}{1 \text{ m}}$$

11.4 Geräte zur optischen Vergrößerung

11.4.1 Der Sehwinkel

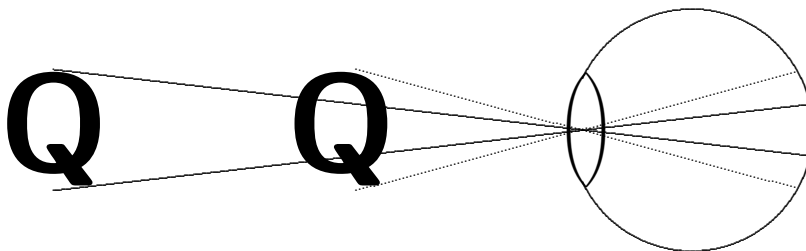
Was bedeutet eine Vergrößerung optisch?

?

Wie wir einen Gegenstand sehen, hängt letzten Endes davon ab, wie seine Lichtstrahlen unsere Netzhäute reizen. Wir sagen, dass wir z.B. eine Kuh *groß* sehen, wenn sie nahe vor uns steht. Dieselbe Kuh sehen wir *klein*, wenn sie weit weg ist. Welchen Unterschied macht die Entfernung für das Bild der Kuh auf der Netzhaut? (In der folgenden Zeichnung sind zur Deutlichkeit nur Mittelpunktstrahlen gezeichnet.)

zu ergänzen:

- Sehwinkel
- Netzhautbilder



Offensichtlich beleuchtet ein Gegenstand in geringerer Entfernung eine größere Fläche der Netzhaut als in größerer Entfernung. Mit dem größeren Netzhautbild erscheint er damit dem Betrachter auch größer. Färbe oben die gereizten Bereiche der Netzhaut für beide Stellungen der Kuh in verschiedenen Farben!

Die Größe, in der ein Gegenstand gesehen wird, lässt sich am besten beschreiben mit dem Winkel, unter dem das Netzhautbild von der Linsenmitte aus erscheint. Zeichne diese **Sehwinkel** für die beiden Stellungen der Kuh oben ein! (in den gleichen Farben wie die entspre-

chenden Netzhautbilder)

Ein Gerät, mit dem man Dinge vergrößert betrachten will, muss also den Sehwinkel vergrößern. Als Vergrößerungsfaktor des Gerätes bezeichnet man dabei das Verhältnis aus dem Sehwinkel mit Gerät und dem Sehwinkel ohne Gerät. Da ein Sehwinkel immer vom Abstand des betrachteten Gegenstandes abhängt, geht man bei der Angabe von Vergrößerungen von dem Abstand aus, in dem man Dinge mit der wenigsten Anstrengung sieht: etwa 25 cm.

11.4.2 Lupe

Eine Lupe ist das einfachste Vergrößerungs-„Gerät“ und besteht einfach aus einer Sammellinse. Man hält die Lupe so, dass der Gegenstand knapp innerhalb der Brennweite liegt, und betrachtet sein vergrößertes virtuelles Bild.

11.4.3 Mikroskop

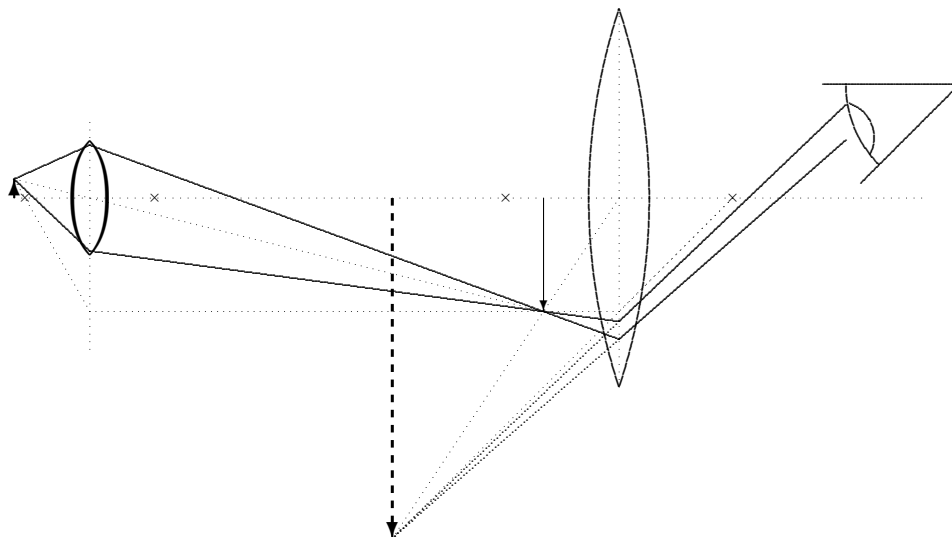
Ein Mikroskop besteht aus zwei Linsensystemen: Das Objektiv ist dem zu betrachtenden Objekt zugewandt, das Okular dem Auge (lat.: *oculus*) des Betrachters. Das Objektiv erzeugt vom winzigen, knapp außerhalb der Brennweite liegenden Gegenstand G ein vergrößertes reelles Zwischenbild Z, dessen virtuelles Bild V mit dem Okular wie mit einer Lupe betrachtet wird.

Das Wort Mikroskop setzt sich aus den beiden griechischen Wörtern *mikron*=klein und *skopein*=sehen zusammen.

zu ergänzen:

- Beschriftung

Wenn Du Dir das Objektiv und die Strahlen vom Objektiv zum Okular aus dieser Zeichnung wegdenkst, siehst Du den Strahlengang durch eine Lupe.



11.4.4 Teleskop

Im Wort Teleskop steckt neben *skopein*=sehen das ebenfalls griechische *telos*=Grenze, Ziel. Da Ziele meist in der Ferne liegen, deutet die Vorsilbe *tele-* meist auf die Ferne hin. Teleskope heißen zu deutsch nun jedoch nicht Fernseher, sondern Fernrohre. (Für Fernsehen benutzt man ja das lateinische *videre*=sehen.) Man unterscheidet bei Fernrohren mehrere Typen:

- **Galilei- oder holländische Fernrohre** mit einer Konvexlinse als Objektiv und einer Konkavlinse als Okular
- **Kepler- oder astronomische Fernrohre** mit Konvexlinsen als Objektiv und Okular; wird im Folgenden beschrieben
- **Spiegel-Teleskope** mit einem Hohl-Spiegel als Objektiv; Um lichtschwache Objekte zu beobachten, muss man mit dem Objektiv viel Licht (d.h. breite Lichtbündel) des Objektes auffangen, und große Spiegel lassen sich leichter herstellen als große Linsen.

Mit einem Fernrohr betrachtet man ja üblicherweise sehr weit entfernte, und eigentlich riesige Gegenstände, wie z.B. einen Planeten. Du kannst Dir daher vorstellen, dass für die folgende Zeichnung ein sehr großer Gegenstand sehr weit nach links entfernt auf der optischen Achse steht.

Wie bisher betrachten wir wieder nur das Licht, das von seiner Spitze ausgeht. Wegen der großen Entfernung treffen die Lichtstrahlen dieser Spitze etwa parallel zueinander bei uns ein.

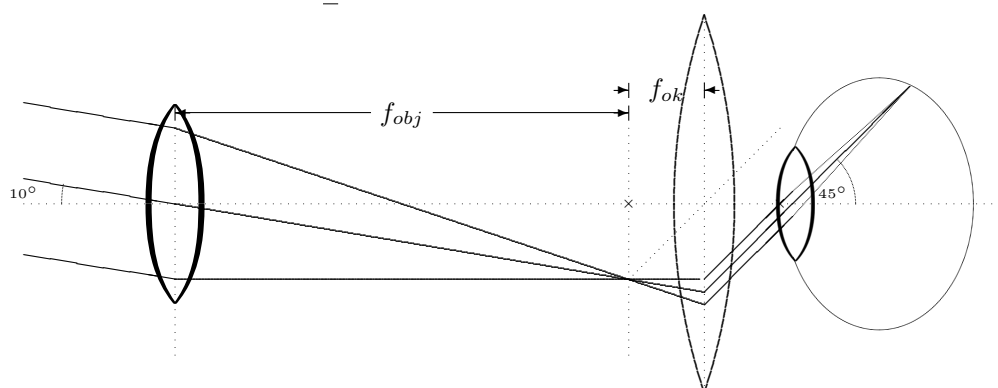
In einem Kepler-Fernrohr fängt man diese Lichtstrahlen mit einem langbrennweitigen Objektiv auf, in dessen Brennebene das Licht fokussiert wird. Das

Okular ist so positioniert, dass seine Brennebene mit der des Objektivs zusammenfällt. Daher verlassen die Lichtstrahlen das Okular wieder parallel zueinander. Da das Okular eine kleinere Brennweite hat, verlaufen die Strahlen danach aber unter einem größeren Winkel zur optischen Achse als vor dem Eintritt ins Fernrohr.

Den 10° weiten Himmelsbereich zwischen der optischen Achse und den einfallenden Strahlen sieht man nach dieser Zeichnung durch das Teleskop auf einen Winkelbereich von 45° auseinandergezogen. Der Sehwinkel zum Gegenstand in diesem Himmelsbereich wird also $4\frac{1}{2}$ -fach vergrößert.

zu ergänzen:

- Beschriftung
- Zwischenbild
- Netzhautbild

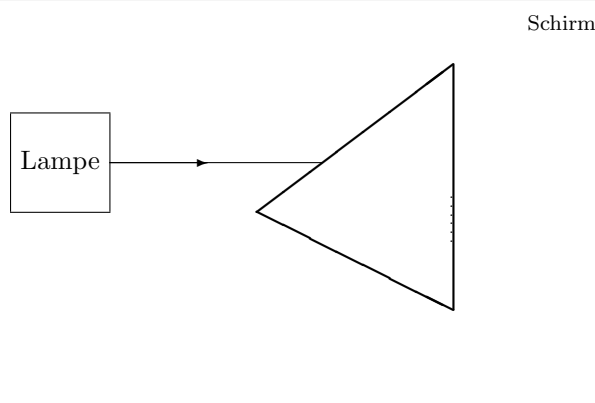


Farben

12.1 Spektren

V 12.1 Weißes Licht auf Glasprisma

Aufbau: Ein weißer Lichtstrahl wird wie abgebildet auf ein Glasprisma gerichtet.



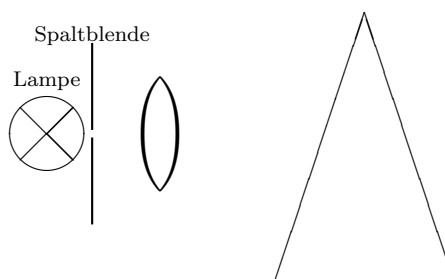
Beobachtung:

zu ergänzen:

- weitere Lichterscheinungen

Das auf dem Schirm zu sehende bunte Band nennt man (vom lateinischen Verb *spectare*=sehen abgeleitet) ein **Spektrum**. Die in ihm enthaltenen Farben heißen **Spektralfarben**. Offensichtlich treten aus dem Prisma verschiedenfarbige Lichtstrahlen in verschiedene Richtungen aus.

Um ein lichtstärkeres Ergebnis zu erhalten, lassen wir tatsächlich nicht nur ein strahlähnliches dünnes Bündel paralleler Lichtstrahlen auf das Prisma treffen, wie laut obiger vereinfachten Versuchsbeschreibung. Statt dessen bilden wir zunächst mit einer Sammellinse einen sehr hell durchleuchteten engen Spalt scharf auf den Schirm ab. In diesen Strahlenweg setzen wir dann das Prisma ein. Das Spektrum ergibt sich dann als Überlagerung aller verschiedenfarbigen, der Farbe entsprechend versetzten Spalt-Bilder.



zu ergänzen:

- Strahlenverläufe

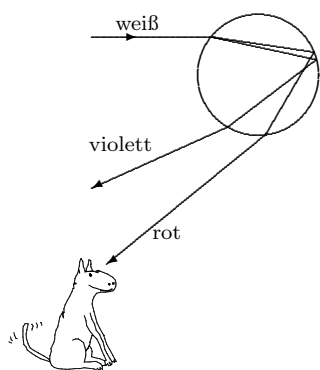
Man kann demnach aus dem weißen Licht z.B. spektralrotes Licht gewinnen, indem man in obigem Versuchsaufbau den Schirm an der rot beleuchteten Stelle durchbohrt: Durch das Loch fällt dann rotes Licht hinter den Schirm und kann dort weiter untersucht werden. Dabei zeigt sich, dass ein spektralfarbiger Strahl sich nicht weiter in andere Farben auffächern lässt.

Lässt man jedoch verschiedenfarbige Lichtstrahlen gleichzeitig auf dieselbe Stelle eines Schirmes fallen, ergeben sich dort neue Farben. Am meisten verblüffendes Beispiel: rotes und grünes Licht ergeben zusammen eine gelbe Beleuchtung. Vereinigt man auf diese Weise wieder alle Spektralfarben, die man aus weißem Licht gewonnen hat, erhält man wieder Weiß.

Es gibt Licht in allen Spektralfarben, das nicht aufgetrennt werden kann. Aus diesen Lichtern ist jedes andere Licht zusammengesetzt. Mit einem Glasprisma (u.a.) lassen sich die spektralen Farbkomponenten eines Lichtes trennen. Das Spektrum aller Spektralfarben reicht mit fließenden Übergängen von rot über orange, gelb, grün und blau zu violett.

Der Grund für diese Auftrennung liegt darin, dass spektralfarbiges Licht je nach Farbe i.A. unterschiedlich stark gebrochen wird. In obigem Versuch wird rotes Licht am wenigsten, violettes Licht am stärksten gebrochen.

Dieser Effekt unterschiedlicher Brechung in Abhängigkeit von der Farbe heißt **Dispersion** und tritt je nach Material mehr oder weniger stark auf. Er ist eine Folge davon, dass sich Licht in einem Medium je nach Farbe verschieden schnell ausbreiten kann.



Auch in Wasser tritt Dispersion auf, z.B. in sonnenbeschienenen Regentropfchen. Das weiße Sonnenlicht kann wie gezeichnet in ein Tröpfchen eintreten, innerhalb des Tröpfchens einmal (total-)reflektiert werden und wieder austreten. Durch die farbabhängigen Brechungen beim Ein- und Austritt verlaufen die verschiedenfarbigen Anteile des Sonnenlichtes danach in verschiedene Richtungen. In der nebenstehenden Zeichnung sind allerdings nur die Randstrahlen des Spektrums eingezeichnet.

Fängt nun ein Beobachter den gezeichneten roten Lichtstrahl auf, so geht der violette über ihn hinweg. Violettes Licht kann derselbe Beobachter aber von tiefer gelegenen Tröpfchen empfangen. Der Beobachter sieht also rot am oberen Rand des Spektrums, violett am unteren.

Die gleichen Erscheinungen spielen sich natürlich nicht nur in vertikaler Richtung ab, sondern in jeder. So entsteht für den Beobachter insgesamt ein intensiv farbiger (Haupt-) **Regenbogen**, der außen rot und innen violett ist.

Neben dem oben gezeichneten Lichtweg gibt es mit anderem Eintrittspunkt und -winkel noch einen weiteren Weg des Lichtes durch ein Wassertropfchen, der zu einem weiteren Regenbogen führt. Dieser Neben-Regenbogen ist allerdings weniger intensiv. Er verläuft mit umgekehrter Farbenfolge außerhalb des Haupt-Regenbogens. Die Ursache für die Umkehr der Farbfolge ist, dass die Lichtstrahlen dieses Bogens innerhalb der Tröpfchen eine zweite Reflexion erfahren.

Ohne nennenswerten Aufwand kannst Du Dich davon überzeugen, dass auch Deine Augenlinsen

Licht unterschiedlicher Farbe unterschiedlich brechen. Blicke mit nur einem Auge auf eine scharfe vertikale Grenzlinie zwischen einer sehr hellen, möglichst weißen und einer dunklen Fläche. (Zum Beispiel auf die Innenkante eines Fensterrahmens als Grenze zwischen dunklem Fensterrahmen und hellem Himmel) Schiebe dann von links oder rechts her Deinen Finger dicht vor dem Auge bis fast an diese Grenzlinie heran. Nun solltest Du an dieser Linie einen dünnen rötlichen (bei Heranschieben des Fingers von der dunklen Seite her) oder bläulichen (von der hellen Seite her) Saum sehen. Kannst Du erklären, wie das zustandekommt?

Untersucht man das Licht verschiedener Quellen auf seine spektrale Zusammensetzung, so stellt man noch fest:

Glühende Lichtquellen enthalten alle Spektralfarben: Sie haben ein **kontinuierliches** Spektrum.

Leuchtstoffröhren u.ä. enthalten nur einzelne Spektralfarben: Sie haben ein **diskretes** Spektrum.

Das bisher besprochene Spektrum ist nur der *für uns sichtbare* Teil eines viel breiteren Spektrums. Die *unsichtbaren* Teile reizen unsere menschlichen Netzhäute einfach nicht zu einer Lichtempfindung.

Unmittelbar an Rot schließt sich **Infra-rot** (IR) an. Infrarote Strahlung ist für uns auch als „Wärmestrahlung“ spürbar.

Unmittelbar an Violett schließt sich **Ultraviolett** (UV) an. UV-Strahlung regt die Bildung von Pigmenten und von Vitamin D in unserer Haut an. Sie kann aber auch organisches Gewebe schädigen.

Während also am Strand die Infrarot-Strahlung für angenehme Wärme sorgt, kümmert sich die Ultraviolett-Strahlung um die Bräune und den anschließenden Sonnenbrand.

12.2 Farbwahrnehmung des Menschen

12.2.1 Die lichtempfindlichen Zellen der Netzhaut

Die Netzhaut enthält zwei verschiedene Sorten lichtempfindlicher Zellen. Aufgrund ihrer Form nennt man die eine Sorte **Zapfen** und die andere **Stäbchen**. Stäbchen und Zapfen reagieren auf Licht unterschiedlicher Farbe unterschiedlich gut.

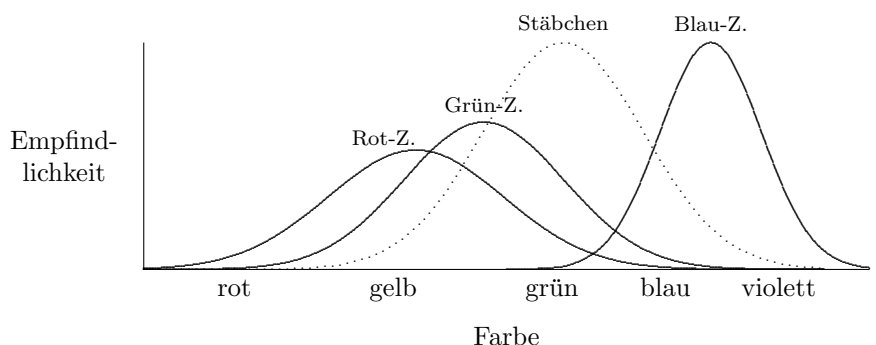
Stäbchen

- Stäbchen reagieren am besten auf blaugrünes Licht.
- Dabei genügt bereits recht wenig Licht, um die Stäbchen zu überreizen. Das heißt, dass schon bei geringer Helligkeit eine weitere Zunahme der Helligkeit den Reiz nicht mehr nennenswert erhöht.
- Stäbchen gibt es kaum im Bereich des gelben Fleckes der Netzhaut.

Zapfen

- Bei den Zapfen gibt es drei Varianten, von denen eine am empfindlichsten auf blaues, eine am besten auf grünes und die dritte stärker als die anderen auf rotes Licht reagiert. Auf andere Farben reagiert ein Zapfen, grob gesagt, umso weniger empfindlich, je weiter diese andere Farbe von der „Hauptfarbe“ des Zapfens im Spektrum entfernt ist.
- Alle Zapfen reagieren aber erst bei relativ großer Lichtintensität.
- Zapfen gibt es in einem ausgedehnten Bereich um den gelben Fleck. Das Gebiet mit Blau-Zapfen reicht dabei am weitesten nach außen, das der Grün-Zapfen am wenigsten weit.

Die folgenden Kurven zeigen, wie empfindlich die 4 Arten von Sehzellen auf Licht verschiedener Spektralfarben reagiert. Da wir aber nicht näher darauf eingehen können, was genau auf den Achsen aufgetragen ist, deutet dieses Diagramm die Zusammenhänge nur sinngemäß an. Es zeigt keine aus Messwerten hervorgehende Kurve. Die gepunktet gezeichnete Kurve für die viel empfindlicheren Stäbchen darf jedenfalls nicht im selben Empfindlichkeits-Maßstab wie die Zapfen-Kurven gesehen werden. Die x-Achse könntest Du noch mit einem Farbspektrum statt Zahlen „beschriften“.



12.2.2 Tages-/Farbsehen

Bei großer Helligkeit – tagsüber – sind alle Stäbchen überreizt. Sie können dem Gehirn folglich keine Information über das Licht liefern, von dem sie getroffen werden. Sie tragen daher nicht zum Sehen bei.

Die Zapfen werden je nach Farbe und Intensität des auf sie fallenden Lichtes mehr oder weniger stark gereizt und informieren das Gehirn über die Reizstärke. Dabei kann aber das Gehirn dieser Information von z.B. einem Rot-Zapfen nicht mehr entnehmen, ob der Rot-Zapfen mit schwachem orangenen Licht, mit intensiverem grünen Licht oder mit noch intensiverem blauen Licht gereizt worden ist.

Warum nicht? Nimm einmal an, dass die Reizstärke sich als Produkt aus der Intensität des Lichtes und der Empfindlichkeit für seine Farbe ergibt. Der Reizstärke als dem Wert des Produktes kann man nicht mehr ansehen, welcher Faktor wie groß war: So wie $2 \cdot 6 = 4 \cdot 3 = 6 \cdot 2 = \dots = 12$, kann schwaches Licht bei hoher Empfindlichkeit die gleiche Reizstärke ergeben wie starkes

Licht bei niedriger Empfindlichkeit.

Wenn also die Rot-Zapfen in einem Netzhautgebiet eine bestimmte Reizstärke melden, sagt dies alleine über die Farben in diesem Bildbereich noch gar nichts aus. Für die folgenden Überlegungen gehen wir aber wenigstens davon aus, dass das ganze Gebiet gleichmäßig hell und einfarbig beleuchtet ist. Die Farbe und die Intensität kennen wir aber noch nicht.

Das Gehirn erhält aus demselben Gebiet auch die Reizstärken der Grün- und der Blau-Zapfen. Ist das dort eintreffende Licht z.B. spektralgelb, so reizt es neben den Rot-Zapfen auch die Grün-Zapfen. Da für gelbes Licht beide Zapfensorten etwa gleich empfindlich sind, werden auch beide Sorten durch das gleiche gelbe Licht etwa gleich stark gereizt. Die Blau-Zapfen werden von gelbem Licht nicht gereizt. Die *Kombination* aus viel rot, viel grün und kein blau wird daher vom Gehirn als „gelb“ erkannt.

Jede menschliche Farbwahrnehmung ist zunächst die Kombination aus einem Rot-, einem Grün- und einem Blau-Anteil.

Die Kombination aus viel rot, kein grün und viel blau kann nicht durch eine Spektralfarbe hervorgerufen werden. Dies kann nur eine Mischung aus rötlichen und blauvioletten Lichtstrahlen, die das Gehirn als „purpur“ erkennt.

Eine Mischung, die alle drei Zapfensorten gleichermaßen reizt, wird als „grau“ empfunden, bei hoher Reizung als „weiß“. Völlig reizlos ist hingegen „schwarz“.

Praktisch jedes Licht, das wir aus unserer Umwelt

auffangen, ist aus vielen Spektralfarben gemischt und spricht alle Zapfensorten zumindest ein wenig an. Werden z.B. alle Zapfensorten ein wenig gereizt, die Rot-Zapfen darüberhinaus jedoch besonders, so ergibt sich im Gehirn aus dem Weiß-Eindruck durch die gemeinsame Reizung und dem Rot-Eindruck durch die hohe Reizung allein der Rot-Zapfen der Gesamteindruck „rosa“. Ein Farbwert kann demnach auch beschrieben werden durch

- den Farbton, quasi die Farbe ohne den Weiß-Anteil
- die Sättigung, die umso geringer ist, je höher der Weiß-Anteil ist
- die Helligkeit, die einer Art Gesamt-Reizstärke entspricht

Beispiel: Das Karminrot in Deinem Wasserfarben-Kasten ist alleine und kräftig auf weißes Papier aufgetragen ziemlich gesättigt. Mischst Du Deckweiß dazu, wird die Farbe vor allem entsättigt. Mischst Du Schwarz dazu, sinkt die Helligkeit.

12.2.3 Verarbeitung der Farbreize im Gehirn

Farbeindrücke hängen nicht nur von den Rot-, Grün- und Blauanteilen ab, die die Zapfen registrieren. Diese sind quasi erst die Eingangsdaten, aus denen das Gehirn Farbeindrücke ermittelt. Dabei spielen für den Farbeindruck an einer Bildstelle sowohl die Farben der *Umgebung* als auch die *Vorgeschichte* eine Rolle.

Das Gehirn ist äußerst leistungsfähig darin, *Gegenstände* zu erkennen. Dazu gehört zunächst, in den Netzhautbildern zusammenhängende Flächen zu erken-

nen und gegen benachbarte Flächen abzugrenzen. Das Gehirn neigt daher dazu, ähnliche Farbeindrücke innerhalb einer (vermeintlich) zusammenhängenden Fläche einander anzugleichen, während es Unterschiede in den Farbeindrücken beiderseits einer (vermeintlichen) Konturlinie eher verstärkt. Es sind nämlich eigentlich nur diese *räumlichen* Kontraste, an denen Formen und damit Gegenstände erkannt werden können.

Zum Erkennen von *Bewegungen* ist es wichtig, *zeitliche* Veränderungen wahrzunehmen. Eigentlich erfordern vor allem die Stellen der Umgebung Aufmerksamkeit, an denen sich etwas bewegt, etwas passiert.

Blickst Du z.B. eine Minute lang mit starrem Blick auf eine gelbe Fläche, so „gewöhnt“ sich das Gehirn an das Gelb an dieser Stelle. Blickst Du dann auf eine weiße Fläche, empfangen dieselben Stellen der Netzhaut, an denen vorher vor allem die Rot- und Grünzapfen erschienen worden sind, auch blaues Licht. Im Vergleich

zu vorher gibt es also nur einen Zuwachs an Blau, der – vom Gehirn entsprechend betont – zu einem blauvioletten sogenannten „Nachbild“ führt.

Ein solches Nachbild kannst Du sogar sehen, wenn Du die Augen schließt. Schon das so verursachte Wegfallen des Gelb wird vom Gehirn als Zugewinn von Blau gewertet. Beim Wechsel von vorher *viel Gelb und kein Blau* zu nun *weder Gelb noch Blau* hat Blau doch immerhin von einem Rückstand auf einen Gleichstand „aufgeholt“!

12.2.4 Dämmerungssehen

Wenn es recht dunkel ist – in der Dämmerung – reagieren die Zapfen wegen der geringen Lichtintensitäten praktisch nicht. Die Stäbchen alleine liefern dem Gehirn Informationen. Aus diesen kann das Gehirn allerdings keine Farbinformation entnehmen. Wie für die Zapfen erklärt enthält ja erst eine Kombination der Informationen aus verschiedenen empfindlichen Zellen Farbinformation.

Es ist daher nicht möglich, in der Dämmerung Farben zu erkennen.

Es ist wegen fehlender Stäbchen im gelben Fleck auch schwieriger, in der Dämmerung Dinge genau zu betrachten. Die Stelle, auf die man den Blick richtet, sieht man nicht so genau wie deren Umgebung!

12.3 Farbmischung

12.3.1 Additive Farbmischung

Strahlt man mit verschiedenfarbigen Lichtern gleichzeitig auf einen weißen Schirm o.ä., so addieren sich diese Farben, sie mischen sich *additiv*. Der weiße Schirm wirft ja alle einfallenden Lichtstrahlen in den Raum zurück, u.a. in unsere Augen. Die Zapfen der Netzhaut erfahren daher alle Reizungen, die die Lichter einzeln hervorrufen würden, zusammen. Deshalb ergeben rotes und grünes Licht zusammen den Farbeindruck viel rot + viel grün + wenig blau, also insgesamt einen (etwas entsättigten) gelben Farbeindruck.

In Techniken zur Bildwiedergabe wie Fernsehen ist es wichtig, möglichst viele Farbwerte mit möglichst we-

nig Aufwand wiedergeben zu können. Es ist daher günstig, möglichst alle Farben aus möglichst wenig Grundfarben mischen zu können.

Mit einem bestimmten Rot, das genau auf die Empfindlichkeitskurven der Zapfen abgestimmt ist, lassen sich die Rot-Zapfen gezielt reizen, ohne die Grün- und Blau-Zapfen mehr als unvermeidlich mitzureizen. Ebenso gibt es ein optimales Grün und ein optimales Blau für die entsprechenden Zapfen. Aus diesen drei Grundfarben lassen sich die allermeisten Farbeindrücke additiv mischen.

Grundfarben der additiven Farbmischung sind rot, grün und blau.

Es lassen sich damit zwar alle Farbtöne mischen, jedoch nicht in voller Sättigung. Unsere Augen und Gehirne lassen sich aber leicht genug täuschen, dass dieser Mangel z.B. beim Fernsehen nie auffällt.

Bei gewöhnlicher Helligkeit im Raum akzeptieren wir neben einer hellen Farbe sogar das Grün des nicht leuchtenden Bildschirms als Schwarz. Ehrlich: Der Bildschirm ist objektiv gesehen nie dunkler als im ausgeschalteten Zustand!

Wenn Du wissen willst, welche Farbtöne diese meistbenutzten additiven Grundfarben haben, schau Dir einfach ein Fernsehbild ganz nah an! Achte insbesondere darauf, wie gelbe Bildflächen dargestellt sind!

12.3.2 Subtraktive Farbmischung

Deine Erfahrungen mit Wasserfarben sind wahrscheinlich anders: dort ergeben rot und grün eher braun. (eine Farbe, die es als Lichtfarbe gar nicht gibt: Dinge sehen braun aus, wenn sie zwar farbig, aber recht dunkel leuchten, zumindest aber neben viel helleren Flächen gesehen werden.)

Ein roter (nicht selbst leuchtender) Gegenstand wirft von einfallendem Licht vor allem die rötlichen Anteile zurück, also z.B. rot und etwas orange. Die übrigen Anteile vom einfallenden Licht werden absorbiert, d.h. „verschluckt“: gelb, grün, blau, violett. Entsprechend: Ein grüner Körper verschluckt alle Farben außer grün und evtl. die im Spektrum angrenzenden gelb und blau.

Ein „roter“ Körper sieht also nur dann rot aus, wenn er mit Licht beleuchtet ist, das rot enthält. In spektral orangenem Licht kann er nur orange aussehen. Und falls er z.B. blau wirklich vollständig verschluckt, sieht er in spektralblauem Licht sogar schwarz aus.

Legt man also z.B. rotes und grünes Filterpapier übereinander auf weißes Papier, so kann vom einfallenden weißen Licht kein Anteil beide Filter passieren, um vom weißen Papier zurückgeworfen zu werden: Die Farben, die den roten Filter passieren, werden vom grünen absorbiert und umgekehrt. So idealisiert sähe man al-

so nur schwarz. Tatsächlich absorbieren die Filter eine Farbe aber nie vollständig, so dass doch ein wenig Licht durch beide Filter zum Papier hin und wieder zurück durchkommt. Diese Art der Farbmischung liegt auch bei Wasserfarben u.ä. vor. Da aus dem einfallenden Licht Anteile weggeschluckt werden, nennt man diese Art der Mischung *subtraktiv*.

Natürlich ist es auch hier günstig, aus möglichst wenigen Grundfarben alle anderen mischen zu können. Damit aber beim Mischen von dem einfallenden, üblicherweise weißen Licht noch etwas übrig bleibt, dürfen die Grundfarbstoffe jeweils nur wenig absorbieren.

Auch hier ist es günstig, sich an den Zapfenempfindlichkeiten zu orientieren. Mit einem Grundfarbstoff, der nur das darauf optimierte Rot absorbiert, einem, der nur das Grün absorbiert, und einem, der nur das Blau absorbiert, lassen sich alle Farbtöne gut herstellen. Die Farbe des rot-absorbierenden Grundstoffes ergibt sich als additive Mischfarbe der übrigen Farben als „cyan“. Die anderen beiden Grundfarben ergeben sich als „gelb“ (=nicht blau) und „magenta“ (=nicht grün).

Grundfarben der subtraktiven Farbmischung: cyan, gelb und magenta

Diese Farben kannst Du Dir in Deinem Farbkasten anschauen. Vielleicht hast Du bisher als Grundfarben zum Mischen rot, gelb und blau gelernt und benutzt. Dies liegt dann aber nur daran, dass man Dir die ungewöhnlicheren Namen cyan und magenta ersparen wollte, die ja doch bläulich bzw. rötlich sind. Du erhältst aber mit cyan, gelb und magenta sattere Mischfarben als mit anderen Farbtönen.

Mischungsbeispiel: „Richtiges“ rot erhält Du aus

magenta und gelb, denn der rote Spektralbereich ist der einzige, den diese beiden Grundfarben gemeinsam durchlassen. Je nach Mischungsverhältnis bekommst Du mit höherem gelb-Anteil auch ein orange.

Im Druckgewerbe verwendet man neben cyan, gelb und magenta zur Herstellung der Farbtöne auch noch schwarz zur Regulierung der Helligkeit, insbesondere, um Kontraste zu erzeugen.

Dies nennt sich dann **Vierfarbdruck**.

Es gibt insgesamt betrachtet nicht *die* Farbe eines Körpers. Was wir als seine Farbe empfinden, variiert zumindest mit der spektralen Zusammensetzung der Beleuchtung, mit den Farben der Umgebung und mit unseren vorhergehenden Wahrnehmungen.

Akustik

13.1 Das Wesen von Schall

13.1.1 Erste Erfahrungen

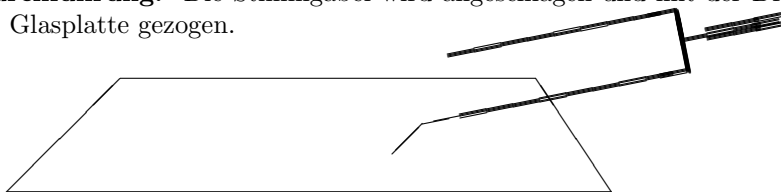
Alle Körper, die Schall erzeugen, schwingen. Du kannst z.B. an Deinem Kehlkopf eine Vibration spüren, wenn Du etwas sagst, singst oder brummst. Bei einer gezupften Saite eines Musikinstrumentes sieht man sogar, dass die Saite nicht stillsteht, sondern sich schnell hin- und

herbewegt. Die Bewegung ist allerdings zu schnell, um mit dem Auge genau verfolgt zu werden. Jede einzelne der schnellen Schwingungen einer angeschlagenen Stimmgabel macht aber der folgende Versuch sichtbar.

V 13.1 Aufzeichnung einer Stimmgabelschwingung

Aufbau: An einem Zinken einer Stimmgabel ist eine Drahtspitze angebracht. Eine Glasplatte wurde angerußt.

Durchführung: Die Stimmgabel wird angeschlagen und mit der Drahtspitze über die Glasplatte gezogen.



Beobachtung: Die Spitze kratzt in die Rußschicht die eingezeichnete Spur, offensichtlich die Spur einer Schwingung.

Von der schwingenden Schallquelle aus muss der Schall erst zu unserem Ohr gelangen, bevor wir ihn hören können. Dabei spielt die Luft eine wichtige Rolle:

V 13.2 Schallquelle im Vakuum

Aufbau: Ein Schallquelle befindet sich unter einer Vakuumpumpe.

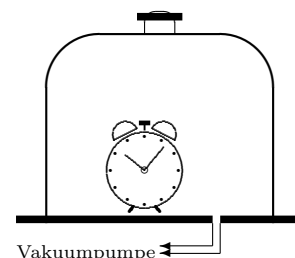
Durchführung: /Beobachtung:

Ergebnis:

zu ergänzen:

- Zugrichtung
- Ruß
- Spur

Schall wird von schwingenden Körpern erzeugt.



Aus der Erfahrung mit Echos wissen wir bereits, dass Schall zum Zurücklegen einer Strecke Zeit braucht. Seine Geschwindigkeit werden wir demnächst messen. Die Existenz von Echos sagt uns auch, dass Schall reflektiert werden kann. Dabei gilt wie in der Optik, dass Ein- und Ausfallswinkel gleich sind. Allerdings können wir beim hörbaren Schall keine schmalen Schall-„Strahlen“ erzeugen, an denen wir das Reflexionsgesetz so ein-

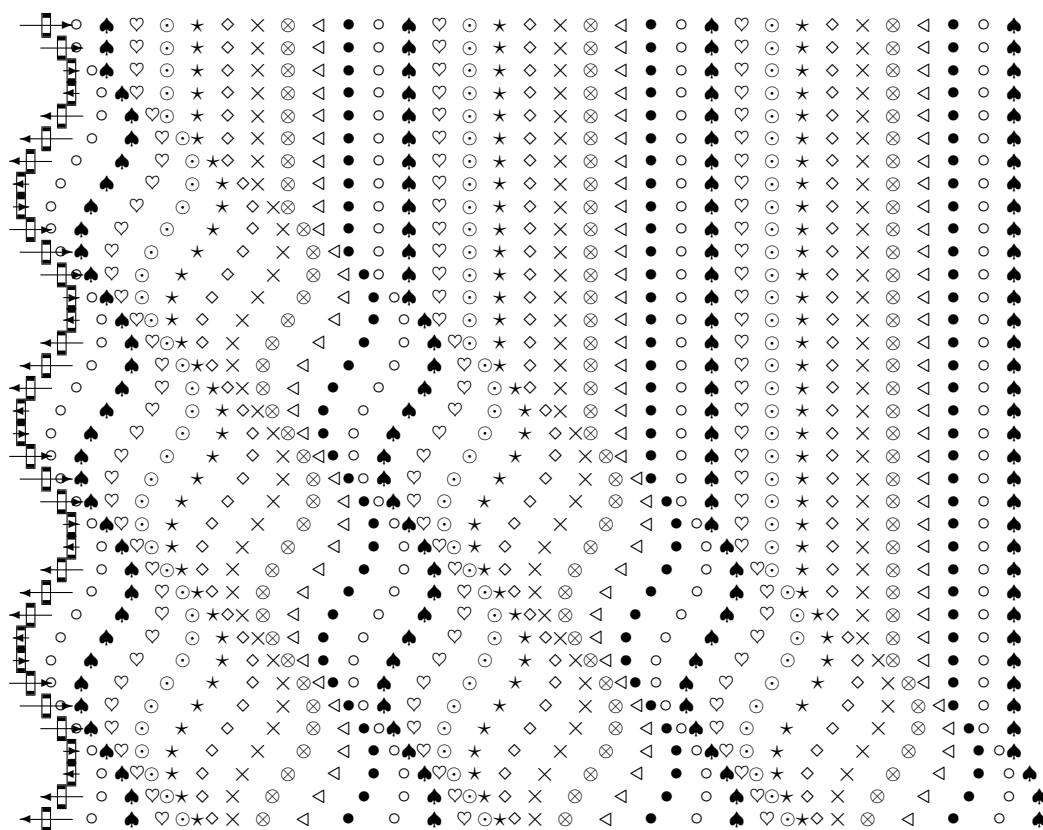
fach nachmessen könnten. Eine Blende wie in der Optik funktioniert nicht: Schall breitet sich auch hinter einer Blende wieder nach allen Richtungen aus, und damit regelrecht „um die Ecke“ statt geradlinig wie Licht. Wie kann man dann aber überhaupt von Ein- und Austrittswinkeln reden? Nun: Für das Reflexionsgesetz bei Schall muss man unter der Ausbreitungsrichtung des Schalls die Richtung verstehen, in der der Schall am lautesten zu hören ist.

13.1.2 Atomistische Vorstellung

Die experimentellen Befunde legen die im folgenden Bild dargestellte Vorstellung nahe: Alle Materie, z.B. auch Luft, besteht aus kleinsten Teilchen, die im Bild als Kringelchen, Kreuzchen, Sternchen etc. dargestellt sind. Das Bild muss nun zeilenweise gelesen werden. Die erste Zeile zeigt ganz links (als Rechteckchen) eine Schallquelle, z.B. eine Spitze einer Stimmgabel. Rechts daneben ist eine Reihe von Teilchen dargestellt. (In Wirklichkeit sind die Teilchen der Luft nicht so schön

gleichmäßig aufgereiht, sondern völlig ungeordnet, was aber für die folgenden Überlegungen nicht so wichtig ist.) Das Pfeilchen durch die Schallquelle deutet deren momentane Bewegung an.

Jede folgende Zeile zeigt dieselbe Teilchen-Reihe einen kurzen Moment später. Man erkennt zunächst, wie die Schallquelle das erste Kringelchen nach rechts drückt.



Die Luft im Bereich des Kringelchens und des Pik (♣) wird also verdichtet, der Luftdruck in diesem Bereich steigt. Das Pik wird dabei – in der 4. Zeile erkennbar – nach rechts gedrückt, wodurch die Verdichtung nach rechts zunächst zum Herzchen und dann so weiter fortschreitet.

Kringelchen dem Druck von rechts in den freigegebenen Raum an der Schallquelle aus, wie ab der 5. Zeile zu sehen. Solange dieser Raum noch nicht wieder gefüllt ist, herrscht in ihm ja Unterdruck. In diesen wird kurz darauf – ab der 7. Zeile zu sehen – auch das Pik getrieben. Dadurch schreitet eine Entspannung der Luft nach rechts fort.

Beim Zurückschwingen der Schallquelle weicht das

So laufen durch die Schwingung der Schallquelle angeregt abwechselnd Verdichtungen und Entspannungen der Luft, Hoch- und Niederdruck-Phasen nach rechts. Eine solche fortschreitende Schwankung nennt man allgemein eine Welle, in diesem Fall eine **Schallwelle**.

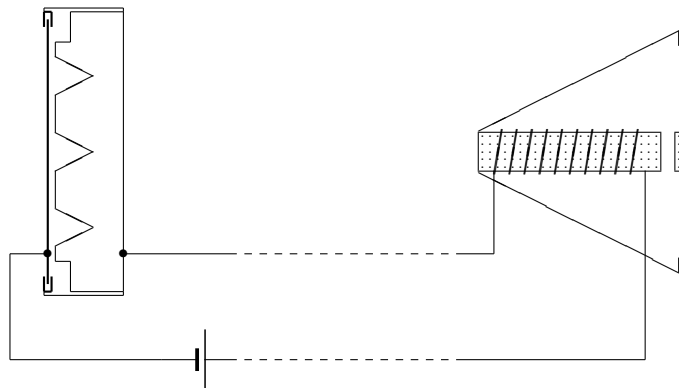
Ein ähnliches Wellenverhalten kannst Du übrigens an einer langen Schraubenfeder beobachten, wenn Du ein Ende davon abwechselnd eindrückst und ausziehst.

An einem Gegenstand rechts von der abgebildeten

Teilchen-Reihe – z.B. an Deinem Trommelfell im Ohr – würden ab der dritt-letzten Momentaufnahme abwechselnd Hoch- und Nieder-Druck-Phasen ankommen und auf den Gegenstand einwirken. (Das Pik würde abwechselnd gegen den Gegenstand drücken und wieder von ihm zurückweichen.) Der Gegenstand würde dadurch zu Schwingungen im Rhythmus der ankommenden Schwankungen angeregt, er würde den Schall *empfangen*.

13.2 Elektrik rund um Schall

13.2.1 Elektrische Schallübertragung



zu ergänzen:

- Beschriftung
- Kohlekörner
- Stromwege
- Schallwellen

Mikrofon: Der eintreffende Schall regt die Membran zum Mitschwingen an. Die Kohlekörner zwischen Membran und Kohleblock werden dadurch im Rhythmus des Schalls verdichtet und gelockert. Bei Verdichtung haben die Körner untereinander, zur Membran und zum Kohleblock mehr und größere Kontaktflächen und leiten daher elektrischen Strom besser. Legt man zwischen Membran und Kohleblock eine Spannung an, erhält man daher einen Strom durch das Mikrofon, dessen Stärke im Rhythmus des Schalls schwankt.

Kohleblock werden dadurch im Rhythmus des Schalls verdichtet und gelockert. Bei Verdichtung haben die Körner untereinander, zur Membran und zum Kohleblock mehr und größere Kontaktflächen und leiten daher elektrischen Strom besser. Legt man zwischen Membran und Kohleblock eine Spannung an, erhält man daher einen Strom durch das Mikrofon, dessen Stärke im Rhythmus des Schalls schwankt.

Legt man zwischen Membran und Kohleblock eine Spannung an, erhält man daher einen Strom durch das Mikrofon, dessen Stärke im Rhythmus des Schalls schwankt.

Lautsprecher: Ein Elektromagnet, der von einem Strom schwankender Stärke durchflossen wird, erzeugt ein in seiner Stärke ebenso schwankendes Magnetfeld. Dadurch wird auch das Weicheisenstück an der Membran im Rhythmus der Stromschwankung angezogen und wieder zurückfedern gelassen. Die Hin- und Herbewegungen der Membran erzeugen Schall, dessen Druckschwankungen den Stromschwankungen entsprechen.

So kann z.B. ein Musikstück viel weiter übertragen werden, als man die Musik direkt hören kann. Die elektrische Übertragung ist mit Verstärkern über sehr weite Strecken möglich, mit der Funktechnik sogar bis zu anderen Planeten. Außerdem können die elektrischen Schwankungen leicht aufgezeichnet und erst später wiedergegeben werden. Schließlich lässt sich das Musikstück auch noch vielfältig elektronisch bearbeiten, wenn es erst einmal in Stromschwankungen übersetzt worden ist.

Allerdings muss man beim Übertragen, Aufzeichnen, Bearbeiten und Wiedergeben recht viel Aufwand treiben, um alle hörbaren Feinheiten des Musikstückes dem Original getreu zu erhalten.

Hohe Wiedergabe-Treue nennt man übrigens auf englisch *High Fidelity*.

13.2.2 Oszilloskop-Bilder

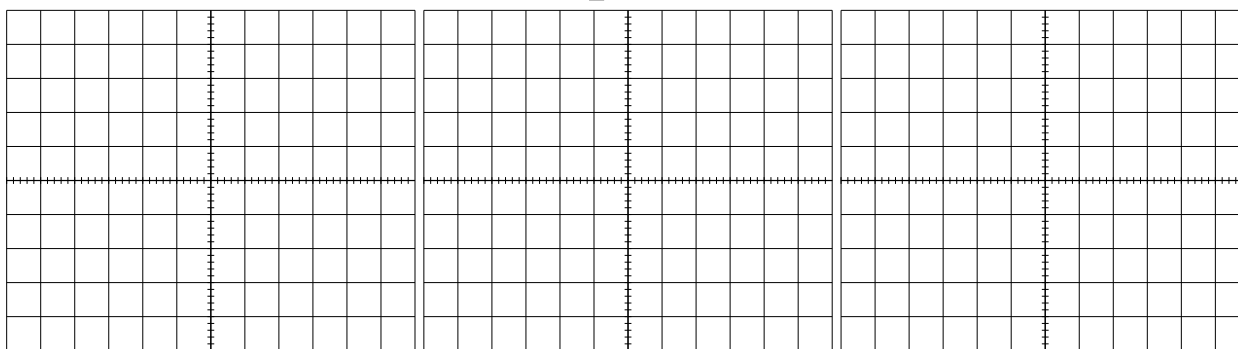
Ein Oszilloskop ist ein Gerät, das elektrische Schwingungen (lateinisch: *oscillare*=schwingen) sichtbar (griechisch: *skopein*=sehen) machen kann.

Dazu wird immer wieder ein Leuchtpunkt von links nach rechts mit in dieser Richtung stets gleicher Geschwindigkeit über einen Bildschirm geführt. Gleichzeitig lässt sich der Punkt mit der interessierenden elektrischen Spannung in vertikaler Richtung ablenken: je nach Vorzeichen der Spannung nach oben oder unten, je nach Betrag der Spannung mehr oder weniger weit. Das Vorzeichen einer Spannung entspricht dabei einfach der Richtung des Stromes, den diese Spannung verursachen würde.

Die Bahn des Leuchtpunktes zeigt also den zeitlichen Verlauf der angelegten Spannung als Schaubild in

einem Koordinatensystem mit nach rechts verlaufender Zeit-Achse und dazu senkrechter Spannungs-Achse. Eine geschickte Elektronik startet jeden von-links-nach-rechts-Lauf des Leuchtfleckes erst dann, wenn die angelegte Spannung ein bestimmtes Verhalten zeigt (z.B. gerade einen Vorzeichenwechsel von $-$ nach $+$ macht). Bei Spannungsverläufen, die mit simpler Regelmäßigkeit ständig wiederkehren, durchläuft der Leuchtpunkt deshalb bei jedem Lauf dieselbe Bahn. Tut er dies schnell genug, sehen wir anstelle des schnellen Pünktchens diese Bahn als durchgehende Linie.

Wir übersetzen nun den von verschiedenen Quellen ausgesandten Schall mit einem Mikrofon in elektrische Spannungs-Schwankungen und führen diese einem Oszilloskop zu.

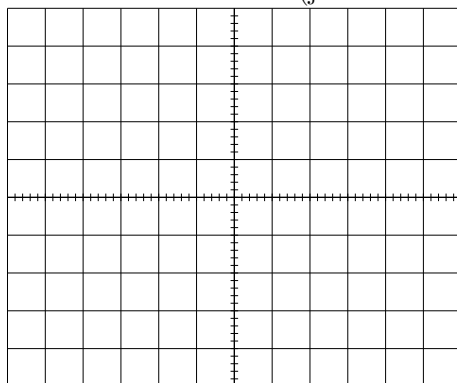


Quelle: _____ Quelle: _____ Quelle: _____

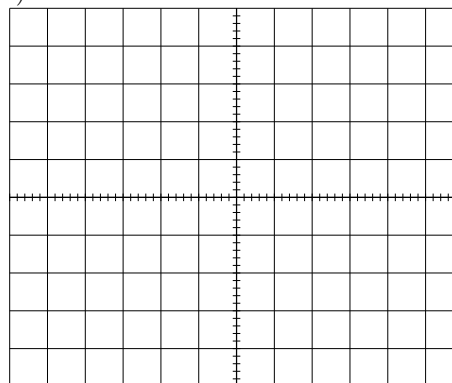
Schall-Art: _____ Schall-Art: _____ Schall-Art: _____

Merkmale: _____ Merkmale: _____ Merkmale: _____

Unterschiedliche Töne: (je zwei in einem Bild)



laut / leise

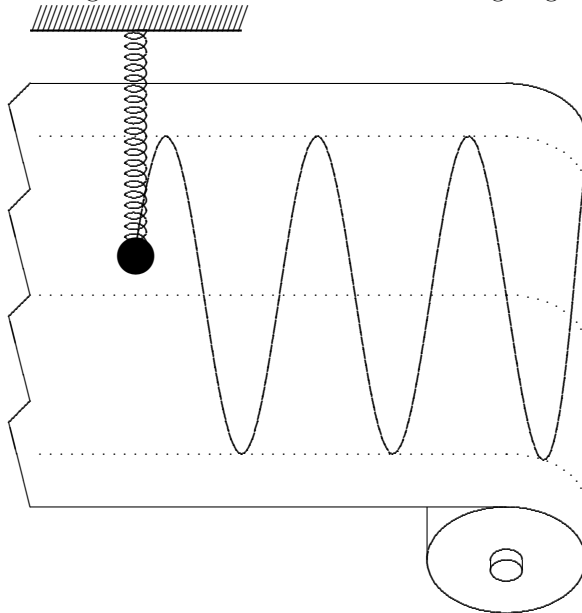


hoch / tief

13.3 Beschreibende Größen

13.3.1 Kenngrößen einer Schwingung

Um Schwingungen zu beschreiben, benutzt man einige Größen, die man an folgendem mechanischen Modell erkennen kann. Ein „Federpendel“ aus einer Schraubenfeder und einem Gewichtsstück schwingt auf und ab. Dabei ist am Gewichtsstück ein Schreiber angebracht, der eine Spur der Schwingung auf eine vorbeilaufende Papierbahn schreibt. Hier erhält man ein Schaubild, das mit horizontaler Zeit-Achse die Auslenkung des Pendels in vertikaler Richtung zeigt.



zu ergänzen:

- Bewegungs- richtung Pendel
- Bewegungs- richtung Papier
- Drehsinn Walze
- Amplitude
- (Schwingungs- dauer)

Amplitude = maximale Auslenkung aus der Mittellage

Schwingungs-
dauer T = Dauer einer vollständigen Einzel-Schwingung, d.h. Dauer des kürzestmöglichen Vorganges, aus dessen Wiederholungen die gesamte Schwingung sich zusammensetzt

Frequenz f = Anzahl der Einzel-Schwingungen pro Zeit;
Einheit: $[f] = \frac{1}{s} =: 1 \text{ Hz} = 1 \text{ Hertz}$

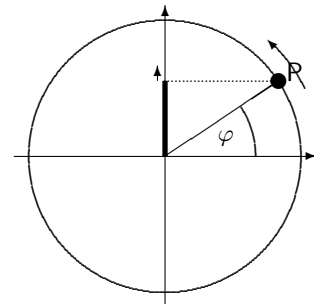
$$f =$$

zu ergänzen:

- Formelterm

Beachte, dass mit „Auslenkung“ und „Amplitude“ nicht immer eine Länge gemeint sein muss. Je nach Art der Schwingung können ganz unterschiedliche Größen im Laufe der Zeit regelmäßig schwanken, also schwingen. Beim Federpendel könnte man statt der Entfernung des Gewichtsstückes von der Mittellage auch z.B. die Spannkraft der Feder als schwingende Größe betrachten. Beim Schall betrachten wir hingegen vor allem den Druck, der schwankt. Ein Oszilloskop als letztes Beispiel zeigt die Schwankung einer elektrischen Spannung.

Auf die Formel $f = \frac{1}{T}$ kommt man sofort, wenn man bedenkt, dass die Anzahl der Schwingungen in einer Schwingungsdauer T ja gerade 1 ist.



Die oben dargestellte Spur des Federpendels ist eine sogenannte **Sinus-Kurve**. Eine solche erhält man auch, wenn man die y-Koordinate eines Punktes P aufzeichnet, der mit gleichbleibender Geschwindigkeit auf einem Kreis um den Ursprung eines Koordinatensystems läuft. Jedem Moment einer Schwingung entspricht daher eine bestimmte Position von P auf seiner Bahn. Da diese Position durch einen Winkel φ wie gezeichnet bestimmt ist, entspricht jedem Moment einer Schwingung auch ein Winkel φ . Diesen Winkel nennt man die **Phase** der Schwingung.

Ergänze, was sich entspricht:

Amplitude der Schwingung	des Kreises
Schwingungsdauer	
Federpendel ganz oben	Phase $\varphi =$
... ganz unten	Phase $\varphi =$
... in der Mitte, von oben kommend	Phase $\varphi =$

13.3.2 Allgemeine Wellengleichung

Eine Welle entsteht, wenn sich eine Schwingung von einem schwingfähigen Teil mit gewisser Verzögerung auf benachbarte, gleichartige Teile überträgt. Und genauso ist ja auf Seite 112 Schall dargestellt: Eine anfängliche Schwingung überträgt sich von Molekül zu Molekül.

Man erkennt in diesem Bild, wie jedes Teilchen im Laufe der Zeit (dargestellt in der Folge der Bilder von oben nach unten) hin- und herschwingt. Verfolgt man in der Abfolge der Bilder z.B. das erste mit Pik (♠) dargestellte Luftmolekül, so erkennt man deutlich, wie dieses Pik ab der dritten Zeile hin- und herschwingt.

Das Sternchen z.B. beginnt etwas später ebenso zu schwingen, nämlich ab der sechsten Zeile. Seine Schwingung läuft dann genauso ab wie die Schwingung des Pik, nur eben zeitlich versetzt. Jede Schwingungsphase (z.B. das Erreichen der größten rechtsseitigen Auslenkung) des Pik findet beim Sternchen 3 Bilder später statt. Beim $\diamond$ findet sie 4 Bilder später statt, beim $\times$ 5, usf.

Man kann dies so sehen, dass jede Phase sich ent-

lang der Molekülreihe mit bestimmter Geschwindigkeit nach rechts bewegt.

Weitere Sichtweise: Es bewegt sich ein Muster von Auslenkungen nach rechts.

Gängigere Formulierung dafür: Es bewegt sich eine Welle mit bestimmter Geschwindigkeit nach rechts.

Das zweite als Pik dargestellte Molekül schwingt im gleichen Takt wie das erste Pik. Es beginnt seine Schwingung zwar erst 10 Zeilen später als das erste Pik, somit aber genau in dem Moment, in dem das erste Pik seine zweite Einzelschwingung beginnt. Die Piks schwingen ab dann *synchron* (griech.: *syn*=zusammen, *chronos*=Zeit). Physiker sagen dazu auch „gleichphasig“ oder „in Phase“.

Allgemeiner schwingen in dem Bild auf Seite 112 jeweils alle diejenigen Moleküle in Phase, die in Abständen von 10 Molekülen auseinander liegen. (Weswegen diese auch jeweils mit demselben Symbol dargestellt sind.)

In einer Welle heißt der kleinste Abstand zweier gleichphasiger Schwinger die **Wellenlänge** der Welle. Die Wellenlänge wird mit λ (gelesen: „lambda“) bezeichnet.

?

Woraus ergibt sich die Wellenlänge?

Vom ersten Pik aus betrachtet ist die Wellenlänge die Entfernung des nächsten Schwingers, der mit dem Pik gleichphasig schwingt. Welcher ist dies?

Der gesuchte Schwinger ist natürlich derjenige, der seine erste Einzelschwingung in dem Moment beginnt, in dem das erste Pik seine erste Einzelschwingung abgeschlossen hat und die zweite beginnt. Der gesuchte Schwinger muss also gerade nach der ersten Einzelschwingung des ersten Pik von der Welle erfasst werden. Da jede Einzelschwingung die Schwingungsdauer T lang dauert, ist der gesuchte Schwinger so weit vom ersten Pik entfernt, wie die Welle mit ihrer Geschwindigkeit v in dieser Zeit vorwärtskommt: $v \cdot T$.

Eine Welle mit Schwingungsdauer T hat bei einer Ausbreitungsgeschwindigkeit v die Wellenlänge

$$\lambda =$$

Üblicher ist es, diesen Zusammenhang mit Hilfe der Frequenz $f = \frac{1}{T}$ der Welle auszudrücken:

Die allgemeine Wellengleichung:

$$v =$$

zu ergänzen:

- Formelterm

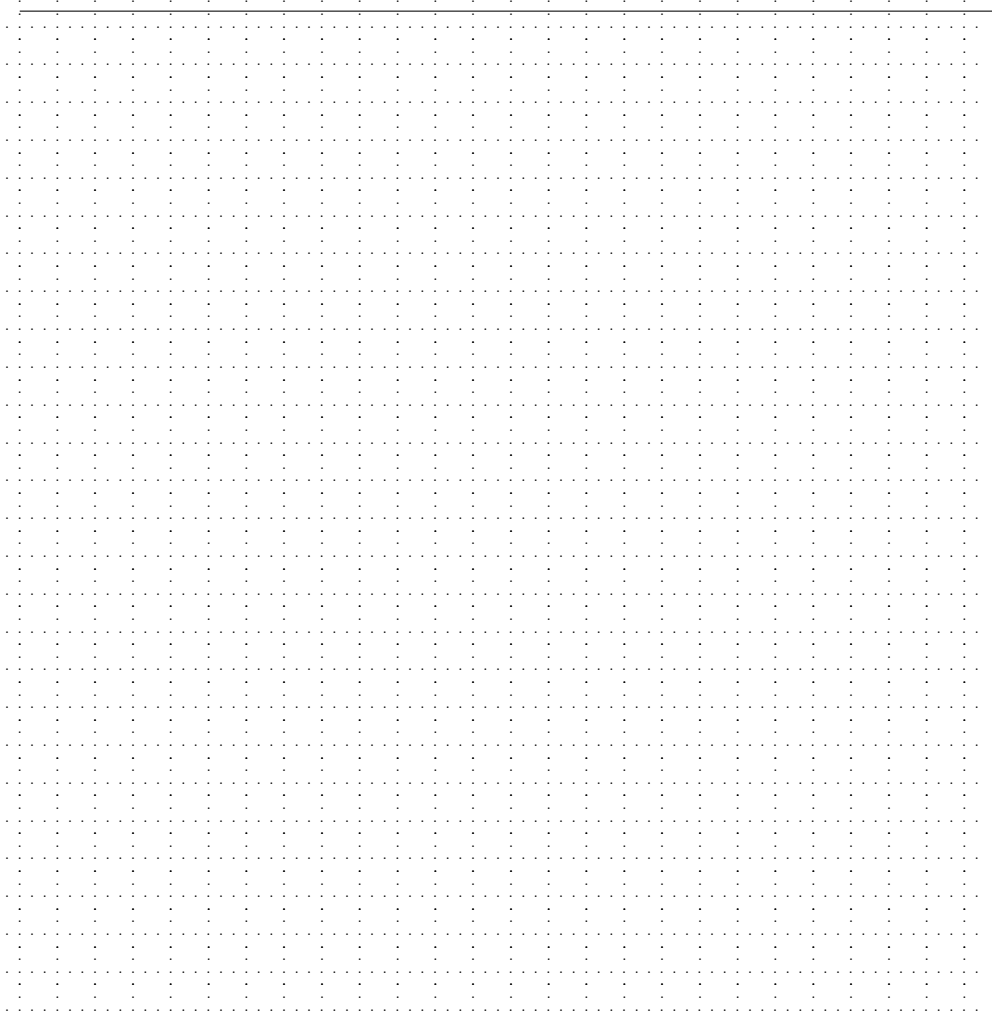
Merkhilfe: Geschwindigkeit ist zurückgelegter Weg pro Zeit. Die Welle legt während jeder Schwingung eine Wellenlänge zurück. So viele Schwingungen, wie die Welle pro Sekunde macht ($\rightarrow f$), so viele Wellenlängen legt sie dann auch zurück.

zu ergänzen:

- Formelterm

13.3.3 Schallgeschwindigkeit

V 13.3 Schallgeschwindigkeit in Luft



In anderen Materialien (hier auch „Medien“ genannt) breitet sich Schall mit anderen Geschwindigkeiten aus.

Medium	$v[\frac{m}{s}]$
Luft	331
Wasser	1497
Meerwasser	1531
Glycerin	1904
Plexiglas	2680
Quecksilber	1450
Blei	2160
Gold	3240
Silber	3650
Kupfer	5010
Eisen	5950
Aluminium	6420

13.4 Musikalische Aspekte

13.4.1 Klänge

Wie wir bereits mit Hilfe eines Oszilloskopes erkennen konnten, haben höhere Töne höhere Frequenzen als tiefere. Wir wissen auch: Laute Töne haben größere Amplituden als leise. Weitere Kenngrößen hat ein einzelner (Sinus-)Ton nicht. In diesem Abschnitt wenden wir uns nun dem Zusammenklang verschiedener Töne zu. Werden gleichzeitig mehrere Töne ausgesandt, erhält man einen sogenannten **Klang**.

Es addieren sich bei einer solchen Überlagerung von Tönen die Schallwellen der einzelnen Töne zu einer Gesamt-Welle: An jeder Stelle des Raumes addieren sich zu jedem Zeitpunkt alle Auslenkungen, die die Summanden-Wellen dort und dann einzeln hervorgerufen würden, zur Gesamtauslenkung.

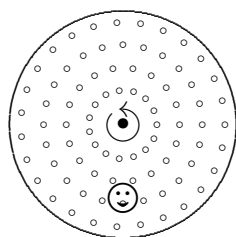
Klänge können sich „gut“ oder „schlecht“ anhören. Die guten nennt man **harmonisch**, die schlechten **dis-harmonisch**.

?

Wann ist ein Klang harmonisch?

Ein einfach gebautes Instrument für eine solche Untersuchung ist eine gespannte **Saite** zum Zupfen. Bei gleicher Länge wird ihr Ton höher, wenn man sie straffer spannt. Bei gleicher Spannung wird ihr Ton höher, wenn die Saite verkürzt wird. Schon in der Antike hat man erkannt, dass man miteinander harmonisierende Töne erhält, wenn man (bei gleicher Spannung) Saiten verwendet, deren Längen sich wie kleine natürliche Zahlen verhalten. Später erkannte man, dass sich dann auch die Ton-Frequenzen wie kleine natürliche Zahlen verhalten.

13.4.2 Diatonische Tonleiter



Dies ist auch im Aufbau einer **Lochsirene** umgesetzt: Eine Scheibe trägt auf konzentrischen Kreisen bestimmte Anzahlen von Löchern. Zum Spielen dieses Instrumentes lässt man die Scheibe schnell rotieren und bläst z.B. mit einem Strohhalm, gezielt gegen einzelne Kreise. (z.B. dorthin, wo im Bild 😊 eingezeichnet ist.)

Bei jeder Umdrehung der Scheibe bläst man dann nacheinander durch alle Löcher des angezielten Kreises. Hinter der Scheibe entsteht jedesmal, wenn man durch ein Loch bläst, ein kleiner Luftstoß hinter diesem Loch. Die rasche Folge von Luftstößen verursacht eine Schallwelle.

Die Frequenz dieser Schallwelle ist natürlich die Frequenz, mit der Löcher vor dem Strohhalm vorbeiziehen. Beispielsweise bei 24 Löchern auf dem Kreis und 10 Umdrehungen der Scheibe pro Sekunde sind dies 240 Hertz. Mit den anderen Kreisen, die andere Lochanzahlen haben, kann man bei gleicher Drehzahl andere Töne spielen.

Durch Vorgabe von Lochanzahlen kann man daher bei fester Drehzahl die gewünschten Frequenzverhältnisse vorgeben. Für das Verhältnis der Töne aus zwei Kreisen 1 und 2 spielt die Drehzahl selbst keine Rolle, solange sie nur gleich bleibt:

$$\frac{\text{Frequenz 1}}{\text{Frequenz 2}} = \frac{\text{Lochanzahl 1} \times \text{Drehzahl}}{\text{Lochanzahl 2} \times \text{Drehzahl}} = \frac{\text{Lochanzahl 1}}{\text{Lochanzahl 2}}$$

Man erhält die natürliche, diatonische C-dur-Tonleiter mit einer Lochsirene, die folgende Lochanzahlen und eine Drehzahl von 11 Umdrehungen pro Sekunde hat:

Ton:	c	d	e	f	g	a	h	c'
Lochanzahl:	24	27	30	32	36	40	45	48
Frequenz[Hz]:								
Multiplikator								

Frequenz=
Lochanzahl × Drehzahl

Da viele dieser Lochanzahlen recht viele gemeinsame Teiler haben, kommen in dieser Tonleiter viele einfache Frequenzverhältnisse vor. In der Musik bezeichnet man Abstände von Tönen als **Intervalle**. Die Namen ergeben sich aus lateinischen Zahlwörtern für die Anzahl der Töne, die das Paar umfasst:

Zwei gleiche Töne umfassen nur diesen einen Ton, bilden eine *Prim*. Ton und folgender Ton, z.B. f und g, umfassen zwei Töne, bilden eine *Sekunde*. Ton und übernächster Ton, z.B. c und e, umfassen drei Töne, liegen eine *Terz* auseinander. Die weiteren Namen sind: *Quart*, *Quint*, *Sext*, *Septime* und *Oktave*.

Wie an den Lochanzahlen für c und c' bereits zu erkennen, verdoppelt sich die Frequenz immer, wenn man eine Oktave höher geht. Der Einfachheit des Frequenzverhältnisses 1:2 entsprechend hören sich Töne, die eine Oktave auseinanderliegen, so harmonisch an, dass man sie sogar gleich bezeichnet.

Bei den übrigen Intervallen sind zunächst nicht so klare Aussagen möglich. Die Quint aus c und g hat das sehr harmonische Frequenzverhältnis $24:36=2:3$, die Quint zwischen d und a jedoch $27:40=0,675$, was zwar nahe an $2:3=0,6$ liegt, sich jedoch gar nicht zu einem Verhältnis kleiner ganzen Zahlen kürzen lässt.

Noch schlimmer fällt der Vergleich aus zwischen der Sekunde c-d mit $24:27=8:9=0,8$ und der Sekunde e-f mit $30:32=15:16=0,9375$.

Diese Unterschiede machen Schwierigkeiten, sobald

ein Musikstück für bestimmte Instrumente in deren höheren oder tieferen Tonbereich übertragen („transponiert“) werden soll. Es können dann nicht einfach alle Töne z.B. um eine Terz erhöht werden, denn es ergeben sich ja dann andere Frequenzverhältnisse innerhalb des Stückes, und damit verzerrte Harmonien, also Missklänge.

Eine Möglichkeit zur Linderung dieser Schwierigkeiten ergibt sich aus den Multiplikatoren zwischen den Frequenzen benachbarter Töne: Die beiden Schritte von e nach f und von h nach c' sind deutlich kleiner als die übrigen fünf Schritte und machen daher die meisten Schwierigkeiten.

Man führte deshalb in den fünf anderen Notenschritten Zwischentöne in die Tonleiter ein, macht dort statt einem „Ganztonschritt“ zwei „Halbtonschritte“. Je nachdem, ob man die Zwischentöne vom tieferen Ton aus nach oben erhöht, oder vom höheren Ton aus nach unten erniedrigt, erhält man die neuen Töne cis, dis, fis, gis und ais, beziehungsweise des, es, ges, as und b.

Je nach Vorgehensweise bei der Erzeugung der neuen Töne mussten beim Erhöhen und Erniedrigen allerdings nicht dieselben Töne herauskommen. Jedenfalls erhält man eine Tonfolge, die in 12 Halbtonschritten (c - cis/des - d - dis/es - e - f - fis/ges - g - gis/as - a - ais/b - h - c) von c nach c' führt, wobei die Halbtonschritte einander schon mehr gleichen als die 7 Multiplikatoren zwischen den ursprünglichen 8 Tönen.

13.4.3 Chromatische Tonleiter

Optimal zu transponieren wäre Musik dann, wenn alle Halbton-Schritte exakt gleich wären, wenn es also nur einen gleichbleibenden Frequenzmultiplikator für jeden Halbton-Schritt gäbe. Um die Frequenzverdoppelung in einer Oktave beizubehalten, müsste dieser einheitliche Multiplikator zwölf mal angewandt insgesamt einen Faktor 2 ergeben. Es ist also die zwölfte Wurzel aus 2, $\sqrt[12]{2}$, näherungsweise = 1.0594631.

Verändert man die Töne der Tonleiter gemäß dieser Überlegung, so erhält man eine temperierte, chromatische Tonleiter: Man hat willkürlich die Frequenz des Tones a beibehalten, jedoch alle anderen Frequenzen

so undefiniert, dass jeder Halbtonschritt nach unten die Frequenz durch $\sqrt[12]{2}$ teilt und jeder Halbtonschritt nach oben sie mit dieser Wurzel multipliziert.

Ein (kleiner) Nachteil dieser Tonleiter liegt nur darin, dass (außer bei der Oktave) überhaupt keine rationalen Frequenzverhältnisse mehr vorkommen. Die $\sqrt[12]{2}$ und alle ihre Potenzen (bis auf jede 12.) sind ja irrationale Zahlen. Die Frequenzverhältnisse liegen jedoch nahe genug bei den Verhältnissen kleiner ganzer Zahlen, dass unsere Ohren sie immer noch als harmonisch empfinden.

13.5 Eigenschwingungen

13.5.1 Eigenfrequenz

Sonnenklar: Bestimmte Vorgänge dauern bestimmte Zeiten. Bisherige Beispiele:

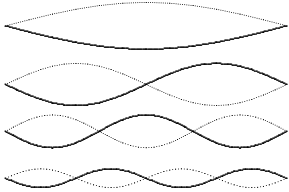
- Das Pendel auf Seite 115 benötigt für ein Auf-und-Ab eine bestimmte Zeit, die sich aus seinem Aufbau (Federhärte und anhängende Masse) ergibt.
- Die Zinke einer Stimmgabel benötigt für ein Hin-

und-Her eine bestimmte Zeit, die wohl von der Form, der Dicke und dem Material der Stimmgabel abhängt.

- Eine gespannte Saite benötigt für eine Schwingung eine bestimmte Zeit, die von der Spannung und der Länge der Saite abhängt.

In allen diesen Fällen überlässt man einen Körper oder ein System nach einer anfänglichen Anregung (Schubs, Anschlag, Anzupfen) sich selbst und beobachtet Schwingungen bestimmter Frequenzen. Nur so kann man überhaupt mit Musikinstrumenten gezielt Töne bestimmter Frequenzen erzeugen.

Eine Frequenz, mit der ein angeregtes System ohne weitere äußere Einwirkung schwingen kann, heißt **Eigenfrequenz** des Systems.



13.5.2 Oberschwingungen

Bei genauerer Untersuchung von Eigenfrequenzen stellt man fest, dass ein System durchaus viele Eigenfrequenzen haben kann. Eine Saite beispielsweise kann nicht nur wie im ersten der nebenstehenden Bilder dargestellt schwingen, sondern auch wie in den anderen Bildern gezeigt. Es ist jeweils dargestellt: Saite bei einer maximalen Auslenkung (durchgezogen) und bei der zeitlich um $T/2$ versetzt vorliegenden entgegengesetzt maximalen Auslenkung (punktirt).

Ab dem zweiten Bild ist zu erkennen, wie die Saite in mehrere Abschnitte aufgeteilt ist, die jeweils wie eine entsprechend kurze eingespannte Saite schwingen. Im zweiten Bild schwingen quasi zwei halbe Saiten im Gegentakt. Der halben Länge entsprechend ist die Frequenz dieser Schwingung doppelt so groß wie im ersten Bild. Im dritten Bild ist die Frequenz dreimal so groß wie im ersten.

Wenn Du ein langes Seil mit einem Ende irgendwo festbindest und das andere mit der Hand hin- und herbewegst, kannst Du die hier dargestellten Schwingungsformen „live“ sehen. Du musst das Seilende in Deiner Hand allerdings mit der jeweils richtigen Frequenz bewegen. Mit etwas Gefühl spürst Du jedoch die richtigen Frequenzen, dann geht es fast wie von selbst. Man nennt ein solches Bewegungsmuster eine **stehende Welle**. Die Stellen, an denen die Saite/das Seil sich nicht bewegt, heißen **Knoten**, die schwingenden Abschnitte dazwischen **Bäuche**.

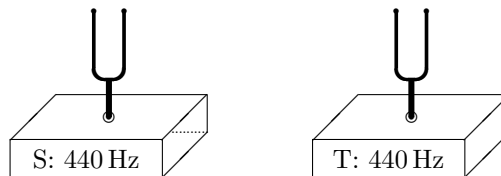
Die einfachste Schwingungsform mit einem einzi-

gen Bauch nennt man **Grundschwingung**, die übrigen **Oberschwingungen**. Die Frequenzen der Oberschwingungen liegen den kürzeren Bauch-Längen entsprechend nämlich *oberhalb* der Grundfrequenz: Es sind die ganzzahligen Vielfachen der Grundfrequenz. Musikalischer ausgedrückt: Grund- und erste Oberschwingung liegen eine Oktave auseinander. Die zweite Oberschwingung (2-fache Grundfrequenz) bildet mit der dritten (3-fache Grundfrequenz) eine Quint. Und so weiter.

Tatsächlich setzt sich die Schwingung einer Saite immer aus allen Eigenschwingungen, also aus der Grund- und den Oberschwingungen zusammen. Eine Saite erzeugt daher – wie jedes Instrument – eigentlich nicht nur einen Ton, sondern immer einen Klang. Die Mischungsverhältnisse zwischen den Eigenschwingungen bestimmen dabei die „Klangfarbe“ des Instrumentes. Sind dem Grundton vor allem niedrige Obertöne beige-mischt, ist der Klang jedenfalls harmonisch.

13.5.3 Resonanz

V 13.4 Resonanz bei Stimmgabeln



Aufbau:

Durchführung: Die linke Stimmgabel S wird angeschlagen und nach einigen Sekunden wieder (durch Anfassen mit der Hand) gedämpft.

Beobachtung:

Durchführung: Eine der Stimmgabeln wird durch Anbringen von Gewichts-Stückchen verstimmt / durch eine Stimmgabel anderer Frequenz ersetzt. Der Versuch wird wiederholt.

Beobachtung:

Die linke Stimmgabel S erzeugt eine Schallwelle der Frequenz f_S , die natürlich auch auf die rechte Stimmgabel T trifft. Auf Stimmgabel T wirken also Druckschwankungen ein, die T's Zinken mit der Frequenz f_S hin und her zu biegen versuchen. Zur Vereinfachung sehen wir nun die Schallwelle als eine Folge von Druckstößen an, von denen jeder die Stimmgabel T einmal „anschlägt“, wenn er bei T eintrifft.

Der erste Stoß alleine würde T zu einer schwachen Eigenschwingung mit der Frequenz f_T anregen. T's getroffene Zinke würde sich zunächst in Stoßrichtung nach rechts biegen, umkehren, nach links zurückfedern und sich nach links biegen, wieder nach rechts umkehren usw.

Trifft der zweite Stoß nun gerade in dem Moment ein, in dem T sich bei seiner Eigenschwingung sowieso wieder in Stoßrichtung bewegen würde, wird diese Bewegung verstärkt. T's Eigenschwingung wird also hef-

tiger, und dies auch wieder bei jedem weiteren Stoß. Damit die Stöße so in den richtigen Momenten ankommen, muss natürlich die Frequenz der Stöße gleich der Eigenfrequenz von T sein: $f_S = f_T$.

Trifft der zweite Stoß in einem anderen Moment ein, weil $f_S \neq f_T$, so wird die Eigenschwingung von T zumindest nicht optimal verstärkt, wenn nicht sogar gebremst. Und selbst, wenn der zweite Stoß noch einigermaßen passt, ist aber der dritte Stoß noch etwas weiter „daneben“ und der vierte noch mehr. Auf Dauer stellt sich daher gar keine Eigenschwingung von T ein.

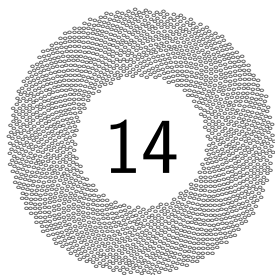
Genauere Untersuchungen und Rechnungen zeigen, dass T auf Dauer lediglich eine relativ schwache Schwingung mit der Frequenz der ankommenden Schallwelle ausführen kann. Deren Amplitude wird dabei kleiner, wenn f_S von f_T mehr abweicht.

Jedes schwingungsfähige System kann leicht zu relativ heftigen Schwingungen angeregt werden, wenn man es mit einer seiner Eigenfrequenzen anregt. Ein so hervorgerufenes heftiges Mitschwingen in einer Eigenfrequenz nennt man **Resonanz**.

Sofern das in Resonanz schwingende System nicht z.B. durch Reibungskräfte ausreichend gebremst wird, kann die Schwingung so heftig werden, dass das System zerstört oder beschädigt wird. Dies versteht man unter einer **Resonanzkatastrophe**. Resonanzkatastrophen zu verhindern ist eine wichtige Aufgabe bei der technischen Entwicklung von Gegenständen aller Art.

Resonanz ist auch im Spiel, wenn wir statt loser Stimmgabeln lieber die auf passende Holzkästen montierten benutzen. Der Holzkasten hat zwei Effekte:

1. Insbesondere die obere Holzfläche, die mit der Stimmgabel verbunden ist, schwingt mit der Gabel mit. Es schwingen nämlich nicht nur die Zinken der Gabel hin und her, sondern auch der Stiel auf und ab. Dadurch wird das Holz mit auf- und abwärts bewegt. Wegen der größeren Fläche des Holzbrettchens wird durch deren Schwingung natürlich eine intensivere (weil breiter angeregte) Schallwelle abgestrahlt als durch die schmale Gabel alleine. Dies alleine ist *kein* Resonanz-Effekt.
2. Resonanz kommt nun ins Spiel: Das Kästchen ist in seinen Abmessungen so konstruiert, dass die darin enthaltenen Luft die Frequenz der Stimmgabel als Eigenfrequenz hat. So wird das Mitschwingen des Kästchens besonders intensiv.



14.1 Temperaturempfindung

Wenn wir einen Körper berühren, so haben wir eine Empfindung darüber, wie warm der Körper ist. Wir unterscheiden in unserer Sprache mehrere Abstufungen von „Warmheit“, etwa: *heiß, warm, lau, kühl, kalt*. Vor allem zu *kalt* und *heiß* sind Steigerungen möglich, z.B. durch Voranstellen geeigneter Bezeichnungen von Tieren.

Den so ausdrückbaren Aspekt des Zustandes eines Körpers nennt man die **Temperatur** des Körpers.

Für physikalische Untersuchungen ist diese verbale

Beschreibung der Temperatur jedoch unzulänglich: Sie ist viel zu grob und mangels Zahlen nicht ohne Weiteres mit Mitteln der Mathematik zu verarbeiten. Ein Maß ohne Zahlenwerte nennt man **qualitativ** (von lat. *qualis*=wie beschaffen) im Gegensatz zu einem **quantitativen** (lat. *quantum*=wie viel) Maß in Zahlen.

Wir können auch nicht unser persönliches Temperaturempfinden zur Bestimmung von Temperaturen einsetzen. Unsere Temperaturempfindung ist nämlich sehr *subjektiv* und schlecht *reproduzierbar*.

?

Was heißt subjektiv?

Laut einem Satz wie *Fritzchen prüft die Wassertemperatur*. tut das Subjekt *Fritzchen* etwas mit dem Objekt *die Wassertemperatur*. Im einfachsten Fall steckt Fritzchen dazu seinen Finger ins Wasser.

Im Idealfall sollte das Ergebnis seiner Prüfung nur vom Objekt, nämlich der Wassertemperatur, abhängen. Jedes andere Subjekt, z.B. Linda, würde am gleichen Objekt mit der gleichen Methode zum gleichen Ergeb-

nis kommen. Eine solche Prüfung würde man dann **objektiv** nennen.

Tatsächlich empfindet Fritzchen aber vielleicht nur als *lau*, was für Linda bereits *warm* ist. Das Ergebnis der Temperaturprüfung hinge dann nicht nur vom geprüften Wasser, sondern auch vom Subjekt ab, das die Prüfung durchführt. Eine solche Prüfung heißt **subjektiv**.

?

Was heißt reproduzierbar?

Fritzchens Wasserprüfung wäre reproduzierbar, wenn Fritzchen bei jeder Wiederholung seiner Prüfung zum gleichen Ergebnis käme. (Vorausgesetzt, das Wasser behält objektiv die gleiche Temperatur.)

Tatsächlich wird sich aber vielleicht Fritzchens nasser Finger nach der ersten Prüfung an winterlicher Luft

stark abkühlen, so dass er das Wasser bei der zweiten Prüfung als *warm* empfinden wird. Vielleicht prüft er auch als nächstes eine heiße Herdplatte und verbrennt sich ein wenig daran. Unmittelbar danach empfindet er deshalb das Wasser bei der dritten Prüfung als angenehm *kühl*.

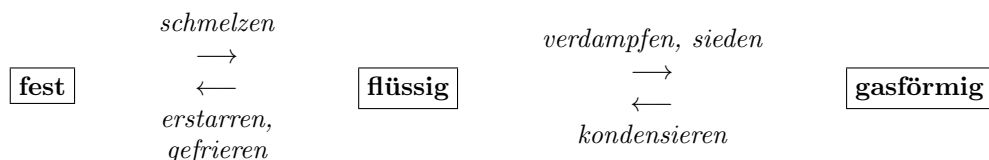
Die gleiche Temperatur empfinden also verschiedene Personen als verschieden warm, und sogar eine einzige Person kann sie unterschiedlich empfinden, je nachdem, was die Person vorher erlebt hat. Unser nächstes Ziel muss daher jedenfalls sein, Methoden zu finden, eine Temperatur quantitativ, objektiv und reproduzierbar zu messen. Anders geht Physik nicht! Wir werden dazu erst einmal klären, welche messbaren Veränderungen überhaupt an einem Körper bei wechselnder Temperatur stattfinden.

14.2 Einfluss der Temperatur

14.2.1 Aggregatzustände

Ein Körper kann fest, flüssig oder gasförmig sein. Diese drei Zustände sind mit dem Begriff *Aggregatzustände* gemeint. Wasser kennen wir in allen drei Zuständen. Meist erleben wir Wasser als Flüssigkeit. Ein Eiswürfel ist Wasser in fester Form. Beim Kochen entweicht aus dem Topf gasförmiges Wasser. Wasser ist in gasförmigem Zustand unsichtbar wie Luft. Der Wasserdampf,

den man über einem Kochtopf sieht, enthält winzige sichtbare Wassertröpfchen. Das gasförmige Wasser, das aus dem kochenden Wasser aufsteigt, kühlt nämlich oberhalb des Topfes an der Luft wieder ab und wird wieder flüssig. Für die Übergänge zwischen den Aggregatzuständen benutzt man folgende Verben:



Für den seltener zu beobachtenden direkten Übergang von fest zu gasförmig benutzt man das Verb *sublimieren*, für den umgekehrten Vorgang *resublimieren*. Ein Beispiel für Sublimation: Im Winter draußen aufgehängene nasse Wäsche gefriert zwar, trocknet aber dennoch langsam. Offensichtlich geht dazu das gefrorene Wasser direkt in den gasförmigen Zustand über.

„Verdampfen“ ist der Oberbegriff für jede Umwandlung von flüssig nach gasförmig. Mit „sieden“ meint man, dass in der Flüssigkeit Dampfblasen entstehen. Beachte, dass ein Verdampfen an der Oberfläche der Flüssigkeit auch ohne Blasen stattfindet. Bei einem sehr langsamen solchen Verdampfen an Luft spricht man auch von „Verdunstung“.

Tatsächlich finden an der Oberfläche einer Flüssigkeit ständig Verdampfung und Kondensation statt, je nach Umständen (Chemie, Druck, Temperatur) mehr oder weniger, so dass das eine oder andere überwiegt.

Solche Übergänge zwischen den Aggregatzuständen können durch eine Änderung der Temperatur ausgelöst werden. Kühlt man flüssiges Wasser weit genug ab, so erstarrt es zu Eis. Erwärmt man Eis genügend, so schmilzt es. Erhitzt man flüssiges Wasser ausreichend, so siedet es. Kühlt Wasserdampf wieder ab, so kondensiert er zu flüssigem Wasser. Aus unserer Erfahrung, insbesondere mit winterlicher Witterung, wissen wir auch, dass Wasser unter gewöhnlichen Bedingungen immer etwa bei der gleichen Temperatur gefriert und wieder schmilzt: Soweit wir unserem subjektiven Empfinden trauen können, fühlt sich z.B. Schneematsch immer gleich kalt an.

Wenn wir sagen: „Wasser gefriert bei 0 °C“, so heißt dies nicht, dass Wasser beim Abkühlen zu gefrieren beginnt, sobald es diese Temperatur *erreicht*. Es ist vielmehr so, dass es beim *Unterschreiten* dieser Temperatur gefriert. Wasser gefriert nicht, wenn nicht in seiner Umgebung schon Temperaturen unter 0 °C herr-

schen. Entsprechendes gilt natürlich für das Schmelzen, Sieden und Kondensieren.

Direkte Erfahrungen mit kochendem Wasser vermeiden wir. Im Bereich so hoher Temperaturen könnten wir auch gar nicht so gut Temperaturen unterscheiden. Wir gehen aber auch hier davon aus, dass Wasser unter gewöhnlichen Bedingungen immer bei der gleichen Temperatur siedet oder kondensiert.

Viele andere Stoffe können ebenfalls in allen drei Aggregatzuständen vorliegen. Zum Beispiel ist Kerzenwachs bei Zimmertemperatur fest und schmilzt in der Nähe von Feuer. Was in einer Kerzenflamme brennt, ist schließlich das im Bereich des Dochtes verdampfte Wachs. Was nach dem Auspusten der Flamme eine kurze Zeit lang weiter aus dem Docht aufsteigt, ist ebenfalls verdampftes und wieder zu feinsten Tröpfchen kondensiertes Wachs. Vielleicht willst Du Dich selbst davon überzeugen, dass dieser Dampf brennt: Zünde eine Kerze an und halte ein brennendes Streichholz bereit! Puste die Kerze aus und halte das brennende Streichholz einige Zentimeter vom Docht entfernt an den vom Docht aufsteigenden Dampf!

Die Temperaturen, bei denen Wachs schmilzt oder siedet, sind andere als die von Wasser. Jeder Stoff hat seine, für ihn charakteristischen Umwandlungstemperaturen. Man nennt diese Temperaturen Schmelz- oder Erstarrungs- bzw. Siedetemperatur.

Bei vielen Stoffen liegen diese Umwandlungstemperaturen so weit außerhalb unseres Erfahrungsbereiches, dass wir persönlich diese Stoffe nicht in allen Aggregatzuständen kennen. Dennoch gibt es sehr wohl z.B. „flüssige Luft“ (eigentlich Stickstoff, der Hauptbestandteil von Luft), festes Kohlendioxid („Trockeneis“), flüssiges Gestein (Lava), flüssiges Eisen (z.B. in Hochöfen) und Gold-Dampf.

14.2.2 Teilchenmodell

Um eine Vorstellung davon zu bekommen, was bei einer Umwandlung des Aggregatzustandes eigentlich passiert, gehen wir von folgender Vorstellung aus:

Jeder Körper besteht aus winzig kleinen Teilchen mit folgenden beiden Eigenschaften:

- Alle diese Teilchen befinden sich ständig in Bewegung. Je heißer ein Körper ist, desto heftiger bewegen sich seine Teilchen.
- Die Teilchen üben andererseits aufeinander (überwiegend) anziehende Kräfte aus. Mit diesen hindern sie sich gegenseitig daran, sich voneinander zu entfernen.

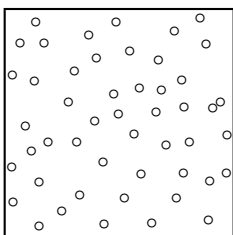
Feste Körper

Ist die Temperatur nun recht niedrig, so bewegen sich die Teilchen recht „zahm“. Aufgrund der anziehenden Kräfte kommt es bei vielen Stoffen zu einer möglichst kompakten regelmäßigen Anordnung aller Teilchen, in der jedes Teilchen von seinen Nachbarteilchen an seinem Platz gehalten wird. Das Teilchen kann nur um diesem Platz ein wenig hin und her schwingen. Der Körper ist in diesem Zustand fest. Man sagt auch: Die Teilchen bilden ein festes (Kristall-)Gitter. Die genaue Form des Gitters hängt von der Art der Teilchen ab.



Flüssigkeiten

Mit steigender Temperatur schwingen die Teilchen immer heftiger um ihre Gitterplätze. Ab bestimmter Temperatur können die Teilchen sogar ihre Plätze verlassen. Die anziehenden Kräfte halten zwar immer noch alle Teilchen aneinander, aber nicht mehr an festen Plätzen. Die Teilchen bewegen sich dicht an dicht, kreuz und quer aneinander vorbei. Die Gitterstruktur des Körpers ist zusammengebrochen. Der Körper ist geschmolzen und nun flüssig. Er passt sich jeder Gefäßform an, hat aber noch ein bestimmtes Volumen.



Gase

Wird die Temperatur noch höher und damit die Teilchenbewegung noch heftiger, so lösen sich die Teilchen aus der gegenseitigen Anziehung. Es fliegen nun einzelne Teilchen ungeordnet durch den ganzen Raum. Sie stoßen gelegentlich aufeinander oder auf die Wände des Raumes und prallen ab. Der Körper ist damit zu einem Gas geworden. Er füllt den ganzen zur Verfügung stehenden Raum aus.

Mit der Temperatur steigt die *mittlere* Geschwindigkeit der Teilchen. Es gibt jedoch immer auch Teilchen, die deutlich schneller oder langsamer sind als der Durchschnitt. Die schnelleren Teilchen können es bereits unterhalb der Siedetemperatur schaffen, die Flüssigkeit zu verlassen, so dass immer zumindest ein wenig Verdampfung stattfindet. Erst bei der Siedetemperatur schaffen es auch die nur durchschnittlich schnellen Teilchen, sich von ihren Mit-Teilchen zu lösen.

Kühlt das Gas wieder ab, bewegen sich die Teilchen nicht mehr heftig genug, um der gegenseitigen Anziehung zu entkommen. Sie klumpen wieder zusammen: Der Körper kondensiert, es bilden sich z.B. immer größere Tröpfchen.

Kühlt die Flüssigkeit weiter ab, stellt sich irgendwann wieder der starre Zusammenhalt in Form eines Gitters ein: Der Körper erstarrt.

Gas + X

Folgende Begriffe darfst Du auch unterscheiden:

Nebel sind Gase, in denen irgendwelche Flüssigkeitströpfchen schweben. Der vom Wetter bekannte Nebel besteht z.B. aus Luft mit winzigen Wassertropfchen darin.

Kurz: Nebel=Gas mit *irgendwelchen* Flüssigkeitströpfchen

Dampf nennt man z.B. das beim Wasserkochen entweichende (an sich unsichtbare) gasförmige Wasser mit den darin befindlichen (sichtbaren) Tröpfchen aus wieder kondensiertem Wasser. Kurz: den Nebel aus Wassergas mit Wassertröpfchen. Im Allgemeinen bezeichnet „Dampf“ jedes Gas, das Kontakt zu der Flüssigkeit aus *derselben* Teilchenart hat. Die Flüssigkeit muss dazu nicht tröpfchenförmig vorliegen. Ein weiteres Beispiel eines Dampfes ist daher auch das Wassergas (ohne bereits wieder entstandene Tröpfchen) unmittelbar oberhalb des noch flüssigen Wassers in einem Kochtopf. Weiteres Beispiel: gasförmiges Kerzenwachs, in dem bereits wieder kondensierte Wachströpfchen schweben
Kurz: Dampf=Gas mit Flüssigkeit(ströpfchen) *derselben* Art

Rauch bezeichnet ein Gas, in dem flüssige Tröpfchen oder feste Körnchen schweben, die aus einer Verbrennung stammen.
Kurz: Rauch=Gas mit festen oder flüssigen Verbrennungsprodukten

14.2.3 Dehnung

In der Regel dehnen sich alle Körper in allen Aggregatzuständen bei Erwärmung aus und ziehen sich bei Abkühlung zusammen.

Die Teilchen eines Körpers bewegen sich heftiger, wenn der Körper erwärmt wird. Man kann sich gut vorstellen, dass sie dann auch mehr Platz beanspruchen.

Ausdehnung von festen Körpern

V 14.1 Kugel durch Ring

Aufbau: Eine Metallkugel und ein Metallring haben solche Ausmaße, dass die Kugel gerade durch den Ring passt. Der Ring wird fest so montiert, dass man die Kugel durch ihn hindurch fallen lassen kann.

Durchführung:

Beobachtung:

Ergebnis:

Durchführung:

Beobachtung:

Ergebnis:

Ausdehnung von Flüssigkeiten

V 14.2 Steigrohr mit Flüssigkeit

Aufbau: Ein Glaskolben ist mit Wasser aus dem Wasserhahn gefüllt und mit einem Stopfen verschlossen. Durch den Stopfen ist ein enges Glasröhrchen geführt. Das Wasser steht bis in das Glasröhrchen hinein.

Durchführung: Der Glaskolben wird über einem Bunsenbrenner erwärmt (und später wieder abgekühlt).

Beobachtung:

Ergebnis:

Anomalie des Wassers

Führt man den vorigen Versuch mit anfangs eiskaltem Wasser durch, so beobachtet man, dass der Wasserspiegel im Steigrohr zunächst sinkt und erst dann steigt. Wasser zieht sich in diesem kalten Temperaturbereich beim Erwärmen zusammen! Es verhält sich damit anders als *normale* Stoffe, die sich ja beim Erwärmen ausdehnen. Man spricht daher von der *Anomalie* des Was-

sers. (kein Druckfehler, kein fehlendes r)

(Im Vertrauen darauf, dass Du die folgende Temperaturangabe schon verstehst, obwohl sie in diesem Skript noch nicht erklärt worden ist:) Diese Anomalie erstreckt sich von 0 °C bis 4 °C. Eine Wassermenge hat demnach bei 4 °C ihr kleinstes Volumen und – gleichbedeutend dazu – ihre größte Dichte.

Ausdehnung von Gasen

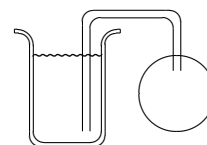
Schon folgender einfache Versuch zeigt die Ausdehnung einer Gasmenge:

V 14.3

zu ergänzen:

- warme Hand
- (Beobachtung)

Aufbau: Aus einem luftgefüllten Glaskolben führt ein Schlauch, dessen anderes Ende man in ein Wasserglas taucht.



Durchführung:

Beobachtung:

V 14.4 Steigrohr mit Gas

Aufbau: Glaskolben mit Steigrohr, jedoch mit Luft gefüllt. Ein Wassertröpfchen verschließt das Steigrohrchen.

Durchführung:

Beobachtung:

Ergebnis:

14.3 Thermometer

14.3.1 Flüssigkeitsthermometer

Thermometer sind Geräte zur Messung von Temperaturen. Der Aufbau im Steigrohr-Versuch mit Wasser gleicht bereits einem Thermometer. Je höher die Temperatur im Glaskolben war, desto höher stieg das Wasser im Röhrchen. Jede Steighöhe gehörte zu genau einer bestimmten Temperatur, die man also an dieser Steighöhe erkennen kann.

Aufgrund seiner Dichte-Anomalie würde Wasser bei Temperaturen um 4°C allerdings jeweils zwei verschiedene Temperaturen mit der gleichen Steighöhe anzeigen. Daher ist Wasser für diesen Temperaturbereich zum Messen ungeeignet. Als Flüssigkeit verwendet man schon von daher lieber einen Alkohol oder Quecksilber. Ein weiterer Grund dafür, nicht einfach Wasser zu verwenden, ist der, dass Wasser schon bei winterlicher Kälte gefrieren würde. Beim Gefrieren dehnt sich Wasser aber aus und würde damit das Thermometer sprengen!

Brauchbare Thermometer sind allerdings viel kleiner als der Glaskolben im vorigen Versuch. Dennoch besteht jedes Flüssigkeitsthermometer auch aus einem verhältnismäßig großen, mit einer Flüssigkeit gefüllten Gefäß mit einem darauf aufgesetztem engen Steigrohrchen.

Um jeder messbaren Temperatur einen Zahlenwert zuzuordnen genügt es, am Steigrohrchen entlang irgendeine Skala anzubringen. Es ist dabei sinnvoll, beim Anbringen einer Skala genauso vorzugehen, wie viele andere Leute es vorher auch schon gemacht haben. Nur so kann gewährleistet sein, dass man auf verschiedenen Thermometern bei gleicher Temperatur auch das gleiche abliest.

In verschiedenen Ländern haben verschiedene Forscher im Laufe der Zeit verschiedene Skalen festgelegt. Die wichtigste dieser willkürlich festgelegten Skalen ist die bei uns benutzte Celsius-Skala.

14.3.2 Die Celsius-Skala – und andere

Große Teile unserer Umwelt und von uns selbst bestehen aus Wasser oder werden von Wasser beeinflusst. Wasser ist zudem fast überall gut zu beschaffen. Die Gefrier- und die Siedetemperatur von Wasser sind technisch recht leicht zu erreichen.

Der Schwede Anders Celsius (1701-1744) hat daher eine Temperaturskala von den Umwandlungstemperaturen des Wassers abgeleitet. Ein Thermometer wird mit folgenden vier Schritten in der Celsius-Skala geeicht:

1. Das Thermometer wird in Eiswasser (Wasser mit Eisstücken darin) getaucht. Eiswasser hat genau die Temperatur schmelzenden Eises / erstarrenden Wassers. Die sich einstellende Steighöhe im Steigrohrchen wird mit 0 beschriftet.
2. Das Thermometer wird in siedendes Wasser getaucht. Siedendes Wasser hat genau die Siedetemperatur von Wasser. Die sich einstellende Steighöhe wird mit 100 beschriftet.

3. Das Stück zwischen der 0- und der 100-Marke wird in 100 gleiche Teile unterteilt. Die Teilstriche werden fortlaufend von 1 bis 99 beschriftet.

4. Unterhalb der 0 und oberhalb der 100 wird die Skala in gleicher Schrittweite fortgeführt mit -1, -2, ... bzw. 101, 102, ...

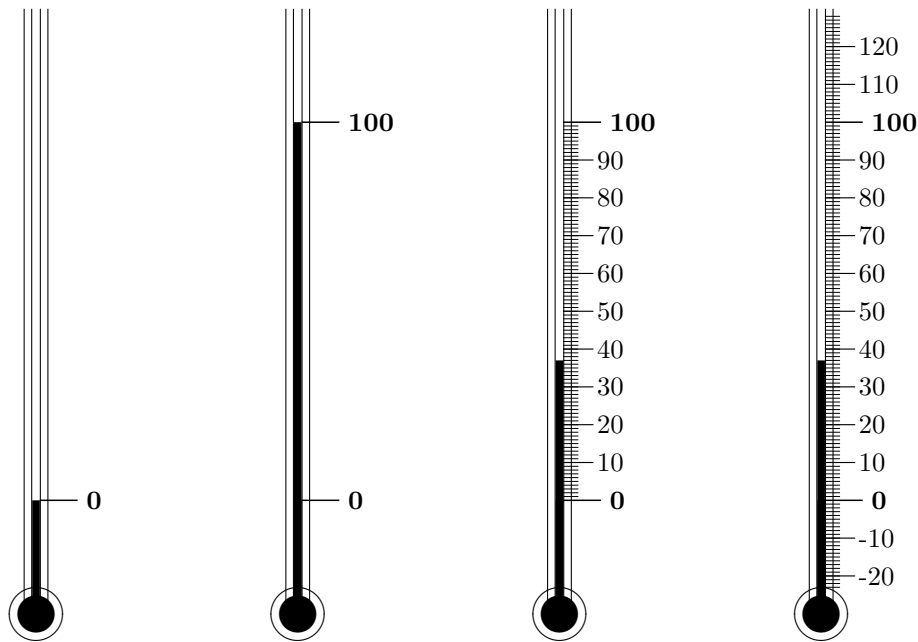
Die folgende Bilderserie zeigt dieses Verfahren. Die Zahlen bekommen noch die Einheit **Grad Celsius** hinzugefügt. Die Abkürzung dafür ist $^{\circ}\text{C}$. Das Wort Grad stammt vom lateinischen Wort *gradus*=Schritt ab. Es erinnert hier daran, dass der Abstand zwischen Gefrier- und Siedetemperatur von Wasser in 100 Schritte eingeteilt worden ist.

Bei einer Temperaturangabe wie z.B. im Satz *Die Körpertemperatur eines Menschen beträgt etwa 37°C* gibt die Maßzahl also an, um wie viele Celsius-Schritte die angegebene Temperatur über der von Eiswasser liegt.

Diese Differenz einer Temperatur zu der Temperatur von Eiswasser wird in Formeln meist mit dem kleinen griechischen Buchstaben ϑ (gelesen: theta) angegeben. Der vorige Beispielsatz könnte also in die Formelsprache übersetzt werden als: $\vartheta_{\text{Mensch}} = 37^\circ\text{C}$

zu ergänzen:

- Eiswasser
- siedendes Wasser



Die Vorgehensweise bei der Celsius-Eichung eines Thermometers lässt sich verallgemeinern. Man benötigt zunächst zwei bestimmte, zuverlässig erzeugbare Temperaturen für die **Fixpunkte** der Temperaturskala. Den Abstand zwischen diesen Fixpunkten nennt man den **Fundamentalabstand** der Skala. Ihn zerteilt man in eine bestimmte Anzahl von Schritten und setzt die Skala mit der gewonnenen Schrittweite nach oben und unten fort.

Bei der Einführung einer neuen Temperaturskala ist es reine Willkür, welche Temperaturen man für die Fixpunkte benutzt, und mit welchen Zahlen man diese Fixpunkte beschriftet. Ähnlich wie Celsius hat René-

Antoine de **Réaumur** (1683-1757) eine Skala in Frankreich festgelegt: Er benutzte die gleichen Fixpunkte wie Celsius, nannte sie jedoch 0°R und 80°R . Die in den angelsächsischen Ländern verbreitete Fahrenheit-Skala von Daniel Gabriel **Fahrenheit** (1686-1736) benutzt andere Fixpunkte: Die niedrigste Temperatur, die Fahrenheit (mit einer bestimmten Mischung aus Wasser, Eis und Salmiak) überhaupt herstellen konnte, nannte er 0°F . Die Temperatur des Blutes in einem Menschen nannte er 100°F . In Celsius-Graden ausgedrückt liegen die Fixpunkte der Fahrenheitskala bei $-17\frac{7}{9}^\circ\text{C}$ und bei $37\frac{7}{9}^\circ\text{C}$. Umgekehrt entsprechen sich 0°C und 32°F , sowie 100°C und 212°F .

Es wäre falsch zu sagen, dass z.B. 0°C *gleich* 32°F wären: Eine Temperaturdifferenz zum Celsius-Nullpunkt ist etwas anderes als eine Temperaturdifferenz zum Fahrenheit-Nullpunkt. Die Gleichheit würde über beiderseitige Division durch 32 außerdem bedeuten, dass 1°F auch gleich 0°C wäre.

Vergleich: Bei einem 32 cm hohen Tischchen mag ja auch die Angabe 0 inch *über dem Tisch* dieselbe Höhenlage bezeichnen wie 32 cm *über dem Fußboden*, ohne dass irgendjemand behaupten würde, 32 cm seien gleich 0 inch. Tatsächlich gilt ja zwischen den Längeneinheiten cm und inch die Umrechnung $1 \text{ inch} = 2,54 \text{ cm}$.

Entsprechend kann man sagen, dass 9 Fahrenheit-Schritte eine ebenso große Temperaturdifferenz sind wie 5 Celsius-Schritte: $\Delta\vartheta = 9^\circ\text{F} \Leftrightarrow \Delta\vartheta = 5^\circ\text{C}$.

Unterscheide also zwischen Temperaturen und Temperaturdifferenzen ebenso wie Du es zwischen Positionen und Strecken und zwischen Zeitpunkten und -spannen tust!

Am besten schreibst Du aber ohne genauere Erläuterung nie eine Gleichung mit unterschiedlichen Temperatureinheiten. Der Leser wird bei Celsius und Fahrenheit nämlich sonst immer davon ausgehen, dass Temperaturangaben bezogen auf den jeweiligen Nullpunkt gemeint seien.

14.3.3 Andere Thermometer

Bimetallthermometer

Flüssigkeitsthermometer nutzen zur Messung der Temperatur aus, dass sich eine Flüssigkeit bei Erwärmung ausdehnt. Es lassen sich aber auch andere Effekte durch Erwärmung ausnutzen, z.B. dass sich verschiedene Metalle unterschiedlich stark ausdehnen.

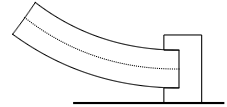
Ein Bimetall-Streifen ist ein Streifen, der aus zwei Schichten aus unterschiedlichen Metallen zusammengesetzt ist. Dehnt sich bei Erwärmung nun z.B. die untere Schicht stärker aus als die obere, so krümmt sich der Bimetallstreifen nach oben. Diese Verformung muss man nur noch in eine Zeigerbewegung umsetzen, eine geeignete Skala anbringen, und schon ist ein Thermometer

weit verbreiteter (weil billiger) Bauart fertig!

normal:



erwärmt:



Meist ist in einem solchen Thermometer eine Spiralfeder aus Bimetall enthalten. Durch die recht große Länge des gewickelten Streifens biegt er sich schon bei geringer Erwärmung um einen recht großen Winkel auf.

Thermocolor

Manche Stoffe ändern ihre Farbe (lat. *color*=Farbe) beim Überschreiten bestimmter Temperaturen. Überzieht man z.B. einen Eisenstab mit einem solchen Stoff, und erhitzt man den Stab an einem Ende stark genug,

so sieht man an der fortschreitenden Verfärbung, wie von diesem Ende her die Wärme sich im Stab ausbreitet.

Pyrometrie

(von griech. *pyros*=Feuer)

Für sehr hohe Temperaturen von z.B. 2000 °C sind die meisten anderen Verfahren nicht mehr geeignet, weil die Thermometer diese Temperaturen gar nicht aushalten würden. Hier kann man ausnutzen, dass ein so hei-

ßes Metallstück je nach Temperatur verschieden farbig glüht. Ein rot glühendes Stück ist sicher nicht so heiß wie ein weiß glühendes.

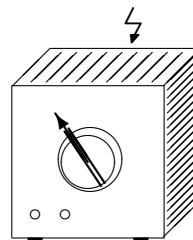
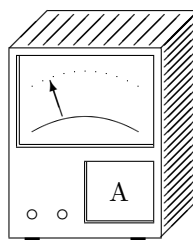
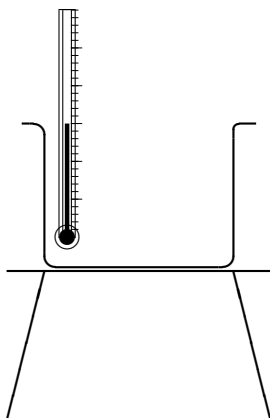
elektrische Thermometer

Auch viele elektrische Größen hängen von der Temperatur ab und eignen sich daher zur Temperaturmessung. Eine dieser Größen ist der elektrische Widerstand, den ein Leiter dem Stromfluss entgegensetzt. Schließt man Leiterstücke an eine feste Spannung an, so ergibt sich eine umso geringere Stromstärke, je größer der Widerstand des Leiterstückes ist. Leiterstücke, die man wegen

ihres Widerstandes in einen Stromkreis einbaut, nennt man ebenfalls kurz *Widerstände*. Legt man an einen Widerstand eine feste Spannung an, so sollte demnach je nach Temperatur mehr oder weniger Strom fließen. Mit dem folgenden Versuch untersuchen wir, wie die Stromstärke durch einen Widerstand von der Temperatur des Widerstandes abhängt.

V 14.5 Widerstandsthermometer

Aufbau:



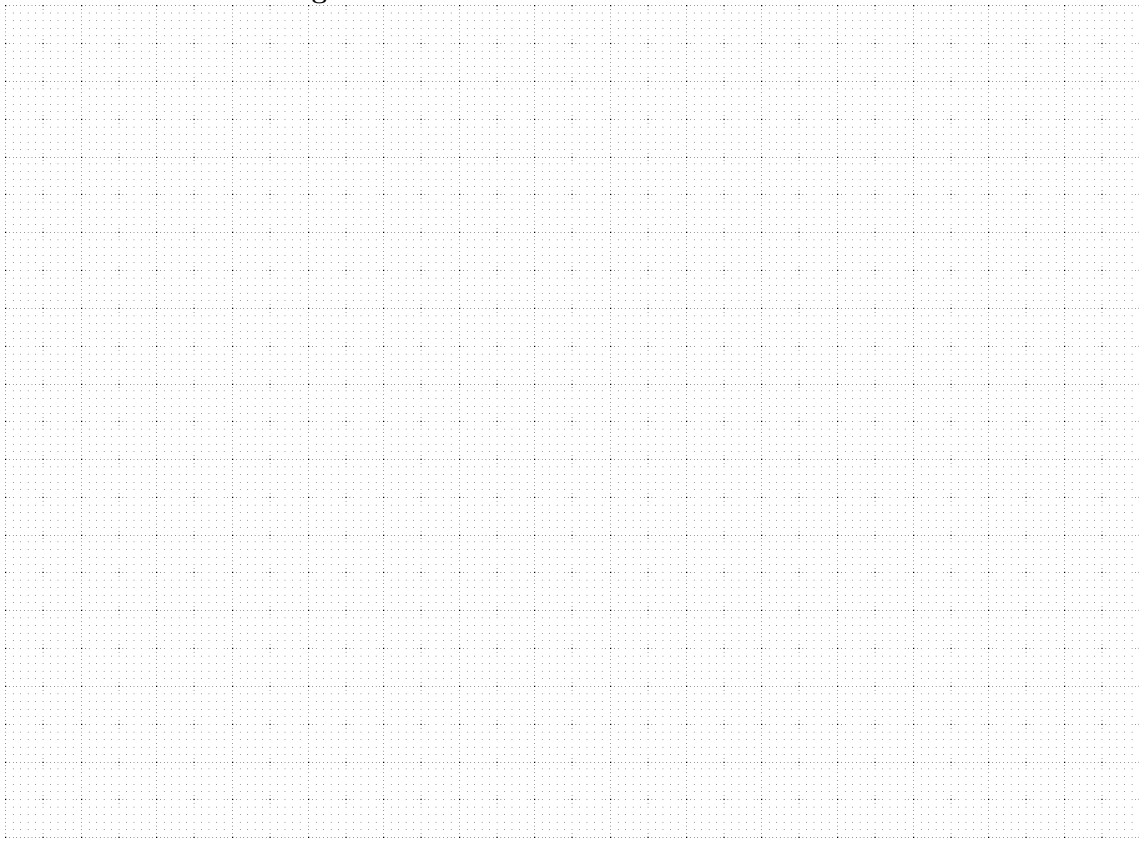
zu ergänzen:

- Wasser
- Heizquelle
- Widerstand
- Verkabelung

Durchführung:

Messwerte:

Auswertung:



Also: Bei dem von uns verwendeten Widerstand _____ die Stromstärke, wenn die Temperatur wächst. Für den Widerstand heißt dies:

Es _____ der Widerstand, wenn die Temperatur wächst.

Würde man den Widerstand zahlenmäßig erfasst gegen die Temperatur auftragen, so erhielte man eine Kurve mit _____ Steigung. Da man diese Steigung auch Temperatur-Koeffizient (englisch: _____

temperature coefficient) des Widerstandes nennt, nennt man unseren Widerstand einen _____TC-Widerstand.

Wie benutzt man nun den Widerstand zur Temperaturmessung?

Man hält den Widerstand dorthin, wo man die Temperatur messen will. Man legt wieder dieselbe Spannung an, mit der man das obige $I(\vartheta)$ -Diagramm aufgenommen hat, und misst die Stromstärke. Nun muss man nur noch aus dem Diagramm die zur gemessenen Stromstärke gehörige Temperatur ablesen.

14.4 Längenausdehnung fester Körper

Körper dehnen sich im Allgemeinen bei Erwärmung aus. Ist diese Dehnung eines Körpers aber bei jedem Grad Temperaturerhöhung gleich? Gemäß der Celsius-Eichung eines Flüssigkeitsthermometers (wie auf Seite 127 beschrieben) ist dies bei der Thermometerflüssigkeit natürlich so, denn die Gradeinteilung wird ja anhand gleicher Zunahmen der Steighöhe definiert. Eine andere Flüssigkeit im Steigrohrchen könnte aber durchaus bei denselben Temperaturerhöhungen um unterschiedliche Höhen ansteigen: z.B. von 29 °C auf 30 °C um 1 mm, von 50 °C auf 51 °C jedoch um 1,02 mm. Genau genommen ist dies in der Regel auch so.

Wir betrachten nun die Verlängerung Δl , die ein Stab der Länge l_0 aufweist, wenn er um $\Delta \vartheta$ erwärmt wird. Es werde z.B. ein 1 m langer Eisenstab von 20 °C um 1 °C auf 21 °C erwärmt. Seine ursprüngliche Länge wächst dabei um 12 µm auf 1,000012 m. Es handelt sich also mit 0,0012% um eine sehr kleine relative Änderung der Länge, die mit bloßem Auge gar nicht zu sehen ist. Auch die Temperaturänderung um 1 °C ist für das Eisen eigentlich recht unbedeutend: Größere Änderungen im mikroskopischen Aufbau des Eisens

(wie z.B. das Schmelzen) finden erst bei viel höheren Temperaturen statt.

Die physikalische Ausgangssituation des Eisens hat sich durch diese Erwärmung also praktisch nicht verändert. Eine weitere Erwärmung um 1 °C auf dann 22 °C findet also etwa unter den gleichen Bedingungen statt wie die vorige Erwärmung. Folglich führt sie auch zur selben Längenzunahme um 12 µm auf nun 1,000024 m. Entsprechend wächst die Länge beim Erwärmen um insgesamt 3 °C um $3 \cdot 12 \mu\text{m}$ und so weiter. Mit anderen Worten: Die Längenänderung ist proportional zur Temperaturänderung: $\Delta l \sim \Delta \vartheta$

Die Längenänderung hängt daneben auch von der Ausgangslänge ab: Ein 3 m langer Stab kann als Verbund dreier 1 m-Stäbe angesehen werden. Da sich jeder der drei 1 m-Abschnitte pro Grad Erwärmung um 12 µm verlängert, wird der gesamte Verbund um $3 \cdot 12 \mu\text{m}$ länger. Kurz: Die Längenänderung ist proportional zur Ausgangslänge: $\Delta l \sim l_0$

Berücksichtigt man noch, dass unterschiedliche Materialien sich verschieden viel dehnen, so lässt sich sagen:

Bei nicht allzu großer Erwärmung um $\Delta \vartheta$ verlängert sich ein fester Körper der Länge l_0 näherungsweise um Δl mit

$$\Delta l =$$

wobei α eine Materialkonstante ist und Längenausdehnungskoeffizient heißt.

zu ergänzen:

- Formelterm

Hierin stecken zwei Vereinfachungen:

1.) Obige Formel liefert für selbes $\Delta \vartheta$ nur dann dasselbe Δl , wenn auch für l_0 immer derselbe Wert eingesetzt wird, z.B. die Länge bei 0 °C. Die momentane Länge des Körpers zu Beginn einer Temperaturänderung hingegen variiert zumindest geringfügig je nach Ausgangstemperatur.

2.) Streng genommen wächst die Länge eines Körpers gar nicht exakt linear mit der Temperatur. Zumindest bei großen Temperaturänderungen zeigen sich recht deutliche Abweichungen.

Die Fehler durch diese Vereinfachungen sind in unseren Betrachtungen jedoch irrelevant, weil wir gewöhnlich schon mit dem Runden von Konstanten und Rechenergebnissen wesentlich größere Fehler produzieren.

V 14.6 Längenausdehnung von Rohren

zu ergänzen:

- Kocher
- Leitungen
- Dampf

Aufbau:



Das Rohr ist links eingespannt und liegt rechts auf einer Schneide auf. Dehnt sich das Rohr, so drückt es die Schneide nach rechts. Schon eine geringfügige Neigung der Schneide ist am Ende des langen Zeigers gut erkennbar. (Bewegung der Zeigerspitze = Bewegung des Rohrendes $\times$ Länge des Zeigers / Höhe der Schneide)

Durchführung:

Ideale Gase

In diesem Abschnitt gehen wir der Ausdehnung von Gasen bei Erwärmung nach. Das Problem der Ausdehnung einer Gasmenge müssen wir jedoch noch genauer fassen, denn ein Gas nimmt ja einfach immer den ganzen ihm zur Verfügung stehenden Raum ein. Wie viel Raum wir dem Gas zur Verfügung stellen, hängt jedoch von uns und nicht von der Temperatur ab.

Erwärmt man eine eingeschlossene Gasmenge, so bewegen sich die Gasteilchen heftiger. Sie stoßen daher auch kräftiger an die Gefäßwände. Mehr Kraft auf dieselben Wandflächen bedeutet, dass die Gasteilchen einen höheren Druck auf die Gefäßwände ausüben. Nur, wenn die Wände diesem erhöhten Druck nachgeben, nimmt das Volumen der Gasmenge zu.

Je nach Art des Einschlusses einer Gasmenge führt eine Erwärmung zu erhöhtem Druck, erhöhtem Volumen oder zu beidem.

Zustandsgrößen
eines Gases:
Druck
Volumen
Temperatur

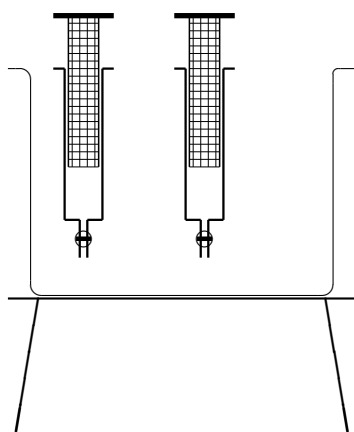
15.1 Erwärmung unter konstantem Druck

Da es immer schwierig ist, mehrere Einflüsse gleichzeitig zu untersuchen, beschränken wir uns zunächst auf Gase bei konstantem Druck. Wie hält man ein Gas unter konstantem Druck? Ganz einfach: Schließt man die Gasmenge in einen Kolbenprober ein, so entspricht ihr Druck jederzeit der Kraft, mit der man den Kolben eindrückt. Lässt man z.B. immer nur den äußeren Luftdruck auf den Kolben einwirken, so herrscht im Inneren des Kolbens auch immer derselbe Druck, nämlich eben der äußere Luftdruck.

Warum? Darum: Drücken das eingespernte Gas von innen und die Umgebungsluft von außen auf den Kolbenquerschnitt, so erfährt der Kolben eine Gesamtkraft in Richtung der größeren Druckkraft. Ist z.B. der Gasdruck größer als der Luftdruck, so wird sich der Kolben nach außen bewegen, so dass der Gasdruck sinkt. Der Kolben wird kräftefrei, wenn Gasdruck und Luftdruck gleich sind. In diese Gleichgewichts-Stellung wird der Kolben sich schließlich einstellen.

V 15.1 Ausdehnung verschiedener Gase

Aufbau:



Füllung 1. Prober: _____

Füllung 2. Prober: _____

zu ergänzen:

- Wasser
- Heizquelle
- evtl. Thermometer

Durchführung:

Beobachtung: / (verallgemeinertes) Ergebnis:



V 15.2 Ausdehnung einer Gasmenge

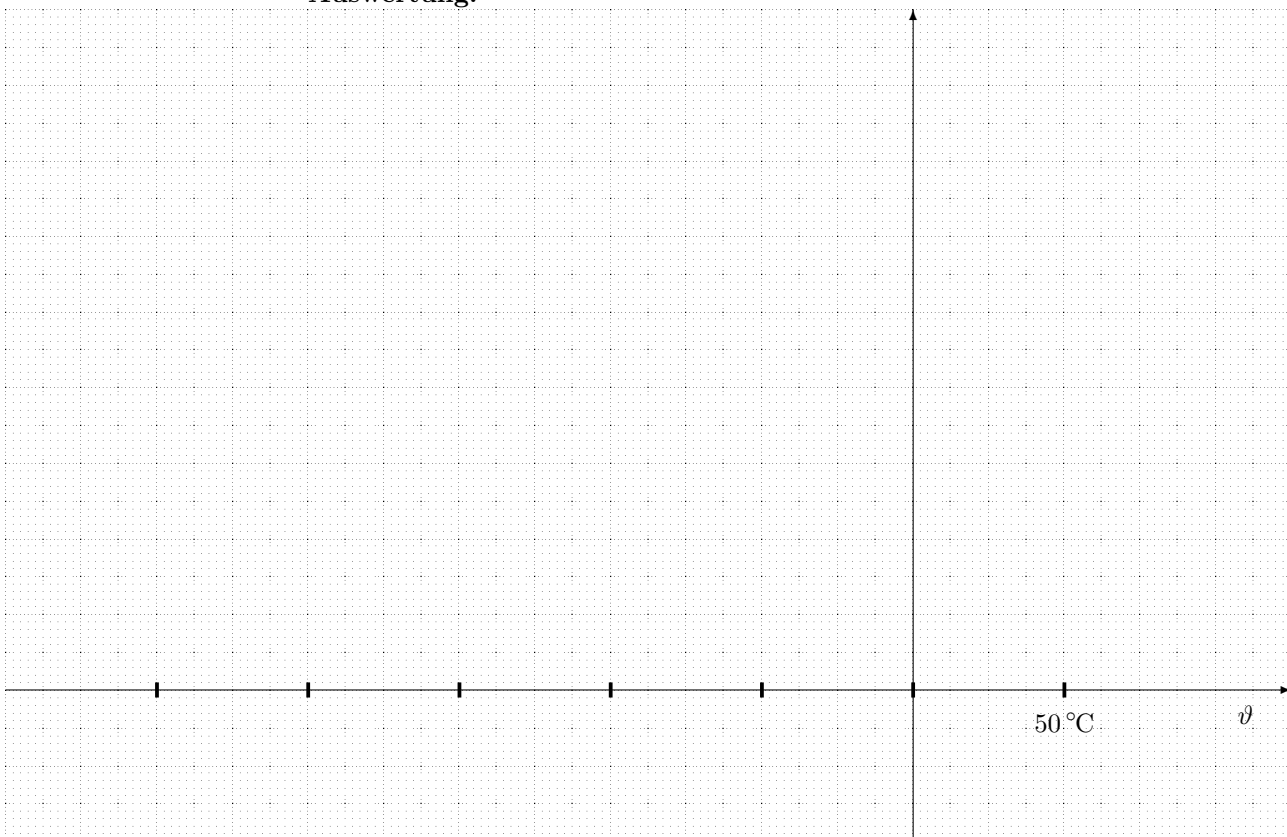
Aufbau: In einem engen, unten geschlossenen Glasröhrchen ist eine Luftmenge durch ein Quecksilbertröpfchen eingeschlossen. (entspricht einem Kolbenprober mit dem Tröpfchen anstelle des Kolbens). Das Röhrchen steht zusammen mit einem Thermometer in einem anfangs kalten Wasserbad.

(Es hat eine Querschnittsfläche von $A =$ _____.)

Durchführung: Das Röhrchen wird im Wasserbad erwärmt. Dabei werden die zu verschiedenen Temperaturen gehörigen Steighöhen gemessen / Volumina der eingeschlossenen Gasmenge bestimmt.

Messwerte:

Auswertung:



Ergebnis:

Da das Volumen V der eingeschlossenen Gasmenge sich als Produkt aus Steighöhe h und der Querschnittsfläche A des Röhrchens ergibt, können wir die Steighöhe auch als Maß für das Volumen ansehen: Bei einer Steighöhe von z.B. 28 mm beträgt

$V = h \cdot A = 28 \text{ mm} \cdot A = 28 \cdot (1 \text{ mm} \cdot A) = 28 \text{ VE}$,
wenn man den Ausdruck $(1 \text{ mm} \cdot A)$ nun als eine neu erfundene Volumeneinheit VE auffasst. Es beträgt bei $h = 28 \text{ mm}$ also auch $V = 28 \text{ VE}$. Die Auswertung an Maßzahlen für Steighöhen könnte demnach ebenso gut eine Auswertung an Maßzahlen für Volumina sein.

Das Röhrchen aus dem vorigen Versuch lässt sich nun als **Gasthermometer** benutzen: Mit der Steighöhe lässt sich mit Hilfe des erhaltenen $V(\vartheta)$ -Diagramms praktisch auch die Temperatur ablesen.

15.2 Die absolute Temperatur

Der vorige Versuch hat gezeigt: Wird eine Gasmenge unter konstantem Druck erwärmt, so nimmt ihr Volumen linear zu, i.a.W. man erhält im $V(\vartheta)$ -Diagramm einen Ausschnitt aus einer Geraden. Die (bei einer sehr genauen Messung) erhaltene Gerade schneidet die ϑ -Achse bei etwa -273°C . Daran ändert sich auch nichts, wenn man den Versuch mit einem anderen Ausgangs-

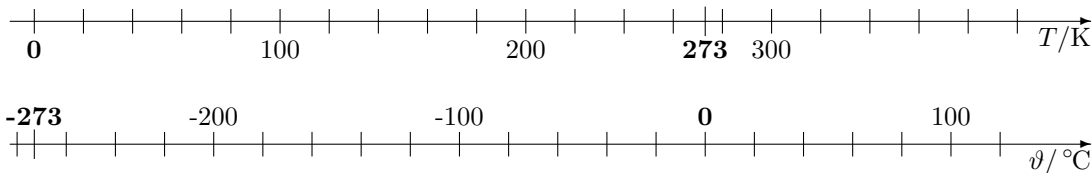
volumen oder einem anderen Gas ausführt.

Im Teilchenmodell stellen wir uns vor, dass die Teilchen eines Gases mit abnehmender Temperatur sich immer weniger heftig bewegen und deshalb um sich herum immer weniger Volumen beanspruchen. Unser voriges Versuchsergebnis legt folgende Vorstellung nahe:

Bei -273°C bewegen sich die Teilchen überhaupt nicht mehr und beanspruchen überhaupt kein zusätzliches Volumen mehr.

Da sich ein Teilchen auch sicher nicht weniger als gar nicht bewegen kann, ist diese Temperatur die tiefstmögliche Temperatur überhaupt: Kälter geht es nicht! (Aus quantenmechanischen Gründen ist diese Temperatur selbst ebenfalls nicht erreichbar, denn ein Teilchen kann nicht völlig ruhig stehen.) Für Physiker ist diese Temperatur (genauer $-273,15^\circ\text{C}$) der ideale Nullpunkt einer Temperaturskala, der sogenannten **absoluten** Temperaturskala:

Der Nullpunkt dieser Skala liegt also bei etwa -273°C . Als ihre Schrittweite hat man die der Celsius-Skala übernommen. Zur Unterscheidung hat man jedoch als Einheitenname **Kelvin** mit der Abkürzung K (kein „Grad“) festgelegt. (zu Ehren von William Lord Kelvin of Largs, 1824-1907, bis 1892 Sir William Thomson) Für Temperaturangaben in dieser Skala schreibt man in Formeln T .



Trägt man das Volumen einer Gasmenge gegen die absolute Temperatur auf, so liegen die Messpunkte auf einer Ursprungsgeraden. Es gilt damit das

Gesetz von Gay-Lussac: Das Volumen V einer Gasmenge ist bei konstantem Druck proportional zu ihrer absoluten Temperatur T . Kurz:

$$V \sim \quad \text{oder} \quad = \text{const}$$

zu ergänzen:

- Terme

Bei früher von uns untersuchten Proportionalitäten haben wir dem konstanten Quotienten einen seiner Bedeutung entsprechenden Namen gegeben. (Geschwindigkeit, Beschleunigung, Federhärte, Dichte, Druck) Der Quotient $\frac{V}{T}$ gibt an, um wie viele cm^3 eine Gasmenge sich bei Erwärmung um 1 K ausdehnt. Man könnte ihn demnach z.B. „thermischen Ausdehnungsfaktor“ nennen. Dieser Faktor hängt jedoch leicht erkennbar sowohl von der Gasmenge als auch von ihrem Druck ab: Doppelt so viele Gasteilchen benötigen bei gleichem Druck immer doppelt so viel Volumen. Dehnt sich die einfache Gasmenge z.B. um 1 cm^3 von 3 auf 4 cm^3 aus, so dehnt sich unter gleichen Bedingungen die doppelte Menge von 6 auf 8 cm^3 um 2 cm^3 aus. Die Volumenausdehnung pro Erwärmung ist also doppelt so groß wie bei der einfachen Gasmenge. Ebenso nimmt dieselbe Gasmenge unter niedrigerem Druck immer ein höheres Volumen ein, so dass auch die Volumenänderungen bei niedrigerem Druck größer sind als bei höherem Druck.

Wegen dieser Abhängigkeit von beidem, der Teilchenzahl und dem Druck, ist der „Ausdehnungsfaktor“ weder typisch für die Gasmenge an sich, noch für die Umstände ihrer Aufbewahrung an sich. Deshalb ist es nicht üblich, den Quotienten $\frac{V}{T}$ eigens zu benennen.

Bei verschiedenen Temperaturen $T_{1,2}$ hat dieselbe Gasmenge unter gleichem Druck also verschiedene Volumina $V_{1,2}$, aber denselben Quotienten $\frac{V_{1,2}}{T_{1,2}}$. Ohne, dass der Quotient selbst als benannte (Zwischen-)Größe darin auftaucht, ist für Berechnungen also letzten Endes folgende Formel anzuwenden:

$$\frac{V_1}{T_1} = \frac{V_2}{T_2}$$

Proportionalität von V und T bedeutet natürlich z.B. auch, dass bei 17-facher Temperatur das 17-fache Volumen eingenommen wird. Manches wird leichter, wenn man sich auf die Änderungsfaktoren konzentriert.

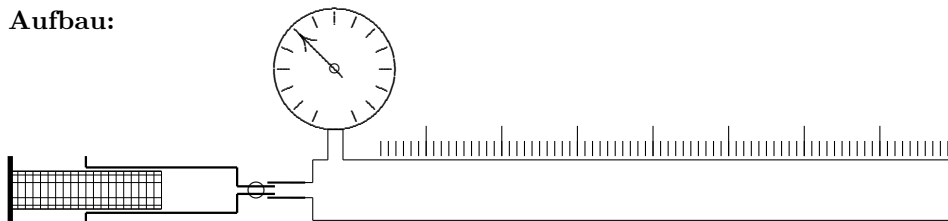
15.3 Gase bei konstanter Temperatur

V 15.3 Boyle-Mariotte

zu ergänzen:

- Stahlkugel
- Beschriftung
- interessierende Gasmenge

Aufbau:



Durchführung:

Messwerte: /Auswertung:

Ergebnis:

zu ergänzen:

- Terme

Gesetz von Boyle-Mariotte: Das Volumen V einer Gasmenge ist bei konstanter Temperatur umgekehrt proportional zu ihrem Druck p .

$V \sim$

oder

$= \text{const}$

Bei verschiedenen Drücken $p_{1,2}$ hat dieselbe Gasmenge bei gleicher Temperatur verschiedene Volumina $V_{1,2}$, aber dasselbe Produkt $p_{1,2} \cdot V_{1,2}$. Ohne, dass das Produkt selbst als benannte (Zwischen-)Größe darin auftaucht, ist für Berechnungen also letzten Endes folgende Formel anzuwenden:

$$p_1 \cdot V_1 = p_2 \cdot V_2$$

Auch bei umgekehrter Proportionalität ist es oft zweckmäßig, die Änderungsfaktoren zu betrachten: Zum 23-fachen Druck gehört ein 23-stel Volumen.

15.4 Zusammenführung beider Gesetze

Nun fehlen natürlich noch Aussagen über ein Gas bei konstantem Volumen und vor allem über ein Gas bei gar keiner konstanten Zustandsgröße. Diese Aussagen können wir jedoch aus den beiden Gesetzen von Gay-Lussac und Boyle-Mariotte mit einem kleinen Kunstgriff herleiten.

Wir betrachten den allgemeinen Fall, dass bei irgendeinem Versuch eine abgeschlossene Gasmenge aus einem Anfangs-Zustand 1, gekennzeichnet durch (p_1, V_1, T_1) , in einen End-Zustand 2 mit (p_2, V_2, T_2) übergeht, ohne dass wir irgendwelche weiteren Annahmen über diese 6 Größen machen. Es können sich also der Druck, die Temperatur und das Volumen *gleichzeitig* ändern. Damit ist erst einmal keines der beiden uns bekannten Gesetze anwendbar, bei denen ja jeweils eine

Größe konstant sein muss.

Man hat aber in allen Versuchen mit Gasen gesehen, dass mit zweien der drei Zustandsgrößen immer auch die dritte festgelegt ist. Beispiel: Bei bestimmtem Druck und bestimmter Temperatur einer Gasmenge ist auch ihr Volumen eindeutig festgelegt, egal wie die Gasmenge vorher behandelt (z.B. verdichtet, abgekühlt, erwärmt) worden ist.

Ob sich nun bei einem Vorgang alle drei Zustandsgrößen einer Gasmenge *auf einmal* ändern, spielt daher für den End-Zustand gar keine Rolle: denselben End-Zustand erreicht man auch *in Schritten*, bei denen man immer eine Größe konstant lässt und nur die beiden anderen ändert, wie es die bereits bekannten Gesetze zulassen.

Statt „Alles-gleichzeitig-Prozess“:

$$\begin{pmatrix} p_1 \\ V_1 \\ T_1 \end{pmatrix} \longrightarrow \begin{pmatrix} p_2 \\ V_2 \\ T_2 \end{pmatrix}$$

lieber mit Zwischenzustand
in 2 Schritten nacheinander:

$$\begin{pmatrix} p_1 \\ V_1 \\ T_1 \end{pmatrix} \xrightarrow{\text{Gay-Lussac}} \begin{pmatrix} p_1 \\ V' \\ T_2 \end{pmatrix} \xrightarrow{\text{Boyle-Mariotte}} \begin{pmatrix} p_2 \\ V_2 \\ T_2 \end{pmatrix}$$

Rechnungen zu beiden Schritten,
Gleichsetzen über Zwischen-Volumen,
Ordnen nach Index:

$$\begin{array}{ccc} \frac{V_1}{T_1} = \frac{V'}{T_2} & \Big| & p_1 V' = p_2 V_2 \\ V_1 \cdot \frac{T_2}{T_1} = V' & = & V' = \frac{p_2}{p_1} \cdot V_2 \\ V_1 \cdot \frac{T_2}{T_1} & = & \frac{p_2}{p_1} \cdot V_2 \end{array}$$

$$\boxed{\frac{p_1 \cdot V_1}{T_1} = \frac{p_2 \cdot V_2}{T_2}}$$

Der Ausdruck $\frac{pV}{T}$ hat also in beiden Zuständen 1 und 2 den gleichen Wert. Da die Herleitung für beliebige Zustände durchgeführt wurde, hat dieser Ausdruck in *jedem* Zustand denselben Wert.

Druck p , Volumen V und Temperatur T einer Gasmenge können sich nur so ändern, dass

$$\frac{p \cdot V}{T} = \text{const}$$

Für Berechnungen bei einer Änderung zwischen zwei Zuständen ist natürlich die Gleichung $\frac{p_1 V_1}{T_1} = \frac{p_2 V_2}{T_2}$ bestens geeignet. Man beachte, dass für konstante Temperatur, d.h. $T_1 = T_2 = T$, diese Gleichung mit T multipliziert (und gekürzt) gerade das Gesetz von Boyle-

Mariotte in der Form $p_1 V_1 = p_2 V_2$ darstellt. Ebenso ergibt sich bei konstantem Druck $p_1 = p_2 = p$ nach Division beider Seiten der Gleichung durch p das Gesetz von Gay-Lussac in der Form $\frac{V_1}{T_1} = \frac{V_2}{T_2}$.

15.5 Die ideale Gasgleichung

15.5.1 Vergleich verschieden großer Gasmengen

Obwohl wir im vorigen Abschnitt schon eine Gleichung gefunden haben, die alle Zustandsgrößen einer Gasmenge miteinander verknüpft, gehen wir in diesem Abschnitt noch einen Schritt in der Verallgemeinerung weiter. Bisher galten alle Überlegungen für *eine* Gasmenge. Was ändert sich nun aber beim Übergang zu *anderen* Gasmengen, insbesondere an dem für eine Gasmenge typischen Ausdruck $\frac{pV}{T}$? Dazu konkreter ein einfaches Problem: Wie ändern sich die Zustandsgrößen, wenn eine Gasmenge halbiert wird? Angenommen, die Gasmenge befindet sich in einem abgeschlossenen Kasten mit Volumen V_0 , stehe unter dem Druck p_0 bei der Umgebungstemperatur T_0 . Wie kann man diese Gasmenge halbieren?

Schiebt man in die Mitte des Kastens eine Trennwand ein, so hat man z.B. in dem „Halbkasten“ links der Wand sicher die halbe Gasmenge wie im ganzen Kasten. Welche Auswirkung hatte das auf die Zustandsgrößen? Das Einschieben der Trennwand hat sicher weder die Temperatur noch den Druck im ganzen Kasten geändert. Die halbe Gasmenge steht also nach wie vor unter dem Druck $p = p_0$ bei der Temperatur $T = T_0$. Ihr Volumen V ist jedoch nur noch das eines Halbkastens, $V = \frac{1}{2}V_0$. Der typische Ausdruck für die halbe Gasmenge ist daher

$$\frac{pV}{T} = \frac{p_0 \cdot \frac{1}{2}V_0}{T_0} = \frac{1}{2} \frac{p_0 V_0}{T_0}$$

also halb so groß wie für die ursprüngliche ganze Gasmenge.

Natürlich kann man eine Gasmenge auch anders halbieren. Man kann z.B. den Kasten mit einem zweiten, gleich großen, jedoch anfangs leeren Kasten verbinden. Das Gas verteilt sich dann gleichmäßig im so entstandenen „Doppelkasten“. Im ursprünglichen Kasten befindet sich dann nur noch halb so viel Gas.

Diese halbe Gasmenge hätte das gleiche Volumen $V = V_0$, das die ganze Gasmenge ursprünglich hatte. Beim Ausdehnen vom ersten Kasten in den Doppelkasten wäre aber sicher der Gas-Druck abgefallen und das Gas hätte sich auch etwas abgekühlt. Beim anschließenden Angleichen der Temperatur an die Umgebungstemperatur wären Druck und Temperatur wieder etwas gestiegen. Es wäre $T = T_0$ geworden und $p = \frac{1}{2}p_0$. (So, wie es uns der Boyle-Mariotte'sche Versuch bei Verdopplung des Volumens zeigt.) Auch hier käme also in den Ausdruck $\frac{pV}{T}$ ein Faktor $\frac{1}{2}$ hinein.

Wir können sogar verallgemeinern: Bei Halbierung der Gasmenge halbiert sich der Ausdruck $\frac{pV}{T}$, egal wie die Halbierung stattfindet. Und was für den Faktor $\frac{1}{2}$ gilt, gilt natürlich auch für jeden anderen Faktor:

Der Ausdruck $\frac{pV}{T}$ einer Gasmenge ist proportional zur Menge an Gas.

15.5.2 Die Stoffmenge

Es muss allerdings noch präzisiert werden, was unter *Menge* zu verstehen ist. Dazu eine kleiner Einschub zur Frage

?

Wie kann man angeben, aus *wie viel* Stoff eine Stoffportion besteht?

Gebräuchlich sind für Mengenangaben schon im Alltag drei Größen:

Volumen V : Man kauft z.B. 3 Liter Milch oder $2\frac{1}{2} \text{ m}^3$ Sand.

Masse m : Man kauft z.B. 100 g Salami oder einen 1 kg-Barren Gold.

Stückzahl n : Man kauft z.B. 12 Zwiebeln oder einen 1000er-Pack Schrauben.

Natürlich kann man jede Angabe in einer dieser Formen in eine der anderen Formen bringen: Mit bekannter Dichte ρ des Materials lassen sich Angaben für V und m gegenseitig ersetzen. Bekanntlich ist $m = \rho \cdot V$. Sofern Du weißt, welche Masse M eine Schraube hat, kannst Du auch ausrechnen, welche Masse m die Schrauben in einem Pack mit n Stück haben: $m = n \cdot M$

Je nach Zweck sind die genannten Größen unterschiedlich gut für Mengenangaben geeignet.

- Ein Volumen lässt sich oft leicht abmessen. Es hat aber den Nachteil, dass ein und dieselbe Stoffportion je nach äußeren Umständen unterschiedliche Volumina hat. Gerade dies untersuchen wir ja gerade bei Gasen. Das Volumen ist also hier für uns nicht geeignet, um die Frage nach dem *Wie viel* einer Gasportion zu beantworten.
- Die Masse ist da schon besser: Für Physiker ist diese Größe *die* Größe zur Beschreibung von Stoffportionen. Man kann Massen leicht abmessen, dieselbe Stoffportion hat unabhängig von äußeren

Umständen (praktisch) immer die gleiche Masse, und von der Masse hängen die wichtigen Körpereigenschaften Gewicht und Trägheit ab.

- Stückzahlen sind immer dann interessant, wenn es um einzelne gleichartige Teilchen geht, aus denen eine Menge besteht. In der Physik und noch mehr in der Chemie sind die interessierenden Teilchen oft die Atome oder Moleküle, aus denen eine Stoffportion besteht. Chemische Reaktionsgleichungen sind auch immer Aussagen über die Anzahlen der Teilchen je einer Sorte, die miteinander reagieren.

Für die in der Chemie üblichen riesigen Stückzahlen hat man eine eigene Einheit eingeführt. Man gibt die Anzahl z.B. der Moleküle in einer Stoffprobe nicht in „Stück“, „Dutzend“ oder sonstigen aus dem Alltag bekannten Anzahls-Einheiten an, sondern in **Mol**. Man denkt dabei auch weniger an *Stückzahlen* von einzelnen Teilchen, sondern an die gesamte von den Teilchen gebildete *Stoffmenge*. Entsprechend führt man als Größe die **Stoffmenge** ein:

Die Stoffmenge n einer Stoffportion gibt an, aus wie vielen Teilchen bestimmter, mitanzugebender Art die Portion besteht. Eine Stoffmenge von 1 Mol besteht aus etwa $6,022 \cdot 10^{23}$ Teilchen. Diese Zahl heißt Avogadro-Zahl N_A .

Warum ist die Avogadro-Zahl gerade diese „krumme“ Zahl?

?

Bei der Einführung dieser Einheit kannte man diese Zahl noch gar nicht. Niemand konnte/kann tatsächlich so viele einzelne Teilchen abzählen. Erstens war es damals unmöglich, überhaupt atomare Teilchen einzeln zu untersuchen. Zweitens würde selbst ein moderner Computer mit einer Taktfrequenz von 1 GHz (etwa 10^9 Takte pro Sekunde) fast 18 Millionen Jahre brauchen, um von 1 bis zur Avogadro-Zahl zu zählen.

Grundidee damals: Es möge z.B. jeweils *ein* Teilchen von Stoff A mit *einem* Teilchen von Stoff B chemisch reagieren können. Wenn nun z.B. eine 5 g-Portion

von Stoff A restlos mit einer 8 g-Portion von B reagiert, dann enthalten die Portionen sicherlich gleich viele Teilchen, auch wenn man diese gleiche Anzahl selbst nicht kennt. Es werden dann auch 10 g von A mit 16 g von B restlos reagieren: Hauptsache, das Verhältnis stimmt.

Da es nur auf Verhältnisse ankommt, ist das Mol auch heute noch eigentlich definiert als Stoffmenge, die aus ebensovielen Teilchen besteht wie 12 g Kohlenstoff (C12). Erst nachträglich konnte man immer genauer angeben, wie viele Teilchen das denn nun sind.

Warum gerade 12 g Kohlenstoff?

?

Warum Kohlenstoff? Aus praktischen, hier uninteressanten Gründen. Warum 12? Um sich das Rechnen leicht zu machen!:

Atome bestehen zu fast 100% ihrer Masse aus sogenannten „Nukleonen“, die alle etwa gleich viel Masse haben. Bei Kohlenstoff-Atomen sind es allermeist 12 Nukleonen. Ein solches C12-Atom hat also die zwölfwache Masse wie ein einzelnes Nukleon.

Folgende dreisatz-artige Überlegung zeigt, wieso die Definition des Mol das Rechnen leicht macht. Als Beispiel wird hier berechnet, wie viel 1 Mol Gold wiegt, dessen Atome jeweils 197 Nukleonen enthalten:

Def.:	1 mol	C12	wiegt	12 g
d.h.	N_A	C12-Atome	→	12 g
	$N_A \cdot$	(12 Nukleonen)	→	12 g
	$N_A \cdot$	1 Nukleon	→	1 g
	$N_A \cdot$	(197 Nukleonen)	→	197 g
	N_A	Gold-Atome	→	197 g
	1 mol	Gold	wiegt	197 g

Mit der Anzahl der Nukleonen in einem Teilchen (Atom, Molekül oder sonstige chemische Formeleinheit) kann man also sofort angeben, wie viel Gramm 1 Mol des Stoffes aus diesen Teilchen wiegt.

Übrigens hat man auch passend dazu die atomare Masseneinheit u definiert: $1 u := 1 \text{ g}/N_A$, so dass die Anzahl der Nukleonen in einem Teilchen eines Stoffes nicht nur die Masse eines Mols des Stoffes in der Einheit Gramm, sondern auch die Masse eines Teilchens in der Einheit u angibt.

Die molare Masse eines Stoffes, dessen Formeleinheit x Nukleonen enthält, beträgt $x \frac{\text{g}}{\text{mol}}$ und jede Formeleinheit wiegt $x u$.

Die Atommassen der Elemente kann man in einem Periodensystem nachlesen: In der Regel gibt in jedem Kästchen die Zahl oberhalb des Elementsymbols die Atommasse des Elementes als Vielfaches von $1 u$ an. Die reine Zahl ohne die Einheit u heißt **relative Atommasse**. Aus drei Gründen sind diese relativen Atommassen gebrochene Zahlen:
Erstens fasst man unter dem Begriff *Nukleon* Proto-

nen und Neutronen zusammen, die zwar ähnliche, aber nicht gleiche Massen haben: jeweils *ungefähr* $1 u$.

Zweitens ändern sich die Massen je nach Zusammensetzung des Atomkerns ein wenig gemäß der Relativitätstheorie.

Drittens und vor allem sind nicht alle Atome eines Elementes gleich gebaut, so dass man in einem Periodensystem nur eine durchschnittliche Atommasse findet.

Anfang des Periodensystems					
1.00794 1 H Wasserstoff					4.002602 2 He Helium
6.941 3 Li Lithium	9.012182 4 Be Beryllium	10.811 5 B Bor	12.011 6 C Kohlenstoff	14.00674 7 N Stickstoff	15.9994 8 O Sauerstoff
22.98977 11 Na Natrium	24.3050 12 Mg Magnesium	26.98154 13 Al Aluminium	28.0855 14 Si Silizium	30.97376 15 P Phosphor	32.066 16 S Schwefel
				35.4527 17 Cl Chlor	39.948 18 Ar Argon

15.5.3 Die ideale Gasgleichung

Man kann nun feststellen, dass zwei Gasmengen denselben Wert von $\frac{pV}{T}$ haben, wenn sie aus gleich vielen Gasteilchen bestehen, wenn also die Stoffmengen gleich sind. Die Art der Gasteilchen spielt idealerweise keine Rolle. Es ist also $\frac{pV}{T}$ proportional zur Stoffmenge n . Die Proportionalitätskonstante $R := \frac{pV}{nT}$ ist eine universelle Naturkonstante.

Das Verhalten von Gasen lässt sich weitgehend beschreiben mit der Gleichung

$$p \cdot V = n \cdot R \cdot T$$

mit der universellen Gaskonstanten $R = 8,31441 \frac{\text{J}}{\text{mol} \cdot \text{K}}$.

Tatsächlich gibt es je nach Gasart in bestimmten Wertebereichen der Zustandsgrößen auch deutliche Abweichungen von diesem idealen Verhalten. Die ideale Gasgleichung passt v.a. dann schlecht, wenn das Gas nahe am Kondensieren ist.

Innere Energie

16.1 Wärme

16.1.1 Mikroskopische Deutung

Die Vorstellung, dass mikroskopische Verhakungen die Ursache für Reibung sind, legt nahe, dass an den reibenden Oberflächen im mikroskopischen Maßstab mechanische Arbeit geleistet wird: Beim Gleiten werden hervorstehende Unebenheiten bei jedem Verhaken verbogen und federn wieder zurück, wenn sich die Verhakung löst.

So erhalten diese winzigen Teile an den Oberflächen Spann- und Bewegungsenergie. Auch benachbarte Teilchen werden mitgerissen, so dass schließlich alle Teilchen in den reibenden Körpern Energie abbekommen. Da diese Energie in lauter Vorgängen und Zuständen innerhalb der Körper besteht, nennt man sie **innere**

Energie der Körper.

Aufgrund dieser inneren Energie können alle Teilchen eines Körpers sich heftig bewegen und schwingen. Da all diese Bewegungen aber im Großen völlig unkoordiniert sind, bewegt sich der Körper als Ganzes nicht. Es ist ja nicht so, dass sich z.B. einmal *alle* Teilchen nach rechts bewegen, so dass der ganze Körper sich nach rechts bewegen würde. Man kann sich den ganzen Körper wie einen betriebsamen Ameisenhaufen vorstellen, der sich als Ganzes ja auch nicht bewegt.

Äußert sich diese innere Unruhe im Körper dann nach außen gar nicht? Doch: Erfahrungsgemäß werden Körper beim Reiben im Ganzen *warm*.

16.1.2 Zusammenhang mit der Temperatur

Um den Zusammenhang zwischen der inneren Energie eines Körpers und seiner Temperatur experimentell zu ermitteln, erhöhen wir die innere Energie eines Körpers durch bestimmte Energiezufuhr und beobachten dabei seine Temperatur.

Die innere Energie erhöhen wir am einfachsten, indem wir Reibungsarbeit am Körper verrichten. Bei der Planung des Versuchs muss man dennoch allerlei beachten:

1. Wir müssen genau bestimmen können, wie viel Reibungsarbeit verrichtet wird. Es muss dazu eine gleichbleibende und bekannte Reibungskraft entlang eines bekannten Weges wirken. Nur dann können wir die verrichtete Reibungsarbeit als Produkt aus Reibungskraft und Weg berechnen.
2. Reibung findet immer an *zwei* aneinander reibenden Körpern statt, so dass dabei immer bei zwei Körpern die innere Energie erhöht wird. (Zum Beispiel: gezogener Klotz und Tisch als Unterlage) Am

einfachsten wäre es, wenn wir sicherstellen könnten, dass nur *einer* der beiden Körper praktisch alleine die gesamte Reibungsarbeit als innere Energie aufnähme.

3. Um einen gut hantierbaren Körper messbar zu erwärmen ist erfahrungsgemäß schon recht viel Reibungsarbeit nötig: Der Holzklotz, der einen Meter weit über den Tisch gezogen wird, erwärmt sich kaum. Man braucht also große Reibungskräfte oder lange Wege.

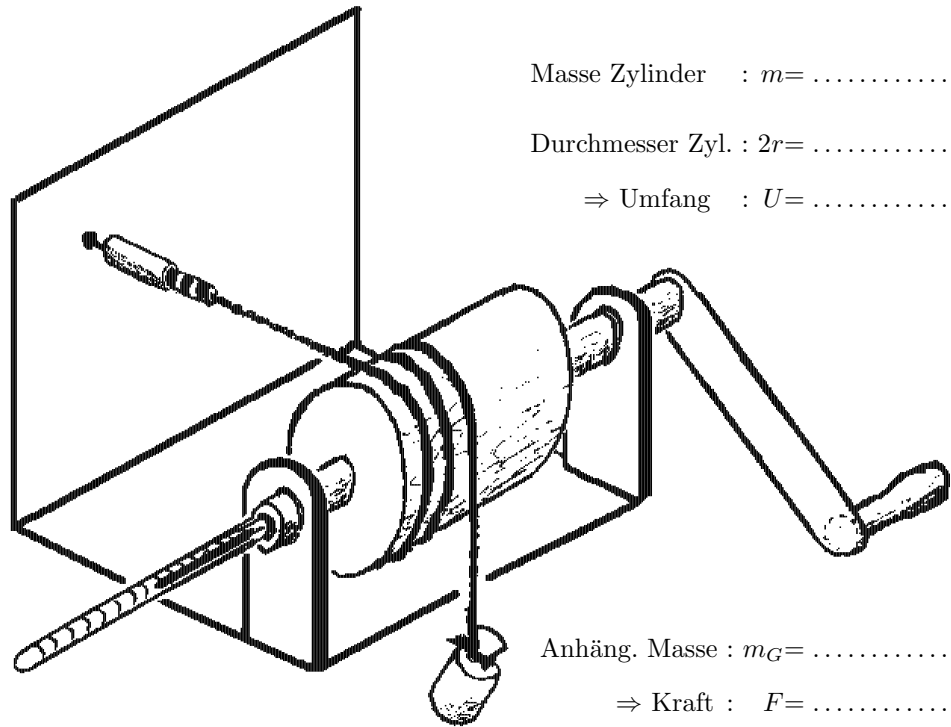
Im folgenden Versuch ist diesen Überlegungen dadurch Rechnung getragen, dass ...

- 1... die Reibungskraft durch die konstante Gewichtskraft eines Massestückes festgelegt wird.
- 2... einer der reibenden Körper, nämlich die Schnur, aufgrund ihrer geringen Masse wohl auch nur geringe Energiebeträge aufnehmen kann.
- 3... einfach durch beliebige Anzahlen an Umdrehungen beliebig weite Reibungswege realisiert werden können.

V 16.1 Temperaturänderung durch Reibung

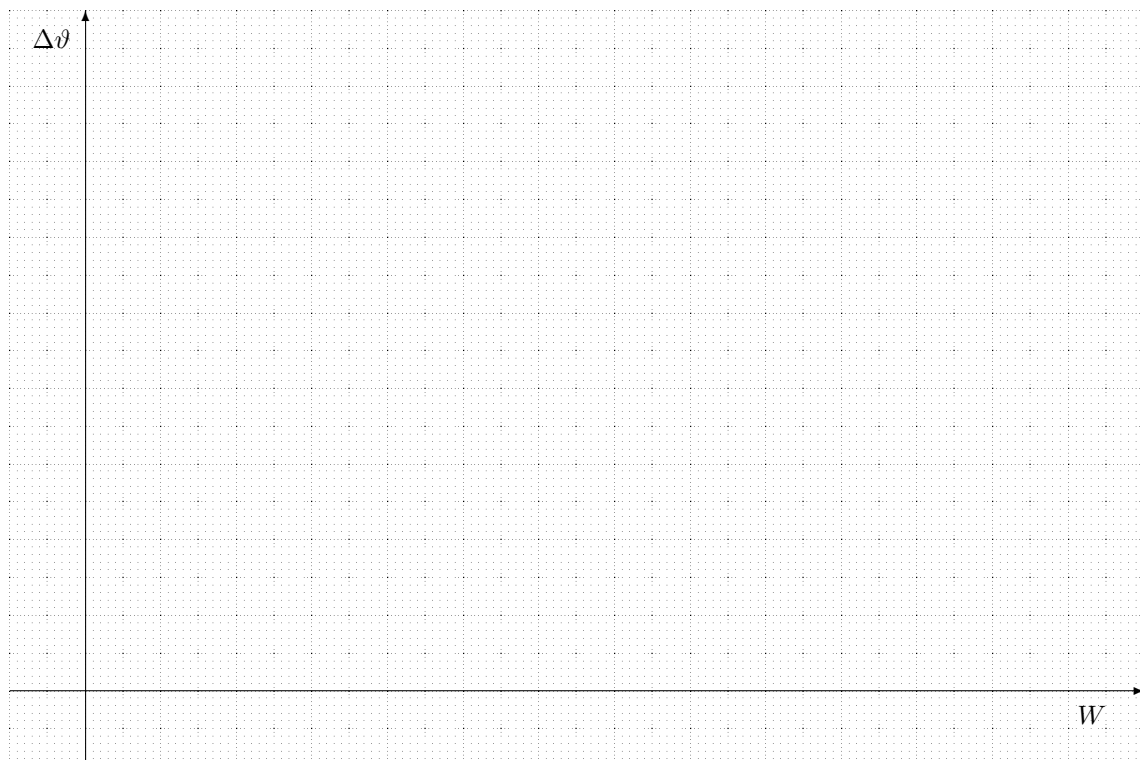
Aufbau: Die Schnur wird so oft um den Zylinder gewickelt, dass der Kraftmesser Null anzeigt, wenn man kurbelt. (Drehsinn noch einzeichnen!)

Durchführung: Wir messen zunächst die Temperatur des Zylinders. Dann drehen wir den Zylinder und lesen nach bestimmten Umdrehungszahlen wieder die Temperatur ab.



Messwerte: / Auswertung:

Umdrehungszahl n	0						
⇒ W_{Reib} (=) [J]	0						
ϑ_n [°C]							
⇒ $\Delta\vartheta$ ($=\vartheta_n - \vartheta_0$) [°C]	0						



Ergebnis:

Der Proportionalitätsfaktor c gibt an, wie viel Energie nötig ist, um den Zylinder um 1°C zu erwärmen. Eine Energie, die zur Änderung der inneren Energie eines Körpers ausgetauscht wird, nennt man jedoch traditionell **Wärme**. Damit formuliert gibt c also an, wie viel Wärme der Zylinder bei der Erwärmung um 1°C aufnimmt. Man nennt diesen Faktor deshalb die **Wärmekapazität** des Zylinders. (Kapazität=Fassungsvermögen, von lateinisch *capere*=(auf-)nehmen; daher auch die Abkürzung c von lateinisch/französisch/englisch *capacitas/-té/-ty*)

Wovon hängt die Wärmekapazität eines Körpers ab?

?

Es erfordert sicher die doppelte Energie, zwei gleichartige Zylinder zu erwärmen. (Den zweiten Zylinder zu erwärmen geht ja genauso wie beim ersten.) Da man dabei insgesamt die doppelte Masse erwärmt, vermuten wir, dass die Wärmekapazität proportional zur Masse des Körpers ist. Dies könnten wir durch einen wiederholten Reibungsversuch mit sonst gleichen Zylindern unterschiedlicher Masse bestätigen.

Also: $c = c_S \cdot m$ mit einem weiteren Proportionalitätsfaktor c_S
Diesen Faktor erhalten wir als $\frac{c}{m}$, also als Wärmekapazität pro Masse.

In unserem Fall beträgt $c_S = \dots\dots\dots$

Wie schon an der Einheit abzulesen ist, gibt c_S an, wie viel Energie nötig ist, um

Wovon hängt c_S ab?

?

Führt man die bisher durchgeführten Reibungsversuche erneut mit Zylindern aus einem anderen Stoff (Material) durch, erhält man einen anderen Wert für c_S . Offensichtlich ist c_S eine *Stoffkonstante* (auch *Materialkonstante*). Wir schreiben daher zukünftig statt S die Stoffbezeichnung als Index, z.B. bei Aluminium c_{Al} . Diese Stoffkonstante heißt **spezifische Wärmekapazität** des Stoffes. Das Adjektiv *spezifisch* stammt vom lateinischen *species*=Art (hier: des Materials).

Wird ein Körper der Masse m aus einem Stoff S um $\Delta\vartheta$ erwärmt, so steigt seine innere Energie um

$$\Delta W =$$

wobei c_S die spezifische Wärmekapazität des Stoffes ist.

zu ergänzen:

- Formelterm

16.1.3 Mischungsversuche

Mit Reibungsversuchen wie dem zuletzt beschriebenen kann man von weiteren Stoffen die spezifische Wärmekapazität bestimmen, allerdings nur von festen Stoffen, aus denen sich ein Zylinder herstellen lässt.

Mit gleicher Vorgehensweise wie im letzten Versuch lässt sich bei bekannter spezifischer Wärmekapazität von Wasser die Kapazität fester Stoffe nun ohne Reibungsversuch bestimmen: Eine bestimmte Menge des interessierenden Stoffes wird im Wasserbad auf bestimmte Temperatur erhitzt und in einer bestimmten Menge kalten Wassers bestimmter Temperatur abgekühlt. Aus der sich einstellenden Mischtemperatur wird schließlich die unbekannte Wärmekapazität errechnet.

Der Ansatz zur Auswertung – ebenso wie zum

Lösen vieler Aufgaben – ist dabei stets der gleiche: Der wärmere Körper gibt dem kälteren solange Energie ab, bis beide Körper die gleiche Temperatur haben. Die Energie, die der wärmere Körper beim Abkühlen auf die Mischtemperatur *abgibt*, ist genau die Energie, die der kältere Körper beim Erwärmen auf die Mischtemperatur *aufnimmt*. In der folgenden Rechnung werden die Masse, die spezifische Wärmekapazität und die Anfangstemperatur des anfangs *wärmeren* Körpers mit dem Index W gekennzeichnet, die Größen des *kälteren* mit K.

$$W_{ab} = W_{auf}$$

$$c_W \cdot m_W \cdot (\vartheta_W - \vartheta_{Misch}) = c_K \cdot m_K \cdot (\vartheta_{Misch} - \vartheta_K)$$

16.1.4 Energieerhaltung

Der im vorigen Abschnitt vorgestellte Ansatz $W_{ab} = W_{auf}$ lässt sich natürlich auch ausdehnen auf Fälle, in denen mehr als 2 Körper beteiligt sind. Es wird einfach angesetzt:

Summe aller abgegebenen Energien	=	Summe aller aufgenommenen Energien
-------------------------------------	---	---------------------------------------

Dies ist eigentlich nur eine andere Formulierung des Prinzips gleichbleibender Energie-Summen auf Seite 65: Wenn die Summe aller Energien im Laufe der Zeit immer gleich bleibt, muss ja jede irgendwo abgegebene Energie sonstwo aufgenommen werden. Man darf

in Energiebilanzen nur keine Energie(-form) vergessen! Unter Berücksichtigung aller (uns bekannten wie der uns noch unbekannten) Energieformen gilt nämlich als einer der Grundpfeiler der gesamten Physik:

Energieerhaltungssatz:
Die Summe aller Energien bleibt alle Zeit erhalten.

Als wir die Reibung untersucht haben, hat es so ausgesehen, als ob über Reibungsarbeit doch Energie verloren gehen könnte. Die Bilanzgleichung für Vorgänge mit Reibung auf Seite 66 können wir aber nun doch als Sonderfall der Bilanzgleichung auf Seite 65 deuten: Die

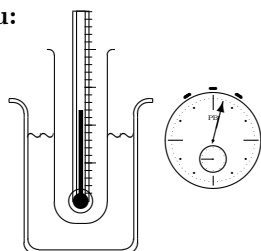
„dazwischen verrichtete Reibungsarbeit“ ist nämlich in Form von innerer Energie dem System erhalten geblieben, ist in dieser Form nur ein Summand in der rechten „Summe aller Energien“.

16.2 Änderung von Aggregatzuständen

16.2.1 Umwandlungsenergien

V 16.3 Temperaturverlauf beim Schmelzen

Aufbau:



In einem Reagenzglas befindet sich ein Thermometer, umgeben von etwas

.....
Das Reagenzglas hängt in einem Wasserbad.

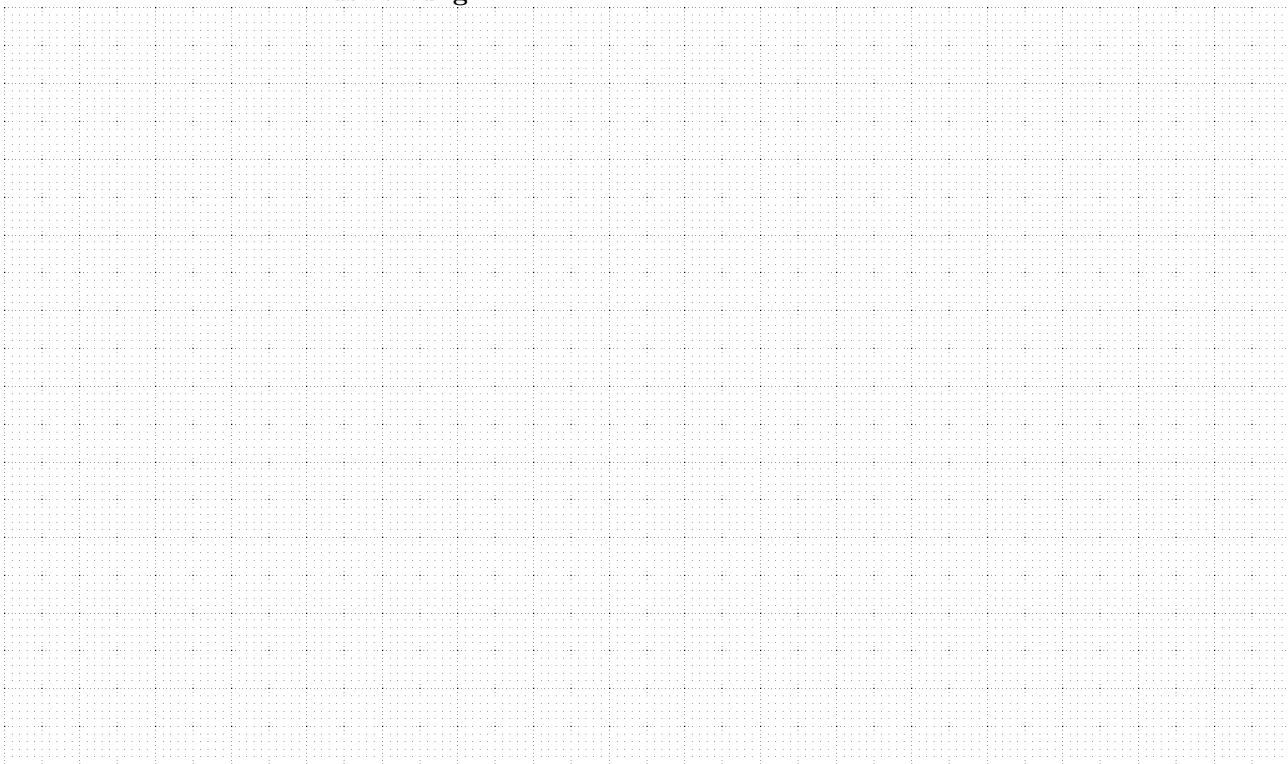
zu ergänzen:

- Brenner
- schmelzende Substanz

Durchführung: Das Wasserbad wird erhitzt.

Messwerte: / Beobachtung: (Beginn und Ende des Schmelzens)

Auswertung:



Ergebnis:

Einige Schmelz-(1) und Siede-(2)-temp. in °C:

Stoff	ϑ_1	ϑ_2	Stoff	ϑ_1	ϑ_2	Stoff	ϑ_1	ϑ_2
Wasser	0	100	Zinn	232	2270	Quecksilber	-39	357
Kalium	64	774	Kupfer	1083	2567	Brom	-7	59
Natrium	98	883	Silber	962	2212	Wasserstoff	-259	-253
Aluminium	660	2467	Gold	1064	2807	Sauerstoff	-218	-183
Eisen	1535	2750	Platin	1772	3800	Stickstoff	-210	-196

Spezifische Schmelzwärme

Wie viel Energie ist zum Schmelzen erforderlich?

?

Die benötigte Energie hängt sicher davon ab, *wie viel* einer Substanz geschmolzen wird: Um z.B. doppelt so viel Wassereis zu schmelzen benötigt man natürlich auch doppelt so viel Energie. Die zum Schmelzen erforderliche Energie hängt – wie zu vermuten – auch davon ab, *welche* Substanz geschmolzen wird. Für den nächsten Versuch präzisieren wir die Frage daher: Wie viel

Energie ist erforderlich, um 1 g Wassereis zu schmelzen?

Diese erfragte Angabe ist die **spezifische Schmelzenergie** von Wasser. In diesem Skript werden spezifische Schmelzenergien mit q bezeichnet. Da man im Zusammenhang mit innerer Energie statt von Energien auch von Wärmen spricht, heißen Schmelzenergien auch Schmelzwärmen.

Die zum Schmelzen einer Stoffportion erforderliche Energie W ist der Masse m der Stoffportion proportional. Die Proportionalitätskonstante $\frac{W}{m}$ heißt spezifische Schmelzwärme q des Stoffes.

$$W =$$

Erstarrt ein Stoff, so wird dieselbe Energiemenge, die zum Schmelzen des Stoffes erforderlich ist, wieder frei.

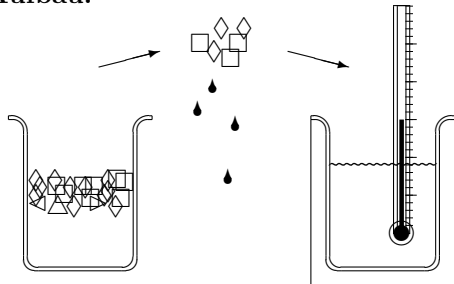
Zur Bestimmung der spezifischen Schmelz-/Erstarrungswärme von Wasser muss eine bekannte Menge Wassereis mit bekannter Energiezufuhr geschmolzen werden. Dies geht am einfachsten in einem Mischungsversuch:

zu ergänzen:

- Formelterm

V 16.4 Spezifische Schmelzwärme von Wassereis

Aufbau:

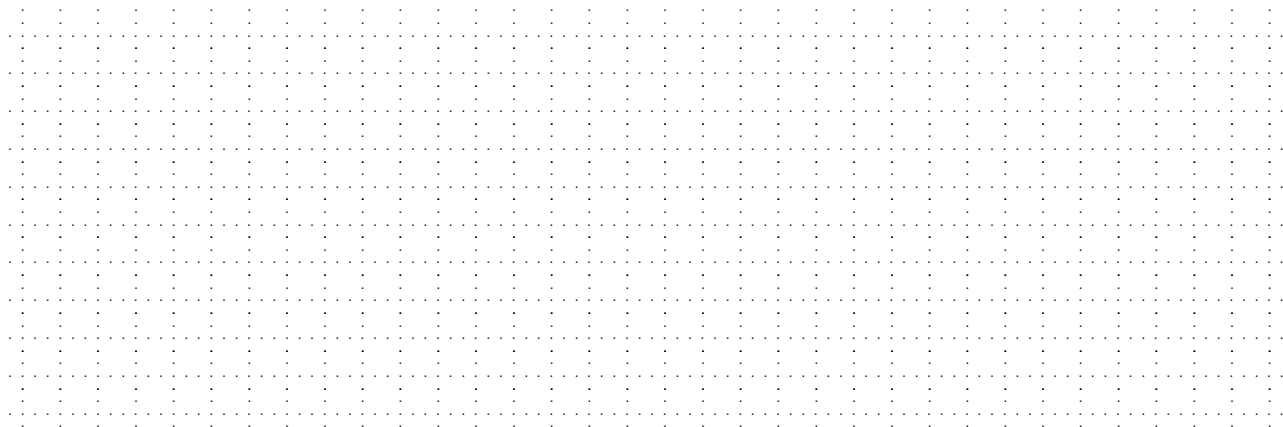


Das linke Gefäß enthält Eiswasser:
Wasser mit Wassereisstückchen.
(Kann so nur bei 0 °C dauerhaft existieren. Warum?)
Im rechten Gefäß (Kalorimeter)
befindet sich eine abgewogene Menge
Wasser.

Durchführung:

Messwerte:

Auswertung: Das relativ warme Wasser im Kalorimeter hat beim Abkühlen auf die Mischtemperatur genau die Energie abgegeben, die erforderlich war, um das Eis zu schmelzen und das Schmelzwasser auf die Mischtemperatur zu erwärmen:



Ergebnis:

Einige spezifische Schmelzwärmen:

Stoff	$q/\frac{\text{J}}{\text{g}}$
Wasser	334
Aluminium	404
Eisen	270
Nickel	300
Quecksilber	11,7
Naphtalin	147
Gold	64,5
Kupfer	207
Blei	24,7

In manchen Aufgaben weiß man nicht von Anfang an, welche Vorgänge überhaupt ablaufen werden. Beispiel:

Man mischt $m_W = 400 \text{ g}$ Wasser von $\vartheta_W = 20^\circ\text{C}$ und $m_E = 120 \text{ g}$ Wasser-Eis von $-\vartheta_E = -15^\circ\text{C}$. Berechne, welche Temperatur und welche Zusammensetzung (wie viel Gramm Wasser bzw. Eis) die Mischung am Ende hat!

Wird das Eis nur erwärmt? Oder kommt es auch zum Schmelzen von Eis? Wenn ja: Schmilzt das Eis komplett, so dass anschließend Schmelzwasser erwärmt werden kann? Oder gefriert ein Teil des anfangs flüssigen Wassers zu Eis?

In jedem Fall müsste die Energiebilanz eigentlich anders zusammengesetzt werden. Beispiel: Man schätzt, dass das Eis vollständig schmilzt und das Schmelzwasser auf die Mischtemperatur erwärmt wird. Entsprechende Bilanz: $c_W m_W (\vartheta_W - \vartheta_M) = c_E m_E \vartheta_E + q m_E + c_W m_E \vartheta_M$

Man erhält daraus die negative Mischtemperatur $\vartheta_M = -4,7^\circ\text{C}$. Erst an diesem physikalisch unmöglichen Ergebnis (flüssiges Wasser unter seinem Gefrierpunkt) erkennt man dann, dass wohl doch nicht das ganze Eis schmilzt. Die Energiebilanz hätte (mit m_S = Masse des geschmolzenen Eises) besser gelautet: $c_W m_W \vartheta_W = c_E m_E \vartheta_E + q m_S$

Solche Fehlversuche beim Bilanzieren spart man sich, wenn man z.B. nach dem folgenden Schema vorgeht: Den gesamten wahren Ablauf zerlegt man gedanklich in einfache Teilprozesse, über deren Energie-Bedarf oder Energie-Freisetzung man Buch führt. Die einzelnen Teilprozesse brauchen in der Wirklichkeit nicht genauso abzulaufen, müssen jedoch zusammen zu ausgeglichener Temperatur und einer ausgeglichenen Energiebilanz führen.

	Portion 1 400 g	Portion 2 120 g	Energie- rechnung	konto
0	Wasser 20 °C	Eis -15 °C		0 J
1	↓ → 0 °C	" "	$+c_W m_W \vartheta_W =$ +33600 J	33600 J
2	" "	↓ ← 0 °C	$-c_E m_W \vartheta_E =$ -3780 J	29820 J
3	" "	↓ ← Wasser "	$-q m_E =$ -40080 J	-10260 J
	Gesamtportion: 520 g			
4	Wasser 0 °C			-10260 J
	↓ →		$W/q =$ 30,7 g	
5	30,7 g Eis + restliches Wasser			0 J

Zeile 0: Ausgangszustand laut Aufgabentext, keine sonstigen Energien

Zeile 1: Wir haben das Wasser abgekühlt. Die dabei freigewordene Energie vermerken wir als „Guthaben“ auf unserem „Energie-Konto“

Zeile 2: Wir haben das Eis erwärmt. Dies hat Energie gekostet. Da wir aber immer noch Energie übrig haben, schmelzen wir das Eis.

Zeile 3: Wir haben das Eis geschmolzen. Alles ist jetzt flüssiges Wasser. Dieser Vorgang hat aber mehr Energie gekostet, als auf dem Konto noch war. Wir haben „Energie-Schulden“.

Zeile 4: Wir lassen gerade so viel Wasser wieder gefrieren, dass wir mit der dabei freiwerdenden Energie unsere Energie-Schulden tilgen können.

Zeile 5: Endzustand erreicht: Temperatur-Ausgleich zwischen allen beteiligten Stoffen und ausgeglichene Energiebilanz.

Natürlich haben wir diesen Zustand in Gedanken auf anderen Wegen erreicht als in der Wirklichkeit, wo niemals das ganze Eis geschmolzen war. Natürlich gibt es auch keine „Energie-Bank“, die so beliebig Kredite gewährt. Da aber mit der korrekten Gesamtbilanz der Energieerhaltungssatz erfüllt ist außerdem wie bei jedem Mischungsversuch Temperatur-Gleichheit erzielt ist, ist der errechnete Zustand auch der tatsächlich eintretende. Die Vorgehensweise war im Groben:

- Solange noch Temperatur-Unterschiede bestehen, müssen Vorgänge stattfinden, die diese Unterschiede verringern. Bei negativem Kontostand wählt man dazu einen Vorgang, der Energie liefert, ansonsten einen Vorgang, der Energie kostet. So bleibt der Kontostand möglichst nahe bei Null und damit auch die Rechnung nahe an der Wirklichkeit.
- Bei erreichter Temperatur-Gleichheit muss eventuell noch das Energiekonto durch einen dazu geeigneten Vorgang ausgeglichen werden, ohne dass wieder Temperatur-Unterschiede entstehen.

Spezifische Verdampfungswärme

In Analogie zu unseren Untersuchungen des Überganges von fest nach flüssig, fragen wir nun beim Übergang von flüssig zu gasförmig nach der **spezifischen Verdampfungswärme**.

zu ergänzen:

- Formelterm

Die zum Verdampfen einer Stoffportion erforderliche Energie W ist der Masse m der Stoffportion proportional. Die Proportionalitätskonstante $\frac{W}{m}$ heißt spezifische Verdampfungswärme r des Stoffes.

$$W =$$

Wie viel Energie ist erforderlich, um 1 g Wasser zu verdampfen?

?

Um diese Verdampfungsenergie experimentell zu bestimmen, muss im Versuch eine bekannte Menge Wasser mit bekannter Energiezufuhr verdampft werden. Es ist beim Erhitzen von Wasser allerdings nicht so (wie oft vereinfacht dargestellt), dass das ganze Wasser zuerst ohne Verdampfung auf 100 °C erwärmt wird und erst dann beim Sieden nach und nach verdampft. Tatsächlich beginnt Wasser bereits deutlich unter seiner Siedetemperatur nennenswert zu verdunsten.

Es lässt sich daher im Versuchsablauf kein Zeitpunkt ausmachen, bei dem die Erwärmung des Wassers auf Siedetemperatur aufhört und die Verdampfung beginnt. Wir können daher beide Vorgänge nur in einer Gesamt-Energiebilanz erfassen. (Ähnlich, wie im vorigen Versuch die praktisch unvermeidbare Erwärmung des Schmelzwassers mit zu bilanzieren war.)

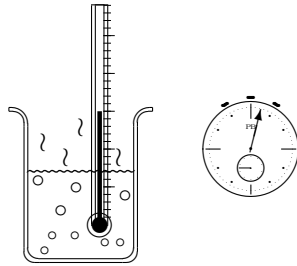
Die Verdampfungswärme von Wasser lässt sich weiterhin nicht einfach in einem Mischungsversuch bestimmen: Schon der Umstand, dass der wärmere Mischungspartner über 100 °C heiß sein müsste, erschwert uns einen Mischungsversuch.

Die Verdampfungsenergie muss dem Wasser also anders zugeführt werden: Man kann mit einer Heizquelle bekannter Heizleistung in bekannter Zeit auch bekannte Energien zuführen. Als Heizquellen kommen Kochplatten, Gasbrenner und Tauchsieder in Frage.

Sofern man die Heizleistung der Quelle noch nicht kennt, (– bei einem Tauchsieder z.B. kann sie vom Hersteller angegeben sein –) muss man sie zunächst bestimmen. Man kann dazu z.B. feststellen, in welcher Zeit die Heizquelle wie viel Wasser um wie viel erwärmt.

V 16.5 Spezifische Verdampfungswärme von Wasser

Aufbau: **Durchführung:**



Auswertung:

Ergebnis:

Einige Spezifische Verdampfungswärmen:

Stoff	$q/\frac{\text{J}}{\text{g}}$	Stoff	$q/\frac{\text{J}}{\text{g}}$	Stoff	$q/\frac{\text{J}}{\text{g}}$
Wasser	2257	Äther	377	Gold	1578
Alkohol	854	Quecksilber	285	Blei	871

16.2.2 Wechselwirkungen mit der Umgebung

Umwandlungen zwischen Aggregatzuständen haben einerseits interessante Auswirkungen *auf* die Umgebung, und andererseits unterliegen sie Einflüssen *aus* der Umgebung.

Mit dem Aggregatzustand ändert sich auch das Volumen einer Stoffportion. Unter sonst gleichen Bedingungen ist in der Regel das Volumen in festem Zustand kleiner als im flüssigen, im flüssigen kleiner als im gasförmigen. Eine wichtige Ausnahme vom ersten Teil dieser Regel ist Wassereis, dessen Volumen beim Schmelzen abnimmt.

Solche Volumenänderungen kannst Du leicht im Alltag beobachten: Beim Umgang mit Eiswürfeln, beim Kochen, beim Abkühlen des flüssigen Wachses in einer ausgeblasenen Kerze. Inwiefern?

Behindert man eine Ausdehnung, indem man den Stoff einschließt, so äußert sich das Ausdehnungsbestreben in recht großen Kräften auf das einschließende Gefäß. Dies trifft sowohl auf die Ausdehnung bei Änderung des Aggregatzustandes zu als auch auf die Ausdehnung bei Temperaturänderung. (Erinnerung: In der Regel dehnen sich Körper beim Erwärmen aus; Ausnahme: Wasser, das sich zwischen 0 °C und 4 °C beim Erwärmen zusammenzieht.)

Auch diese enormen Kräfte kannst Du im Alltag er-

leben: Vielleicht hast Du ja schon 'mal eine geschlossene Sprudelflasche o.ä. in die Gefriertruhe gelegt ;-) Gefrierendes Wasser in Felsspalten sprengt den Fels. Gleiche Wirkung kann auch gefrierendes Kühlwasser im Automotor haben.

Dagegen vergleichsweise harmlos sind die häufigen Knack-Geräusche, die Temperaturänderungen an Geräten, Dachstühlen, Fenstern etc hervorrufen: Je nach Richtung der Änderung dehnen sich verschiedene Teile unterschiedlich stark aus oder ziehen sich zusammen. Dadurch entstehen zwischen aneinander haftenden Teilen wachsende mechanische Spannungen, die sich von Zeit zu Zeit sprunghaft auflösen.

Es ist eine große Herausforderung an Techniker, Dinge so zu entwerfen, dass solche Spannungen am besten gar nicht erst entstehen, sich aber nötigenfalls rechtzeitig gefahrlos auflösen können, bevor sie zerstörerische Stärken erreichen.

Noch ein paar Beispiele zur Kraft siedenden Wassers:

Küche: Sicher weißt Du, dass in Schnellkochtöpfen beachtliche Drücke entstehen.

Auto: Damit siedendes Kühlwasser keinen Schaden an den Schläuchen des Kühlkreislaufes anrichtet, hat das Kühlwasser-Ausgleichsgefäß ein Überdruckventil im Deckel.

Die gleichen Druckkräfte, die eine sich umwandelnde Stoffportion auf ihre Umgebung, z.B. ihr Gefäß, ausübt, übt (nach dem grundlegenden Reaktionsprinzip „Kraft=Gegenkraft“) auch das Gefäß oder die Umgebung auf die Stoffportion aus. Eine mit Ausdehnung verbundene Umwandlung wird aber natürlich gehemmt, wenn die Ausdehnung gegen einen von außen einwirkenden Druck erfolgen muss.

Zum Beispiel siedet Wasser erst bei höherer Temperatur, wenn es unter höherem Druck steht. Genau dies ist die Grundidee des Schnellkochtöpfes: Der darin eingeschlossene Wasserdampf verursacht einen erhöhten Druck im ganzen Topf. Gegen diesen Druck können sich erst bei höherer Temperatur (bei heftigerer Teilchenbewegung) neue Dampfblasen in der Flüssigkeit aufblähen. Diese führen dann aber zu einer weiteren Zunahme des Druckes, wodurch die zum Sieden nötige Temperatur weiter steigt. So kann man mit über 100 °C heißem Wasser heißer und damit schneller als sonst kochen. Man erreicht in Schnellkochtöpfen etwa 115 °C.

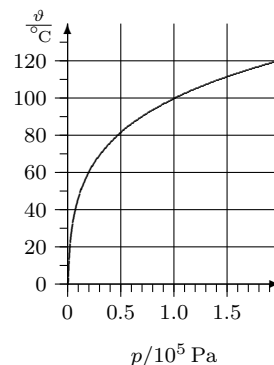
Beim Gefrieren von Wasser ist es ähnlich: Bei Wasser ist Erstarren mit einer Volumenzunahme verbunden. Erhöhter Druck bewirkt hier also, dass Wasser erst bei tieferer Temperatur gefriert als sonst. Oder umgekehrt: Wassereis schmilzt bei erhöhtem Druck bereits unter 0 °C. So gleiten Schlittschuhläufer eigentlich nicht auf dem (viel zu rauen) Eis, sondern auf der dünnen Schmelzwasserschicht, die sich bei dem hohen Druck unter den Kufen bildet.

Es gibt also eigentlich nicht *die* Schmelz- und *die* Siedetemperatur eines Stoffes. Beide hängen vom Druck ab. Ohne Angabe des Druckes sind in Tabellen immer „Normalbedingungen“ vorausgesetzt, d.h. hier: normaler Luftdruck.

Bei niedrigerem Druck siedet Wasser auch schon weit unter 100 °C. Das nebenstehende Diagramm zeigt die Siedetemperatur von Wasser in Abhängigkeit vom Druck.

Prinzip des kleinsten Zwanges

(Henry Le Chatelier, 1888): Übt man auf ein im Gleichgewicht befindliches System durch Änderung der äußeren Bedingungen einen Zwang aus, so verschiebt sich das Gleichgewicht derart, dass es dem äußeren Zwange ausweicht.

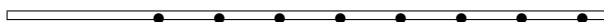


16.3 Übertragung von Wärme

Wärme kann bekanntlich von einem Körper zu einem anderen sowie innerhalb eines Körpers übertragen werden. Für eine solche Übertragung gibt es jedoch verschiedene Mechanismen.

16.3.1 Wärmeleitung

V 16.6 Wachsklumpen



zu ergänzen:

- Flamme

Aufbau:

An einem Metallstab sind einige Wachsklumpchen angedrückt.

Durchführung:

Beobachtung:

Mikroskopische Deutung: Die Atome in dem Teil des Stabes, der von der Flamme direkt erwärmt wird, schwingen recht heftig. (Hohe Temperatur $\hat{=}$ heftige Bewegung der Teilchen) Ein heftig schwingendes Teilchen bringt jedoch auch seine vorher „zahmen“ Nachbar- teilchen zu heftigen Bewegungen, sei es, weil es die Nach- barn über Anziehungskräfte mitreißt oder weil es mit ihnen zusammenstößt.

Dabei gibt das heftiger schwingende Teilchen dem zahmeren Nachbarn Energie ab, so dass sich die Energie gleichmäßiger auf alle Teilchen verteilt. Dies ist gleich- bedeutend damit, dass die Temperatur sich ausgleicht. Wird also einer Stelle dauernd Energie zugeführt, ver- teilt diese sich von dort aus: die Wärme breitet sich von dort her aus.

Geben Teilchen eines Körpers durch Stöße oder Mitreißen Energie an ihre Nachbar- teilchen weiter, so spricht man von **Wärmeleitung**.

Über diesen Mechanismus kann Wärme sich inner- halb eines Körpers verteilen und auch über Kontakt- flächen von einem Körper zu einem anderen Körper übertragen werden. Im letztgenannten Fall spricht man jedoch von **Wärmeübergang**.

Aus Erfahrung weißt Du bereits, dass verschiedene Stoffe die Wärme verschieden gut leiten. Wärme gut zu leiten heißt dabei, die Wärme schnell zu lei- ten, schnelle Temperatursausgleiche zu ermöglichen. Gu- te Wärmeleiter sind Metalle, vor allem Silber, Kupfer

und Aluminium. Schlechte sind z.B. Porzellan, Holz und Kunststoffe. Nach dem Wiedemann-Franz'schen Gesetz hängt bei vielen Metallen die gute thermische mit der guten elektrischen Leitfähigkeit zusammen.

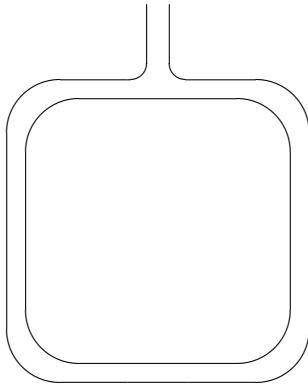
Von der unterschiedlich guten Wärmeleitfähigkeit verschiedener Stoffe kannst Du Dich recht leicht über- zeugen, indem Du z.B. mit der rechten Hand einen nicht allzu langen Metallstab, mit der linken einen vergleich- baren Glasstab in eine Kerzenflamme hältst. Welchen Stab wirst Du wohl länger halten können? ;-)

Für die Schnelligkeit einer Wärmeübertragung ist neben dem Leitungsmaterial auch Folgendes entscheidend: Es wird zwischen einem warmen und einem kalten Gebiet umso mehr Wärme pro Zeit übertragen ...
... je größer die Querschnitts-/Kontaktfläche für den Energiestrom ist,
... je größer der Temperaturunterschied ist,
... je kürzer die Entfernung zwischen den Gebieten ist.

16.3.2 Konvektion

V 16.7 Rechteckrohr

Aufbau:



Durchführung: Das Rohr wurde mit Wasser gefüllt und an einer Stelle mit etwas Farbstoff versehen. Eine der unteren Ecken des Rohres wird mit einem Bunsenbrenner erwärmt.

Beobachtung:

zu ergänzen:

- Thermometer
- Flamme
- Färbung
- Strömungspfeile

Erklärung:

Wärmeübertragung durch Strömung von erwärmtem flüssigen oder gasförmigen Material heißt **Konvektion**. Oft bilden sich Strömungen schon dadurch aus, dass das Material sich an erwärmten Stellen ausdehnt, dadurch eine geringere Dichte bekommt und aufsteigt.

Nach diesem Prinzip – jedoch unterstützt von Pumpen – arbeiten auch übliche Heizungsanlagen: Im Heizungsraum wird z.B. mit einem Öl- oder Gasbrenner Wasser in einem Heizkessel erwärmt. In einem geschlossenen Rohr-Kreislauf strömt dieses Wasser durch alle Heizkörper im Hause. Das heiße Wasser gibt durch die Heizkörper Wärme an die kühlere Raumluft ab. Die Raumluft wird wunschgemäß wärmer, das Heizungswasser kühlt ab. In den Heizkessel zurückgeflossen nimmt das Wasser wieder Wärme auf und transportiert sie erneut zu den Heizkörpern.

In größerem Maßstab findet Konvektion in der Erd-Atmosphäre statt (Winde), in Meeren (z.B. Golfstrom) und auch im flüssigen Erdinneren. In kleinerem Maßstab sind Konvektion z.B. das Wallen in einem Kochtopf oder die Umwälzungen des flüssigen Wachses rings um einen brennenden Kerzendocht.

16.3.3 Wärmestrahlung

Die wichtigste Wärmequelle für das Leben auf der Erde ist die Sonne. Ihre Wärme kann jedoch weder durch Konvektion noch durch Wärmeleitung zu uns gelangen. Grund: Zwischen Erde und Sonne gibt es ja kaum Materie, die mit ihrer Strömung (*Konvektion*) Wärme zu uns bringen könnte oder die Wärme zu uns *leiten* könnte.

Die Wärme der Sonne erreicht uns ebenso wie ihr Licht in Form von Strahlung, die keinen materiellen Überträger braucht. Die Sonne sendet Energie in Form dieser **Wärmestrahlung** aus. Jeder Körper, der diese Strahlung aufnimmt, nimmt mit ihr Energie auf und wird erwärmt.

Strahlung auszusenden bezeichnet man auch als Strahlung **emittieren**, Strahlung aufzunehmen auch als **absorbieren**.

Die Wärmestrahlung ist physikalisch mit Licht verwandt: Jedes sichtbare Licht lässt sich – z.B. mit einem Prisma – in „Regenbogen“-farbige Anteile zerlegen. Wärmestrahlung schließt sich bei einer solchen Zerlegung unmittelbar an den roten Rand des sichtbaren Bereiches an. Man nennt die Wärmestrahlung daher auch **Infrarot**.

Wärme kann ohne Beteiligung von Materie in Form von Wärmestrahlung (=Infrarot) übertragen werden.

Die Verwandtschaft der Wärmestrahlung zu Licht zeigt sich in einigen Erscheinungen, die wir auch von Licht her kennen:

V 16.8 Reflexion

Aufbau: 2 Parabolspiegel, in den Brennpunkten

links:

rechts:



zu ergänzen:

- Strahlen

Beobachtung:

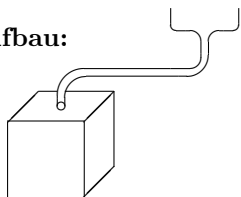
Ergebnis:

Mit sichtbarem Licht vergleichbar ist Wärmestrahlung auch in puncto Absorption. Zunächst zum Licht: Von dem Licht, das einen Körper trifft, wirft der Körper nur bestimmte Anteile zurück, den Rest absorbiert der Körper. Für jede Lichtfarbe hat jeder Körper typische Anteile, die er zurückwirft. Ein blauer Körper z.B. wirft von bläulichem Licht recht große Anteile zurück, von rötlichem Licht absorbiert er mehr. Bei Bestrahlung mit weißem Licht (– das sich aus allen Farben zusammensetzt –) sieht er gerade deshalb blau aus, weil von ihm bevorzugt bläuliches Licht zurückgeworfen wird.

Ein schwarzer Körper sieht deshalb schwarz aus, weil er fast alles Licht absorbiert und kaum etwas davon zurückwirft. Inwieweit auch für Wärmestrahlung ähnliches gilt, zeigt der folgende Versuch:

V 16.9 Leslie-Würfel, Absorption

Aufbau:



Nebstehend abgebildeter „Leslie-Würfel“ ist hohl, dicht und aus Metall. Eine seiner Außenseiten ist schwarz gefärbt, eine andere blank poliert.

zu ergänzen:

- Wärmestrahler
- Wassertropfen
- schwarz/blank

Durchführung: Der Würfel wird auf der blanken Seite bestrahlt, wieder abkühlen gelassen und dann auf der schwarzen Seite bestrahlt.

Beobachtung:

Ergebnis:

Inwieweit die Farbe eines Körpers Einfluss auf seine Fähigkeit zur Emission von Wärmestrahlung hat, zeigt der folgende Versuch:

V 16.10 Leslie-Würfel, Emission

Aufbau: Ein Leslie-Würfel ist mit heißem Wasser gefüllt, seine Außenseiten folglich auch heiß.

Durchführung: Wir halten ein empfindliches Thermometer einmal vor die blanke, einmal vor die schwarze Seite des Würfels, jeweils in gleichem Abstand zum Würfel.

Beobachtung:

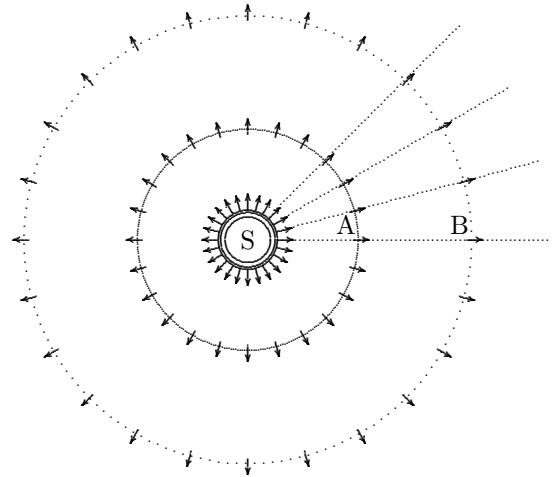
Ergebnis:

Die Sonne strahlt Energie in den Weltraum ab, u.a. auch in Form von Wärmestrahlung. Diese Energie stammt aus Kernfusionen (=Verschmelzung von Atomkernen), die im Inneren der Sonne ablaufen: Dort werden in jeder Sekunde 600 Millionen Tonnen Wasserstoff zu Helium umgewandelt. Das „Kraftwerk“ Sonne erzielt damit eine Leistung von $4 \cdot 10^{26}$ Watt: In jeder Sekunde strahlt die Sonne $4 \cdot 10^{26}$ – eine 4 mit 26 Nullen – Joule an Energie nach allen Richtungen in den Weltraum.

?

Wie viel davon bekommen wir ab?

Die Energie, die die Sonne (S) in einem Moment abstrahlt, bewegt sich radial in alle Richtungen von der Sonne weg. (Pfeilchen) Sie verteilt sich daher auf immer größer werdende Kugeloberflächen. Dieselbe Energie, die in einem Augenblick die Oberfläche der Kugel A passiert, passiert etwas später die Oberfläche von Kugel B.



Die Erde ist von der Sonne etwa 150 Millionen km entfernt. Bis zu dieser Entfernung ist die Sonnenstrahlung soweit auseinander gelaufen, dass jeder Quadratmeter, den man der Sonne zuwendet, in jeder Sekunde etwa 1400 J abbekommt. 1400 J pro Sekunde und pro m^2 lässt sich kurz angeben als eine Bestrahlungsstärke von $1400 \frac{\text{W}}{\text{m}^2}$.

Dieser Wert gilt natürlich nur, wenn der Strahlung „nichts dazwischen kommt“. Bis zur Erdoberfläche muss die Strahlung aber erst durch die Erdatmosphäre, wo bereits ein großer Teil der Energie absorbiert wird. Bestenfalls kommen am Erdboden etwa $1 \frac{\text{kW}}{\text{m}^2}$ an. Weiterhin ist zu bedenken, dass ein Stück Erdboden meist nicht optimal zur Sonne hin ausgerichtet ist – insbesondere nachts gar nicht. So erklärt sich, dass man z.B. in Deutschland mit einer mittleren Bestrahlungsstärke in der Größenordnung von $150 \frac{\text{W}}{\text{m}^2}$ auskommen muss. (über ein ganzes Jahr, Tag und Nacht, gemittelt)

Genutzt werden kann die Sonnenenergie – auch in Deutschland ökonomisch sinnvoll – auf mehrere Arten:

- Pflanzen bauen mit ihrer Hilfe energiereiche chemische Verbindungen auf. Die Wärme, die wir beim Verbrennen von Holz oder Pflanzenölen ($\rightarrow$ „Bio-Diesel“) wieder freisetzen können, stammt so aus Sonnenenergie.
- Die Solarzellen in Photovoltaik-Anlagen können Sonnenenergie direkt in elektrische Energie umsetzen.
- In Sonnenkollektoren absorbieren dunkle Flächen die Sonnenstrahlung, werden heiß und erwärmen Wasser, das man dann z.B. zur Heizung von Gebäuden verwenden kann. Typischer Aufbau: Rohrsystem mit großer schwarzer Oberfläche, von Wasser durchströmt, hinter Glas.

Verbrennungsmotoren

17.1 Dampfmaschine

Beim Verdampfen von Wasser in einem geschlossenen Gefäß entsteht erhöhter Druck, mit dem – wie bei einem Kolbenprober – ein Kolben kraftvoll aus seiner zylindrischen Hülse ausgefahren werden kann. Um eine Drehbewegung anzutreiben, wird diese geradlinige Bewegung über eine Pleuelstange auf eine Kurbelwelle übertragen. Dieses Prinzip kennst Du sicher vom Fahrrad her, wo die Abwärtsbewegung des Oberschenkels über den Unterschenkel als Pleuelstange auf die Tretkurbel übertragen wird.

Im Vergleich mit dem Fahrrad erkennst Du auch schon das nächste Problem: Der ausfahrende Kolben kann die Kurbelwelle nur eine halbe Drehung weit antreiben. Will man Antrieb während der ganzen Drehung, gibt es dafür zwei Lösungen, die man einzeln oder kombiniert einsetzen kann.

Erstens kann man einen zweiten Kolben an derselben Kurbelwelle um eine halbe Drehung versetzt angreifen lassen, der mit dem ersten Kolben abwechselnd die Kurbelwelle antreibt. Diese Lösung findest Du an Deinem Fahrrad und in jedem gängigen Kraftfahrzeug.

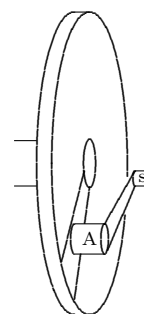
Als zweite Lösung kann man den Kolben so bauen, dass man ihn abwechselnd von zwei Seiten her unter Druck setzen kann, so dass er abwechselnd an der Kurbelwelle drückt und zieht. Diese Lösung erkennst Du in der folgenden Bildergeschichte über eine volle Umdrehung im Getriebe einer Dampflokomotive. Natürlich fehlen noch ein paar Beschriftungen und der Dampf! (Die Geschichte ist spaltenweise zu lesen.)

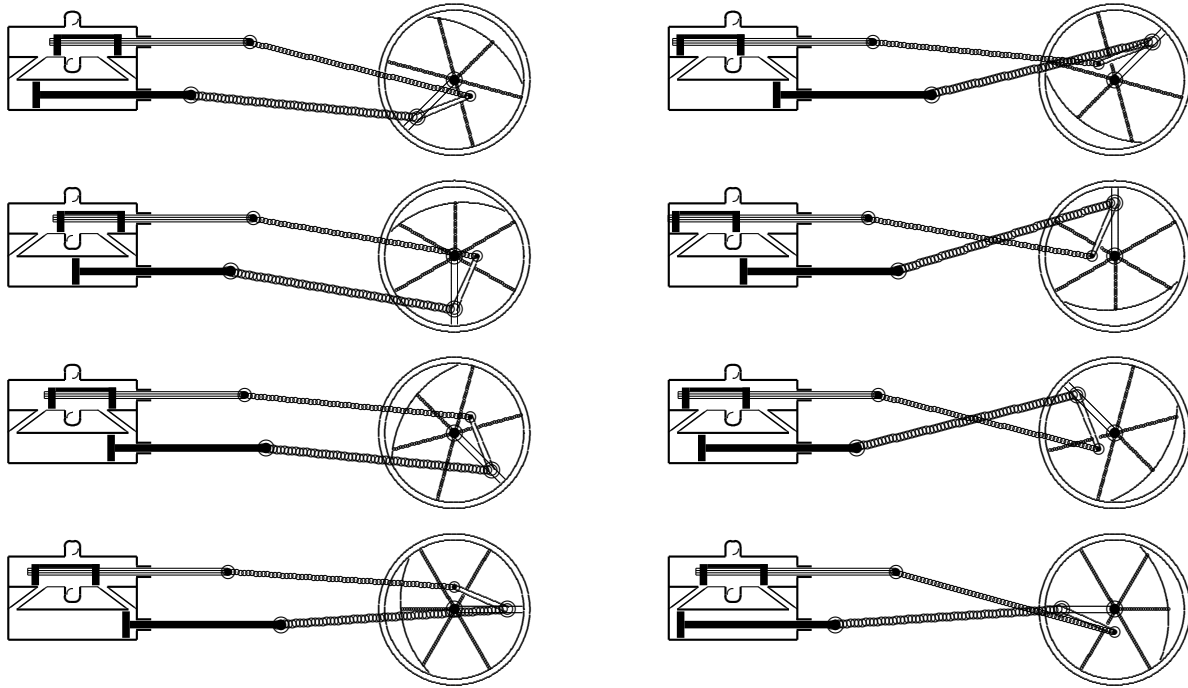
Das wahrscheinlich am schwersten vorstellbare Bauteil ist das angetriebene Rad mit den Zapfen für die Antriebs- und die Steuerpleuelstange. Nebenstehend ist das Rad alleine gezeichnet. Alles was Du hier siehst, ist ein einziges so zusammenschweißtes Stück Stahl.

Nun zu den folgenden Bildern. Im Schieberkasten links in jedem Bild erkennst Du eine obere und eine untere Kammer. Je nach Stellung des Schiebers im oberen Teil strömt der Dampf von der Einlassöffnung am oberen Bildrand durch den linken oder rechten Kanal in die untere Kammer und drückt folglich von links bzw. rechts auf den dort eingebauten Kolben. Über eine Stange wird die Kraft auf diesen Kolben jedenfalls nach rechts aus dem Schieberkasten heraus auf die Antriebs-Pleuelstange übertragen. Diese überträgt die Kraft weiter auf denjenigen Zapfen des Rades, der im großen Rad-Bild mit A beschriftet ist.

Das Rad bewegt über die am Zapfen S angeschlossene Steuer-Pleuelstange den Schieber in der oberen Kammer des Schieberkastens. Zapfen S ist gegenüber Zapfen A um etwa 90° versetzt am Rad angebracht. Durch diesen Versatz gibt der Schieber immer im richtigen Moment die richtigen Kanäle für den Dampf frei: sowohl für den frischen Dampf, der den Kolben für die nächste halbe Umdrehung antreiben soll, als auch für den verbrauchten, entspannten Dampf, der in der vorigen halben Umdrehung gearbeitet hat und nun durch die Öffnung zwischen den Kammern entweicht.

Unter der Internet- Adresse
softfrutti.de
 (→ „Materialien“)
 wurde die Datei
 DAMPFLOK.GEO
 zum Download bereit-
 gestellt und kann mit dem
 Programm *Euklid*
 (→ www.dynageo.de)
 benutzt werden.



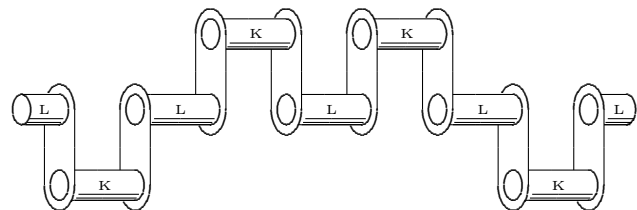


17.2 Ottomotor

Nikolaus Otto,
 ★ 14.6.1832
 in Holzhausen,
 † 26.1.1891 in Köln

Die meisten Automobile werden von einem Ottomotor angetrieben. In Ottomotoren wird jeder Kolben dadurch angetrieben, dass ein Benzin-Luft-Gemisch bei seiner Verbrennung einen hohen Druck im Brennraum oberhalb des Kolbens erzeugt. Jeder Kolben kann damit nur bei seiner Abwärtsbewegung arbeiten. Schon von daher ist klar, dass man mehr als einen Kolben verwendet. Meist sind es vier Kolben, die an der Kurbelwelle angreifen. Da jeder Kolben sich in einer zylindrischen Hülse bewegt, spricht man von einem 4-Zylinder-Motor.

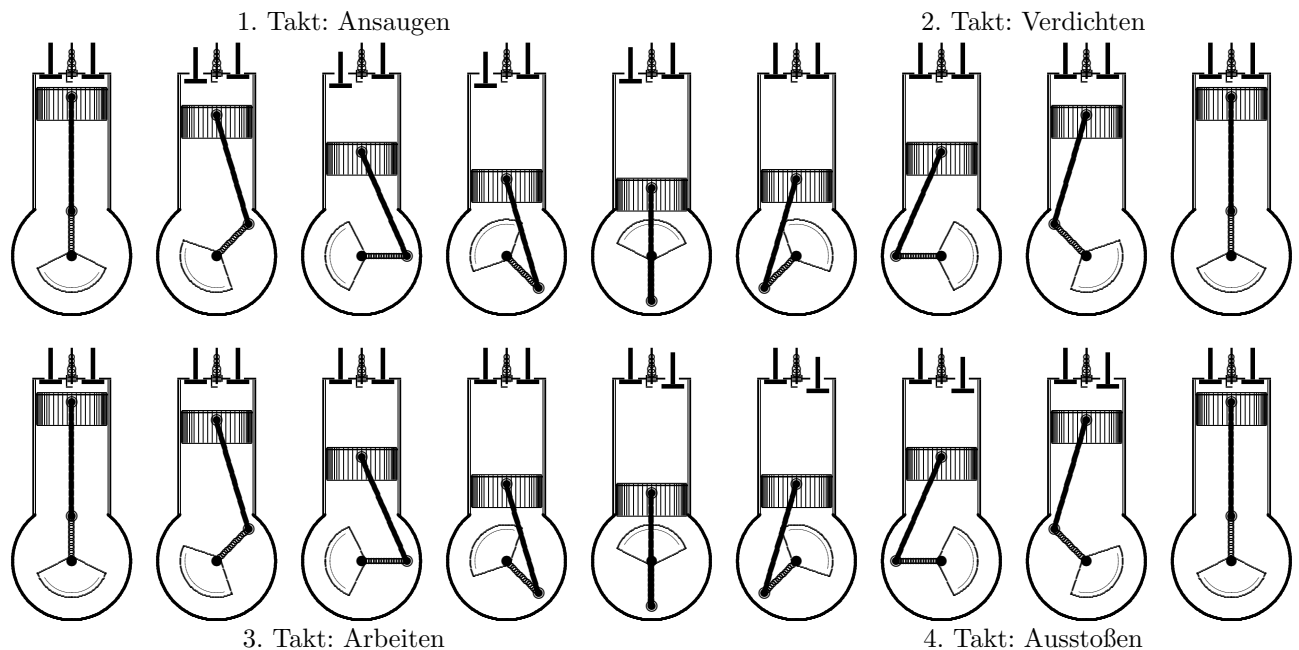
Die Kurbelwelle sieht dabei etwa so aus wie nebenstehend gezeichnet. Dieses Bild könnte zu einem schematischen Querschnitt durch den gesamten Zylinderblock ergänzt werden.



Die mit L markierten Abschnitte liegen genau entlang der Drehachse. An ihnen kann die Kurbelwelle gelagert sein. Ihre Verlängerung überträgt das Drehmoment aus dem eigentlichen Motorblock in die Kupplung, von wo die Motorleistung über das Schaltgetriebe, das Differentialgetriebe und die Antriebswellen zu den

Rädern übertragen wird.

An den mit K beschrifteten Abschnitten greifen die Pleuelstangen der Kolben an. Im Vergleich mit einem Fahrrad wären dies die Pedale, auf die die Füße einwirken.



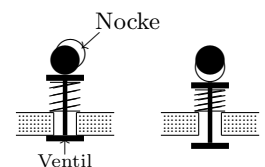
Diese Bildergeschichte zeigt die Arbeitsweise eines Zylinders. Ein kompletter Zyklus umfasst zwei volle Umdrehungen der Kurbelwelle und damit insgesamt vier Auf- und Ab-Bewegungen des Kolbens. In diesen vier Phasen laufen im Zylinder vier verschiedene Vorgänge ab, weswegen man auch von einem 4-Takt-Motor spricht.

Im ersten Takt saugt der abwärts gehende Kolben durch das geöffnete Einlassventil das Benzin-Luft-Gemisch ein. Im zweiten Takt verdichtet der aufwärts gehende Kolben bei geschlossenen Ventilen dieses Gemisch. Am Ende dieses Taktes erzeugt die Zündkerze einen elektrischen Funken, der das Gemisch entzündet. Durch die rasch fortschreitende Verbrennung steigt der

Druck im Brennraum stark an und treibt im dritten Takt den Kolben nach unten. Am Ende dieses Taktes wird das Auslassventil geöffnet, so dass die Verbrennungs-Abgase vom aufwärts gehenden Kolben ausgestoßen werden können.

Die Steuerung der Ventile und der Zündung erfolgt dabei über die Nockenwelle. Diese wird über einen Zahnriemen von der Kurbelwelle angetrieben, ähnlich wie beim Fahrrad das Hinterrad über eine Kette von der Tretkurbel. Die Nockenwelle dreht sich jedoch nur halb so schnell wie die Kurbelwelle. (Wie müssen sich also die Radien der Zahnräder beider Wellen verhalten?)

Die Nockenwelle dreht sich also während vier Takten nur einmal. An der Nockenwelle ist für jedes der zu steuernden Ventile an der richtigen Position eine seitliche Verdickung (Nocke) angebracht, mit der die Welle das Ventil gegen die Kraft einer Rückstellfeder öffnet. (Meist jedoch nicht direkt wie gezeichnet, sondern über Hebel vermittelt.) Ähnlich lassen sich zur Zündung der richtigen Zündkerze im richtigen Moment elektrische Kontakte und Spannungen herstellen.

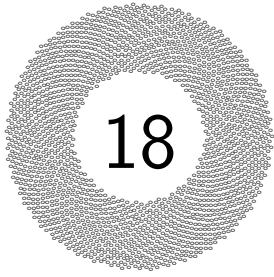


17.3 Dieselmotor

Rudolf Diesel, ★ 18.3.1858 in Paris, † 29.9.1913 im Ärmelkanal

Ähnlich wie Ottomotoren sind Dieselmotoren aufgebaut. Ein Dieselmotor saugt im ersten Arbeitstakt statt eines Benzin-Luft-Gemischs nur Luft ein. Im zweiten Takt verdichtet er diese Luft, etwa doppelt so stark wie ein Benzinmotor. Zu Beginn des dritten Taktes spritzt eine Einspritzpumpe Diesel in die verdichtete Luft. Da diese Luft bei der hohen Verdichtung schon sehr heiß ($700 - 900^\circ\text{C}$) geworden ist, entzündet sich der Diesel ohne weitere Maßnahmen von selbst. Ein Dieselmotor hat daher keine Zündkerzen.

Dieselmotoren haben gegenüber Ottomotoren den Vorteil, dass sie geringere Anforderungen an den Treibstoff stellen. Diesel ist bis auf steuerliche Aspekte mit Heizöl identisch. Statt Diesel könnten in vielen Motoren aus technischer Sicht z.B. auch pflanzliche Öle verwendet werden. Ein weiterer Vorteil von Dieselmotoren ist ihr höherer Wirkungsgrad. Während ein PKW mit Benzin-Motor etwa 25% der Energie des Treibstoffes nutzt, kommt ein Diesel-PKW bis über 40%. Große Dieselmotoren wie in Schiffen kommen auf über 50%.



18.1 Erste Erfahrungen

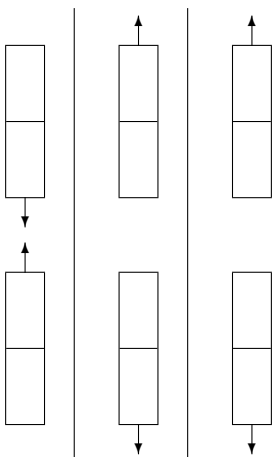
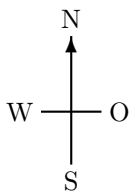
Die ersten Erfahrungen, die Menschen mit Magneten gesammelt haben, sind bereits sehr alt: Schon vor über 2500 Jahren soll man bei der Stadt **Magnesia** (Kleinasien) Steine gefunden haben, die untereinander und auf Eisenstücke Kräfte ausüben. Nach dem Fundort wurden diese Steine **Magnete** genannt. Wir untersuchen nun den Magnetismus mit Hilfe technisch hergestellter Magnete. Solche gibt es in verschiedenen Formen, z.B. als Stabmagnete, Hufeisenmagnete oder Magnetnadeln.

18.1.1 Magnetpole

Ein Magnet zieht Gegenstände aus Eisen an, z.B. Nägel und Eisenfeilspäne. Die stärkste Anziehung beobachten wir an den beiden Enden eines Magneten. Diese Stellen stärkster Kraftwirkung nennt man die **Magnetpole**.

Eine Magnetnadel, die wir frei drehbar aufstellen, stellt sich so ein, dass einer der Pole nach Norden zeigt, der andere nach Süden. Auch an anderen Standorten stellt sich immer wieder derselbe Pol nach Norden ein. Da sich die beiden Pole also immer in dieser Weise unterschiedlich verhalten, können wir die beiden Pole unterscheiden. Man nennt den Pol, der nach Norden zeigt, den **Nordpol**, den anderen den **Südpol**. In der Regel wird der Nordpol eines Magneten rot, der Südpol grün gefärbt.

Diese Eigenschaft von Magneten nutzt man schon seit langem in **Kompassen** aus. Ein Kompass ist im Wesentlichen eine frei drehbare Magnetnadel, mit der sich die Nord-Süd-Richtung – und damit alle Himmelsrichtungen – feststellen lassen.



Wenn wir mit zwei Stabmagneten hantieren, spüren wir, dass die beiden Magnete Kräfte aufeinander ausüben. Je nachdem, mit welchen Polen man die beiden Magnete einander nähert, stoßen sich die Magnete ab oder ziehen sich an. Die nebenstehende Abbildung zeigt die drei verschiedenen möglichen Kombinationen von Polen und die Richtungen der dabei auftretenden Kräfte. Du musst allerdings noch die Stabmagnete passend färben.

Der Nordpol eines Magneten und der Südpol eines anderen Magneten ziehen sich gegenseitig an. Zwei gleichnamige Pole stoßen sich ab.

Auch die Erde ist ein riesiger Magnet mit Nord- und Südpol. Im Laufe der Jahrhunderte hat sich allerdings eine irreführende Bezeichnung festgesetzt: Der geografische Nordpol, zu dem ja Nordpole von Kompassnadeln angezogen werden, muss demnach magnetisch gesehen der Südpol des Magneten „Erde“ sein: Der magnetische Südpol der Erde liegt (grob betrachtet etwa) am geografischen Nordpol und umgekehrt.

18.1.2 Dipol

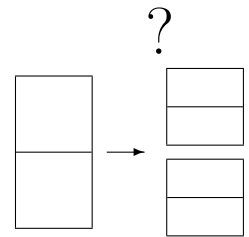
Alle bisher von uns benutzten Magnete haben einen Nord- und einen Südpol.

Gibt es nicht auch einzelne Nordpole und einzelne Südpole?

Man kann versuchen, einzelne Pole herzustellen, indem man einen Stabmagneten in der Mitte zwischen seinen Polen zertrennt. Man wird allerdings enttäuscht feststellen: Jedes Teilstück hat wiederum einen Nord- und einen Südpol. An jedem Teilstück entsteht an der Trennstelle ein Pol der vermeintlich abgetrennten Sorte. Auch weiteres Zerteilen bringt immer nur Stücke mit jeweils Nord- und Südpol hervor.

Der umgekehrte Weg funktioniert ebenfalls: Setzt man zwei Stabmagnete mit ihren ungleichnamigen Polen aneinander, erhält man keinen Magneten mit vier Polen, sondern wieder einen Magneten mit *einem* Nord- und *einem* Südpol. Die beiden zusammengefügte Pole in der Mitte verschwinden.

Es gilt ganz allgemein: Einzelne Magnetpole gibt es nicht, es gehören immer zwei Pole, Nord und Süd, zusammen. Man sagt:



zu ergänzen:

- Farben der Magnethälften

Es gibt keine magnetischen Monopole, (griech.: *monos*=einzig)
Magnete kommen als Dipole vor. (*duo*=zwei)

Magnete und Föne: Auch im Alltag gibt es Gegenstände mit ähnlicher Dipol-Eigenschaft, z.B. Föne: Jeder Fön hat einen Pol, äh, pardon: eine Seite, aus der er Luft hinausbläst, und eine Seite, in die die Luft hineinströmt. Man könnte bei einem Fön von einem Einström- und einem Ausströmpol sprechen.

Würde man zwei Föne so miteinander verbinden, dass der eine direkt in den anderen hineinbläst, hätte auch dieser Doppelfön wieder nur eine Seite für einströmende und eine Seite für ausgeblasene Luft. (Nicht ausprobieren!)

Geht man von einem solchen Doppelfön aus und trennt diesen, erkennt man, dass auch dabei an der Trennstelle ein Einströmpol und ein Ausströmpol entstehen.

18.2 Innerer Aufbau von Magneten

18.2.1 Elementarmagnete

Lassen sich Magnete beliebig zerteilen?

?

Um dies zu klären, müssen wir vorher klären: Lässt sich überhaupt ein Körper beliebig oft weiter zerteilen?

In wie kleine Teile kannst Du ein Streichholz zerbrechen? Brich dazu ein Streichholz in der Mitte durch, dann eine der Hälften wieder und so fort. Wie oft kannst Du halbieren?

Es ging gar nicht so oft, stimmt's? Mit Hilfsmitteln (Zangen, Messer) ginge es sicher doch noch ein paar Male mehr. Beliebiger oft ginge es aber sicher auch damit nicht.

Man geht heute davon aus, dass man auch mit bestmöglichen Hilfsmitteln einen Körper nicht beliebig klein zerteilen kann. Schon Demokrit (um 460 v. Chr. geboren) ging davon aus, dass jeder Körper aus kleinsten, nicht mehr weiter teilbaren Teilchen unterschiedlicher Art zusammengesetzt sei. Er nannte diese Teilchen **Atome**, was vom griechischen Wort für *unteilbar* kommt.

Die Teilchen, die wir heute Atome nennen, können wir zwar spalten, aber auch nach unserem Verständnis besteht die Welt aus unteilbaren, allerdings noch viel kleineren Teilchen.

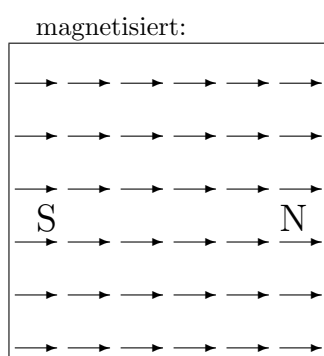
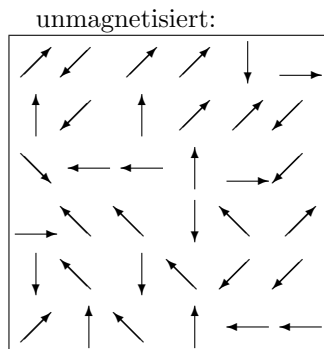
Wir nehmen an, dass es auch kleinste, nicht mehr weiter teilbare Magnete gibt. Da alle Magnete aus diesen winzigsten Magnetchen aufgebaut sind, nennen wir diese Winzlinge **Elementarmagnete**.

18.2.2 Magnetisierung

Mit diesen Elementarmagneten lässt sich nun auch erklären, warum ein Magnet gewöhnliche Eisenstücke anzieht. In einem Eisenstück gibt es ebenfalls Elementarmagnete. Diese sind jedoch bei einem unmagnetischen Eisenstück ungeordnet, so dass sich ihre Wirkungen auf die Außenwelt gegenseitig aufheben.

Nun bringen wir einen Magneten z.B. mit seinem Nordpol in die Nähe des Eisenstückes. Die Südpole der Elementarmagnete werden nun zum Nordpol des Magneten hin ausgerichtet. Damit erhalten alle Elementarmagnete im Eisenstück die gleiche Ausrichtung, und ihre Wirkungen verstärken sich. Das Eisenstück als Ganzes wirkt nun wie ein Magnet, ist magnetisiert worden.

So kommt es auch immer zu einer Anziehung zwischen dem angenäherten Magneten und dem zunächst unmagnetischen Eisenstück. Die nebenstehende Abbildung stellt die Elementarmagnete in einem Stück Eisen als Magnetnadel-Pfeilchen dar. Die Pfeilspitzen sind die Nordpole.



Ebenso wie Körper aus Eisen lassen sich Körper aus den Metallen Nickel und Kobalt magnetisieren. Stoffe mit ähnlich guter Magnetisierbarkeit wie Eisen nennt man nach lateinisch *ferrum*=Eisen **ferromagnetisch**.

„Metall“ ist der Oberbegriff für Stoffe mit gewissen physikalischen Eigenschaften, die Du in den nächsten Jahren kennenlernen wirst. Metallische Stoffe sind z.B. Eisen, Gold, Silber, Kupfer, Blei, Zinn, Aluminium und viele andere, eben auch Nickel und Kobalt.

Man kann Metalle schmelzen und dann meist auch recht gut miteinander mischen. Solche Mischungen nennt man **Legierungen**. Je nach genauer Zusammensetzung lässt sich eine Eisen-, Nickel- oder Kobalt-haltige Legierung mehr oder weniger leicht und dauerhaft magnetisieren. Die gewöhnlichen Eisenstücke, die wir bisher verwendet haben, lassen sich leicht magnetisieren, verlieren aber auch ihre Magnetisierung schnell wieder, wenn kein Magnet mehr in der Nähe ist. Solche Legierungen nennt man **weich-magnetisch**.

Die von uns verwendeten Stabmagnete, Magnetnadeln u.ä. behalten ihre Magnetisierung lange. Ihre Elementarmagnete geraten erst nach langer Zeit wieder in Unordnung, oder wenn man den Magneten stark erschüttert oder erhitzt. Es ist allerdings aus denselben Gründen auch aufwändiger, solche Dauermagnete herzustellen: Ihr Bestandmaterial ist **hart-magnetisch**.

18.3 Magnetfelder

18.3.1 Feldlinienbilder

Hält man eine beliebige Kompass-Nadel an eine Stelle in der Nähe anderer Magnete, richtet sie sich in einer bestimmten Richtung aus. Diese ausrichtende *Wirkung* der anderen Magnete auf die Nadel können wir sogar beschreiben, ohne weiteres über die anderen Magnete zu wissen: wo sie sich befinden, wie viele es sind, wie stark sie sind, welche Form sie haben usw.

Die Wirkung auf eine Magnetnadel hängt natürlich auch davon ab, *wo* man die Nadel hält. Im Allgemeinen wird man an verschiedenen Stellen in der Umgebung von Magneten verschiedene Ausrichtungen der Nadel feststellen.

Man stellt sich vor, dass Magnete den Raum um sie herum so beeinflussen, dass andere Magnete dann in diesem Raum Kräfte erfahren. Dies formuliert man mit dem Begriff **Magnetfeld** oder **magnetisches Feld**:

Wirken in einem Raumgebiet auf eine Kompass-Nadel (aufgrund ihrer Magnetisierung) Kräfte, so sagt man, es herrsche in diesem Gebiet ein Magnetfeld.

Man legt weiter fest: Die *Richtung* eines Magnetfeldes an einer Stelle ist die Richtung, in die sich der Nordpol einer frei drehbar dort aufgestellten Magnetnadel einstellen würde. Die *Stärke* des Feldes entspricht der Stärke der Kräfte auf die Magnetnadel.

V 18.1 Aussondieren eines Magnetfeldes

Aufbau: Ein Stabmagnet liegt auf dem Tisch.



Durchführung: Wir stellen drehbar gelagerte Magnetnadeln an verschiedene Stellen um den Stabmagneten.

Beobachtung: Die Magnetnadeln stellen sich so ein wie abgebildet.

zu ergänzen:

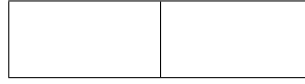
- Magnetnadeln
- Färbung der Pole

In obigem Bild fällt auf, dass die im Feld ausgerichteten Nadeln ein einfaches Muster von Linien zwischen den Polen des Stabmagneten andeuten. Noch deutlicher werden diese Linien, wenn man statt der Magnetnadeln Eisenfeilspäne verwendet. Diese werden im Magnetfeld des Stabmagneten magnetisiert, stellen sich in Feldrichtung ein und reihen sich zu Ketten aneinander. Diese Ketten bilden gut sichtbare Linienabschnitte.

Mit diesen Linien lässt sich ein Magnetfeld gut darstellen. Man nennt sie **Feldlinien**. Damit auch die Richtung des Feldes dargestellt werden kann, gibt man den Feldlinien eine Orientierung: Die Feldlinien zeigen immer die Richtung an, in die sich der Nordpol einer Magnetnadel einstellen würde. Die Feldlinien eines Magneten verlaufen daher immer von seinem Nord- zu seinem Südpol.

Beachte: Feldlinien sind gedachte Linien. Durch jede Stelle in einem Magnetfeld ließe sich eine Feldlinie zeichnen. Dies würde allerdings zu einer vollständig eingefärbten Zeichnung führen. Um überhaupt noch einzelne Linien erkennen zu können, werden daher nur einige ausgewählte Feldlinien gezeichnet. Bei der Auswahl der Feldlinien ist lediglich zu beachten: In Bereichen, in denen das Feld stärker ist, zeichnet man dichter liegende Feldlinien.

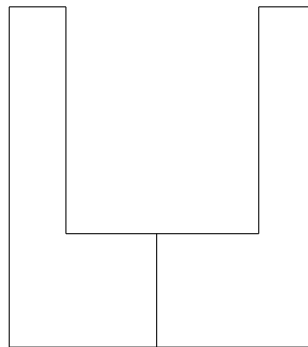
Feldlinien eines Stabmagneten:



zu ergänzen:

- Feldlinien
- Färbung der Pole

Feldlinien eines Hufeisenmagneten:



?

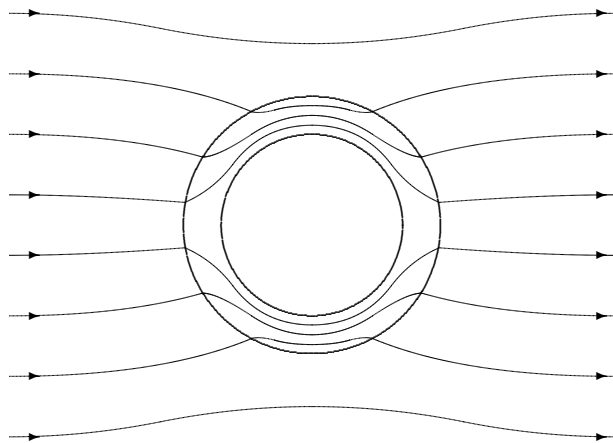
Können Feldlinien sich kreuzen?

Nehmen wir an, zwei Feldlinien würden sich kreuzen. Das hieße: Eine Magnetnadel am Kreuzungspunkt würde sich dann in zwei verschiedenen Richtungen gleichzeitig einstellen. Dies geht natürlich nicht. Also: Feldlinien kreuzen sich nie!

18.3.2 Abschirmung

Bringt man zum Beispiel ein Stück Weicheisen in ein Magnetfeld M1, so richten sich die Elementarmagnete im Eisen bekanntlich diesem Feld entsprechend aus. Wir erkennen daran schon einmal, dass Magnetfelder nicht nur durch die Luft, sondern auch durch andere Stoffe hindurch verlaufen können.

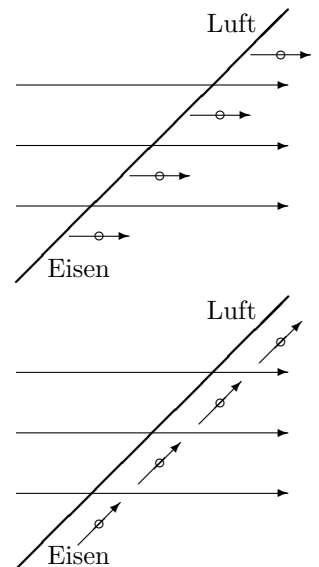
Feldlinien verlaufen sogar bevorzugt durch ferromagnetische Stoffe: Unser Eisenstück wird ja durch die Ausrichtung seiner Elementarmagnete im Feld M1 selbst im Ganzen zu einem Magneten, dessen Feld M2 das Feld M1 verstärkt. Dem stärkeren Gesamtfeld im und am Eisenstück entsprechen dabei natürlich dichter liegende Feldlinien. Von allen Feldlinien verläuft also ein großer Anteil durch das Eisen.

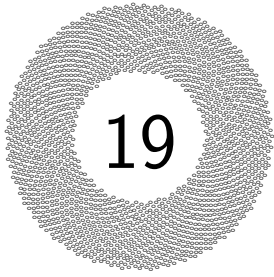


Diese „Vorliebe“ der Feldlinien für ferromagnetische Stoffe kann man zur Abschirmung von Magnetfeldern nutzen.
← Dieser Querschnitt zeigt, warum Magnetfelder kaum in das Innere z.B. einer eisernen Hohlkugel eindringen.

Beim Übergang von Luft in Eisen knicken die Feldlinien so ab, dass sie möglichst entlang der Oberfläche des Eisens verlaufen. Dies kann man sich mit der folgenden Überlegung plausibel machen. Jeder einzelne Elementarmagnet im Eisen würde sich zunächst so ausrichten, wie das von außen eintretende Feld ursprünglich verläuft. (→ oberes der Bilder nebenan)

An der Oberfläche des Eisens richten sich benachbarte Elementarmagnete dann jedoch gegenseitig aneinander aus. (→ unteres Bild) Beispiel: Der Südpol vom dritten eingezeichneten Elementarmagneten und der Nordpol des zweiten ziehen sich gegenseitig an, so dass die beiden Magnete sich entlang der Grenzfläche ausrichten. So verläuft die Magnetisierung des Eisens bevorzugt entlang seiner Oberflächen.





Vorstellung von elektrischem Strom

19.1 Erste Erfahrungen

Erfahrungen mit Elektrizität haben Menschen eigentlich schon sehr lange. Besonders eindrucksvoll konnte man schon immer Elektrizität bei Gewitter-Blitzen erleben, jedoch nur aus der Ferne. Einen solchen Blitz hautnah zu erleben ist nämlich häufig die letzte Erfahrung im Leben.

Aus dem Alltag wissen wir, dass elektrische Geräte nur funktionieren, wenn sie – wie man sagt – „Strom kriegen“. Als **Stromquellen** kennen wir Steckdosen,

Fahrrad-Dynamos, Batterien und Akkus. Im Physikunterricht benutzen wir außerdem sogenannte Netzgeräte als Stromquellen.

Strom aus Steckdosen ist bei unsachgemäßem Umgang lebensgefährlich. Unternimm daher keinesfalls auf eigene Faust Versuche mit Steckdosen! Versuche mit Taschenlampen-Batterien (Monozellen, Babyzellen, Mignonzellen, ...) sind ungefährlich.



V 19.1 Betrieb eines Lämpchens mit einer Batterie

Durchführung: Elektrische Geräte werden bekanntlich mit **Kabeln** an Stromquellen angeschlossen. Wir versuchen, mit einer (Flach-)Batterie und Kabeln ein Lämpchen zum Leuchten zu bringen. Es gelingt mit nebenstehender Schaltung:

Ergebnis: Ein einziges Kabel genügt nicht, um das Lämpchen leuchten zu lassen. Ein Anschluss der Batterie muss mit dem Lampengewinde verbunden werden, der andere mit der Metallkuppe unten am Lämpchen. Erst dann leuchtet es.



zu ergänzen:

- Kabel

Dieses Versuchsergebnis mag zunächst überraschen: Alle Haushaltsgeräte schließt man doch mit nur *einem* Kabel an die Steckdose an! Schneidet man bei einem nicht mehr gebrauchten – und vor allem von der Steckdose getrennten! – Gerät das Kabel durch, klärt sich der Widerspruch auf: Das aufgeschnittene Kabel enthält 2 oder sogar 3 Teilkabel. Man nennt solche Teilkabel **Adern**, das Gesamtkabel entsprechend 2- bzw. 3-adrig. Eine einzelne Ader oder ein einadriges Kabel nennt man auch eine **Leitung**. Auch die Steckdose hat zwei Löcher, in die jeweils ein Stift des Gerätesteckers gesteckt wird. Es wundert uns nun auch nicht, dass zu jedem Stift eine eigene Ader gehört. Zu der eventuell vorhandenen dritten Ader kommen wir später.

Jede Stromquelle und jedes elektrische Gerät hat immer zwei Anschlüsse. Zum Betrieb eines Gerätes muss jeder Anschluss des Gerätes mit jeweils einem Anschluss der Stromquelle durch eine Leitung verbunden werden.

Verlegt man die beiden Leitungen so, dass sie sich nicht kreuzen, ähnelt eine so hergestellte Schaltung einem Rundweg: von der Quelle über die eine Leitung zum Gerät und über die andere Leitung zurück zur Quelle. Man nennt eine solche Schaltung deshalb einen **Stromkreis**.

19.2 Elektrischer Strom

Der Vergleich zwischen einem Stromkreis und einem Rundweg lässt sich über die Ähnlichkeit in der Form noch weiter fortführen: Auf einem Rundweg läuft man im Kreise herum, z.B. als 5000m-Läufer auf der Laufbahn eines Stadions. Man entwickelte demnach die Vorstellung, dass auch durch einen Stromkreis etwas rundläuft, anders gesagt: durch die Leitungen strömt. Und deshalb nennt man überhaupt Strom „Strom“. Dem strömenden Etwas hat man den Namen **elektrische Ladung** gegeben. Die beiden Anschlüsse einer Stromquelle sind demnach Stellen, an denen elektrische Ladung aus der Quelle *heraus*- bzw. in die Quelle *hinein*strömt. Man nennt die beiden Anschlüsse **Pole**.

Wir kennen das Wort „Pol“ vor allem in den Zusammensetzungen Nordpol und Südpol, im geografischen oder magnetischen Sinne benutzt. In jedem Falle handelt es sich um Stellen, an denen etwas durch eine Fläche hindurchtritt: Die Erdachse tritt als gedachte Gerade an den Polen durch die Erdoberfläche, und Magnetfeldlinien treten vor allem an den Polen durch die Oberfläche eines Magneten.

Wir können den elektrischen Stromfluss durch ein Kabel mit unseren Sinnesorganen nicht direkt wahrnehmen: nicht sehen, nicht hören, nicht erfühlen, schmecken oder riechen. Wir können im Folgenden lediglich einige Versuche durchführen, die mit der Vorstellung eines fließenden Etwas in Einklang stehen.

Wenn die Strom-Vorstellung zutrifft, erwarten wir feststellbare Unterschiede zwischen den beiden Polen einer Batterie. Schließlich strömt Ladung aus dem einen Pol *heraus*, in den anderen *hinein*.

Eine besondere Lampenart, die **Glimmlampe**, verhält sich im Gegensatz zu normalen Glühlampen tatsächlich unterschiedlich, je nachdem, mit welchen Polen man ihre beiden Anschlüsse verbindet. Eine Glimmlampe besteht in der einfachsten Ausführung aus einem Glasröhrchen mit Metallkappen, von denen aus jeweils ein Draht in das Innere des Röhrchens ragt. Die Drähte berühren sich nicht. Das Röhrchen ist mit Neon-Gas gefüllt. In etwas aufwändigerer Version tragen die Drähte dreieckige Plättchen und die Lampe hat einen Fuß zum Aufstellen. In jeder Ausführung nennt man die in das Innere hineinragenden Teile **Elektroden**.

V 19.2 Glimmlampe an einem Netzgerät

Durchführung: Eine Glimmlampe wird an ein Netzgerät angeschlossen wie nebenstehend gezeichnet. Dann wird das Netzgerät eingeschaltet.

Beobachtung: Um die eine Elektrode herum

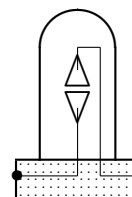
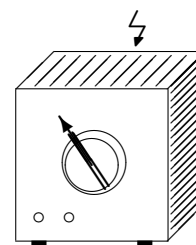
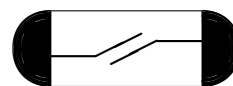
Durchführung: Das Netzgerät wird abgeschaltet, die Anschlüsse werden vertauscht, das Netzgerät wird wieder eingeschaltet.

Beobachtung: Es leuchtet nun nur um die andere Elektrode.

Ergebnis: Die Anschlüsse des Netzgerätes haben unterschiedliche elektrische Eigenschaften: Nur um die Elektrode an dem einen Anschluss leuchtet die Glimmlampe, um die Elektrode an dem anderen Anschluss nicht.

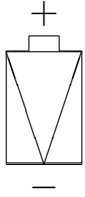
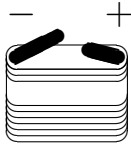
Man nennt den Pol einer Stromquelle, an den die leuchtende Seite angeschlossen ist, den **Minus-Pol**. Den anderen Pol nennt man den **Plus-Pol**.

Eine Batterie bringt eine Glimmlampe nicht zum Leuchten: sie ist zu schwach. Es gibt jedoch sogenannte **Leuchtdioden**, die schon von einer Batterie zum Leuchten gebracht werden können. Leuchtdioden werden in vielen elektrischen Geräten zur Anzeige von Zuständen des Gerätes benutzt: Bei vielen Geräten ist das rote oder grüne Lämpchen, das beim Einschalten des Gerätes angeht, eine Leuchtdiode.



zu ergänzen:

- Leitungen
- leuchtender Bereich



Eine Leuchtdiode leuchtet jedoch nur, wenn ihre Anschlüsse mit den richtigen Polen der Stromquelle verbunden werden. Wir können daher auch die Pole einer Batterie feststellen:

- Bei einer Flachbatterie ist der kürzere Blechstreifen der Pluspol, der längere der Minuspol.
- Bei einer Taschenlampenbatterie ist der Boden der Minuspol, die hervorstehende Stelle oben der Pluspol.

Im Wort *Diode* findet sich wie bei *Dipol* die Vorsilbe di- für *zwei*. -ode kommt von dem griechischen Wort *hodos* für *Weg*. Es steht hier – wie auch in *Elektrode* – für Strom-Wege, für Anschlüsse. Dioden wurden *Dioden* genannt, weil sie zwei unterscheidbare Anschlüsse haben.

Unsere bisherigen Beobachtungen decken sich mit der Vorstellung eines Stromes von einem Pole der Quelle zu dem anderen. Wir fassen zusammen:

- Ein elektrischer Strom kann erst fließen, wenn über Leitungen und Geräte eine durchgehende Verbindung zwischen den Polen der Stromquelle besteht.
- Die ganze Schaltung nennt man einen geschlossenen Stromkreis.
- Die Stromquelle „pumpt“ elektrische Ladung in diesem Kreis herum. Nebenbei: Stromquellen pumpen unterschiedlich stark. Die Pump-Stärke bezeichnet man als **Spannung** und misst sie in der Einheit Volt. Doch dazu später mehr.
- Wird der Stromkreis irgendwo unterbrochen, hört der Stromfluss sofort auf.

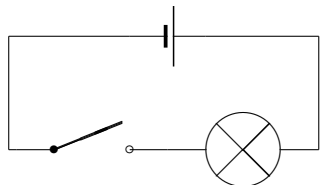
Zum Öffnen(=Unterbrechen) und Schließen eines Stromkreises verwendet man **Schalter**. Stromkreise kann man zwar auch dadurch öffnen, dass man darin irgendwelche Stecker herauszieht (oder Leitungen durchschneidet), aber es ist sicherer und übersichtlicher, Schalter zu verwenden: Es liegen nach dem Öffnen keine losen Leitungsenden herum. Natürlich werden im Schalter selbst zum Schließen und Öffnen auch nur bewegliche Leitungsenden aneinander gedrückt bzw. voneinander getrennt.

19.3 Schaltsymbole und Schaltpläne

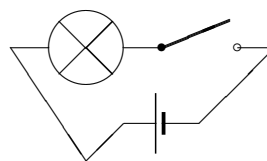
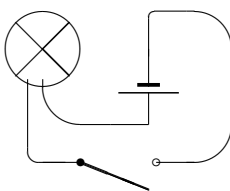
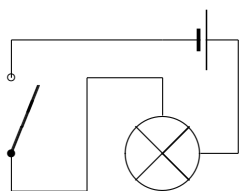
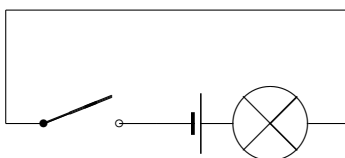
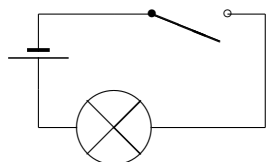
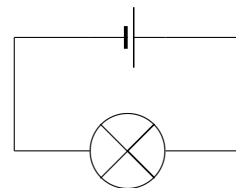
Um elektrische Schaltungen eindeutig, einfach und übersichtlich darzustellen gibt es für elektrische Bauteile spezielle Symbole. Anstatt die Bauteile und Leitungen einer Schaltung möglichst wirklichkeitsgetreu zu malen, zeichnet man einfach die entsprechenden Symbole für die Bauteile und verbindet sie, den Leitungen entsprechend, mit Linien.

	Leitung
	Glühlampe oder anderes mit Strom betriebenes Gerät, auch „Verbraucher“ genannt
	Batterie; kurzer Strich ist Minuspol, langer Strich ist Pluspol (Merkhilfe: aus dem langen Strich kann man zwei kurze machen, aus denen ein + besteht)
	feste Kontaktstelle
	leicht lösbare Kontaktstelle
	offener Schalter
	geschlossener Schalter
	Wechselschalter; verbindet die feste Kontaktstelle immer mit einer der beiden leicht lösbaren (keine Mittelstellung)

Einen einfachen Stromkreis aus Batterie und Glühlämpchen kann man demnach darstellen wie am Rand zu sehen. Baut man noch einen Schalter zwischen Minuspol der Batterie und Lämpchen ein, ergibt sich folgender Schaltplan:



Die genaue Position der Symbole in der Zeichnung spielt dabei keine Rolle. Entscheidend ist allein, welcher Bauteil-Anschluss mit welchem anderen Anschluss verbunden ist. Die folgenden Schaltpläne sind daher elektrisch gesehen völlig gleichwertig zu dem vorigen:



Natürlich sollte man Schaltpläne möglichst übersichtlich zeichnen. Dazu sollte man nur waagrecht oder senkrecht verlaufende Leitungen und möglichst wenig Ecken zeichnen.

19.4 Elektronen

Elektrischen Strom haben wir schon auf Seite 167 als Fluss elektrischer Ladung beschrieben, ohne eine Vorstellung davon zu haben, was eigentlich dabei fließt.

Jeder Körper besteht aus vielen winzigen Teilchen verschiedener Sorten mit verschiedenen Eigenschaften. Manche dieser Eigenschaften sind Dir auch aus dem Alltag geläufig, zum Beispiel der Durchmesser oder die Masse eines Teilchens. Alltäglichen Gegenständen wie z.B. Fußbällen schreibt man daneben weitere Eigenschaften zu wie Farbe, Rauigkeit, Härte, Form, Volumen.

Jede Eigenschaft (in den Naturwissenschaften auch *Größe* genannt) wird mit Hilfe typischer Worte oder mit Zahlenwerten und Einheiten beschrieben: Die Farbe beschreiben wir z.B. mit den Worten rot, orange, gelb, grün, blau und violett. Bezüglich der Härte unterscheiden wir harte und weiche Körper. Ein Volumen geben wir z.B. mit 4711 cm^3 an. So beschreiben wir, welche Erfahrungen und Beobachtungen wir mit dem Gegenstand machen: wie stark er seine Unterlage belastet, wie er aussieht, wie er sich anfühlt, wie er sich verhält, wenn

man ihn quetscht, wie viel Platz er braucht usw.

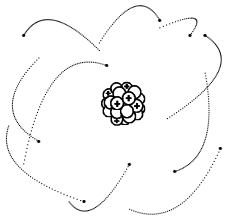
Um das Verhalten von Teilchen in elektrischen Versuchen beschreiben zu können, schreibt man ihnen auch die Eigenschaft *elektrische Ladung* zu. Vorläufig genügt es uns, dass diese Ladung entweder *positiv*, *negativ* oder gar nicht vorhanden sein kann. Später werden wir Ladungen auch quantitativ (=zahlenmäßig) beschreiben. Die Ladung eines Teilchens macht sich unter anderem wie folgt bemerkbar: Zwei gleichnamig geladene – z.B. zwei positive – Teilchen stoßen sich ab, zwei ungleichnamige – ein positives und ein negatives – ziehen sich an. Positive und negative Ladungen können sich in ihrer Wirkung auf ihre Umgebung aufheben.

Manche der Teilchen, aus denen unsere Welt besteht, sind nun positiv, andere negativ, die übrigen gar nicht elektrisch geladen. Jede Teilchensorte hat eine charakteristische Ladung. Damit Du Dir eine tragfähige Vorstellung vom Aufbau und den Vorgängen in einem Körper machen kannst, solltest Du folgende drei Sorten von Teilchen kennen:

Neutronen sind elektrisch ungeladen (*neutral*).

Protonen sind elektrisch positiv geladen und etwa genauso groß und schwer wie Neutronen.

Elektronen sind elektrisch negativ geladen und sehr viel kleiner und leichter als Neutronen und Protonen. Die elektrische Ladung eines Elektrons ist jedoch vom Betrage her genauso groß wie die eines Protons.



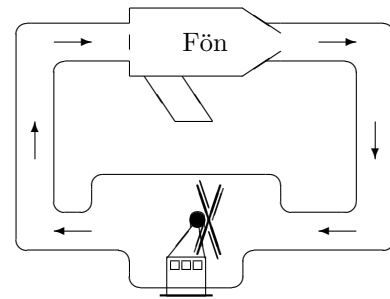
Aus diesen drei Sorten von Elementarteilchen sind Atome aufgebaut: Um einen dicht gepackten Kern aus Protonen und Neutronen schwirren weitläufig Elektronen. Ein Atom enthält dabei in seiner Elektronen-Hülle genauso viele Elektronen wie Protonen im Kern, so dass es als Ganzes elektrisch neutral wirkt.

Atome können untereinander Verbindungen eingehen und bilden so die Körper unserer Umgebung, und uns selbst natürlich auch. Je nachdem, wie fest sich die Atome gegenseitig festhalten, ist der Körper als Ganzes mehr oder weniger fest: In einem Metallstück ist der Zusammenhalt recht groß, in Wasser halten die Teilchen weniger stark zusammen, in der Luft fast gar nicht.

In **Metallen** sind aus jedem Atom ein bis zwei Elektronen herausgelöst. Solche „freien“ Elektronen schwirren zwischen den festsitzenden verbleibenden „Atomrümpfen“ praktisch ungehindert, kreuz und quer durch das ganze Metallstück. Sie können nur nicht ohne Weiteres das Metallstück verlassen.

In einem metallenen Stromkreis treibt die Stromquelle einen Kreislauf aus leicht beweglichen Elektronen durch die Leitungen. Aus ihrem Minuspol drückt sie Elektronen in die Leitung hinein, in ihren Pluspol zieht sie gleichzeitig Elektronen aus der Leitung heraus.

Vergleich: Die Elektronen füllen den Leerraum zwischen den verbleibenden Atomrümpfen eines Metallstückes so aus, wie Luft alle Leerräume zwischen den Gegenständen in einem Zimmer ausfüllt. Man spricht deshalb auch von einem „Elektronengas“ im Metall. Ein elektrischer Stromkreis lässt sich also mit einem geschlossenen Rohrsystem vergleichen. In dem hier abgebildeten Rohrkreislauf zirkuliert Luft, von einem Fön als Pumpe angetrieben.



Bei diesem Vergleich entspricht ...

der Stromquelle: _____	den Elektronen: _____
dem elektr. Gerät: _____	einem Schalter: _____
einer Leitung: _____	einer Diode: _____

Mit diesem Vergleich wird auch klar, dass ein dauerhafter Stromfluss nur bei durchgehend geschlossenem Stromkreis möglich ist: Kann die Luft an einer Stelle des Rohrkreislaufes nicht weiterströmen, kommt es zu einem anwachsenden Stau. Die Pumpe kann gegen den Stau immer schlechter weitere Luft in das Rohr hin-

eindrücken und bekommt auch auf ihrer ansaugenden Seite immer schlechter Nachschub an Luft. Schließlich kommt der Luftstrom ganz zum Erliegen. Ob ein solcher Anstau-Effekt auch in einem Stromkreis möglich ist, zeigt der folgende Versuch.

V 19.3 Löffeln von Ladung

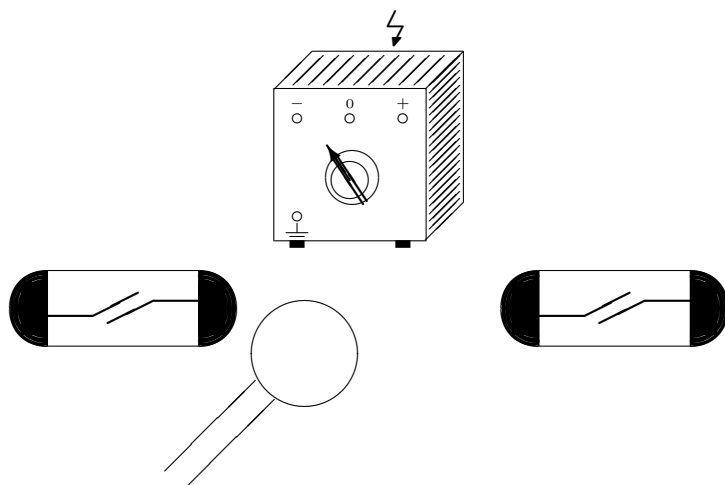
zu ergänzen:

Aufbau: starke Stromquelle, zwei Glühlampen wie gezeichnet angeschlossen; eine große Metallkugel an einem Plastikstab in der Hand

Durchführung: Wir berühren mit der großen Kugel nacheinander die freien Enden der beiden Glühlampen, zuerst die linke.

Beobachtung: Beim Berühren der beiden Glühlampen blitzt die berührte Glühlampe kurz auf, und zwar in beiden Fällen an der

_____ Elektrode.



Ergebnis: Beim Berühren fließt kurzzeitig ein Strom durch die leuchtende Glimmlampe. (die andernfalls ja nicht leuchten würde.)

Erklärung: Während der Berührung kann die Stromquelle durch die entsprechende Glimmlampe ein wenig Ladung in die Kugel hineindrücken bzw. aus ihr abziehen.

Genauer: Es werden zunächst Elektronen durch die linke Glimmlampe auf die Kugel gedrückt. Diese Elektronen werden dann mit der Kugel von Hand zur rechten Glimmlampe gebracht. Bei der dortigen Berührung zieht dann die Stromquelle Elektronen von der Kugel ab. Bei beiden Berührungen sind Elektronen von der linken zur rechten Elektrode der beteiligten Glimmlampe geflossen.

Elektronen lassen sich also in einem Metallstück wie der großen Kugel anstauen. Um diesen Effekt zu sehen, mussten wir allerdings eine große Kugel (mit viel Platz für Elektronen), eine starke Stromquelle und die empfindlichen Glimmlampen zum Nachweis ganz schwacher Ströme verwenden.

In den Versuchen, die wir sonst durchführen, bemerken wir keine Staueffekte: Die „Elektronen-Füllung“ des gesamten Stromkreises bewegt sich praktisch als Ganzes oder gar nicht.

19.5 Stromrichtung

In einem metallenen Stromkreis bewegen sich die frei beweglichen Elektronen außerhalb der Stromquelle von deren Minuspol zum Pluspol. Innerhalb der Stromquelle werden sie natürlich vom Pluspol zum Minuspol geschafft. Insgesamt muss sich ja ein Kreislauf ergeben.

Da in den meisten Stromkreisen der Strom ein Strom von Elektronen ist, läge es nahe, deren Bewegungsrichtung als die Richtung des Stromes zu nehmen. Man muss jedoch bedenken, dass mit den Elektronen *negative* elektrische Ladung fließt. An Stellen, wo Elektronen z.B. *wegfließen*, fließt daher rechnerisch (positive) Ladung *hin*. Dies kennst Du aus der Mathematik: Soll z.B. von 100 die negative Zahl -10 *subtrahiert* werden, so *addiert* man praktisch die (positive) Zahl 10.

Die Stromrichtung (außerhalb einer Stromquelle) ist daher die Richtung vom Pluspol zum Minuspol.
Kurz: Strom fließt von Plus nach Minus.

Hätte man zu Beginn der Erforschung der Elektrizität bereits mehr über die Vorgänge in einem Leiter gewusst, so hätte man bestimmt die Vorzeichen der Pole anders herum festgelegt, und die Ladung der Elektronen dazu passend als positiv.

Wenn also nichts anderes gesagt ist, ist mit der Stromrichtung die Richtung von Plus nach Minus gemeint. Um in Zweifelsfällen ganz deutlich zu machen, welche Richtung man meint, bezeichnet man diese rechnerische Stromrichtung als die **technische Stromrichtung**, weil Techniker in ihren Rechnungen nur diese Richtung benutzen. Die Bewegungsrichtung der Ladungsträger nennt man die **physikalische Stromrichtung**, weil Physiker sich ja mehr für die tatsächlichen Bewegungen im Leiter interessieren. Rechnen tun Physiker aber ebenfalls mit der technischen Richtung.

19.6 Leiter und Nichtleiter

Um Stromkreise herzustellen, haben wir bisher immer metallene Leitungen mit Kunststoffmänteln (– =Kabel! –) verlegt. Für einen elektrischen Kontakt genügt es dabei nie, wenn sich Kabel am Kunststoff berührten, sondern es musste Metall mit Metall in Kontakt kommen. Strom kann offensichtlich gut durch Metalle fließen, aber nicht (oder zumindest viel schlechter) durch den Kunststoff.

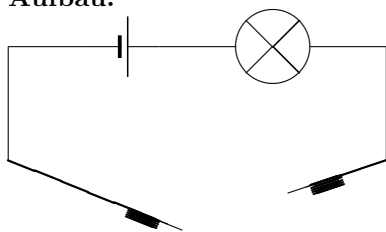
Stoffe, die den Strom (gut) leiten, nennt man **Leiter**. Stoffe, die Strom (sehr schlecht oder) nicht leiten, nennt man **Nichtleiter**, oder lateinischer: **Isolatoren**.

?

Welche Stoffe sind Leiter, welche sind Isolatoren?

V 19.4 Leitfähigkeitstests

Aufbau:



Batterie und Lämpchen verbunden wie nebenstehend gezeichnet, mit Steckern von Experimentierkabeln als Enden des offenen Stromkreises.

Durchführung: Wir halten die beiden Stecker an verschiedene Materialproben, um festzustellen, ob das Material leitend ist. (Ist ein Material leitend, schließt es den Stromkreis und das Lämpchen sollte brennen.)

Beobachtung: Bei manchen Materialproben leuchtet das Lämpchen auf, bei anderen nicht.

es leuchtet bei	es leuchtet nicht bei

Ergebnis:

↑ Diese Stoffe sind Leiter ↑

|| ↑ Diese Stoffe sind Isolatoren ↑

Wie gut ein Material leitet, hängt davon ab, wie viele wie frei bewegliche Ladungsträger, z.B. Elektronen, es enthält.

19.7 Steckdose

Jede Stromquelle muss einen Plus- und einen Minuspol haben. Wir fragen daher erst einmal:

Wo sind bei einer Steckdose Plus- und Minuspol?

?

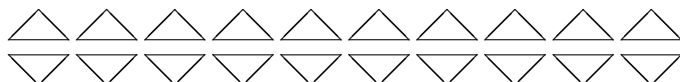
V 19.5 Bestimmung der Pole einer Steckdose mit Dreiecks-Glimmlampe

Durchführung: Eine Glimmlampe wird (nur vom Lehrer!) an die Steckdose angeschlossen.

Beobachtung: Um beide Elektroden herum leuchtet es, wie am Rand gezeichnet. Die Steckdose scheint zwei Minus-Pole zu haben. Nach unserer Vorstellung von Strom kann dies jedoch nicht sein.

Durchführung: Im abgedunkelten Raum wird die Glimmlampe zügig hin- und hergeschwenkt. Dadurch sehen wir die Glimmlampe zu verschiedenen Zeiten an verschiedenen Stellen.

Beobachtung: Entlang des Weges der Glimmlampe sehen wir die obere und die untere Elektrode *abwechselnd einzeln* hell. Die folgende Abbildung zeigt die Glimmlampe an verschiedenen Stellen ihres Weges: (und damit auch zu verschiedenen Zeiten)



Ergebnis: Auch eine Steckdose hat jederzeit einen Plus- und einen Minuspol. Die Pole wechseln einander jedoch ab. Eine Zeit lang ist z.B. der linke Anschluss Pluspol und der rechte Minuspol. Danach ist aber ebensolang der linke Pol Minuspol und der rechte Pluspol.

zu ergänzen:

- die leuchtenden Bereiche

Die Polung wechselt an einer Steckdose so rasch, dass wir mit unseren Augen den Wechsel nicht verfolgen können. In einer Sekunde wechselt die Polung 50 mal hin und her, jeder Anschluss ist 50 mal Plus- und 50 mal Minuspol.

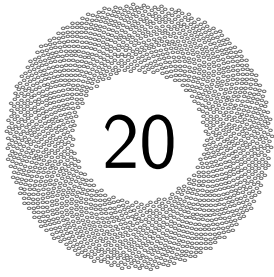
Mit der Polung wechselt auch der Stromfluss seine Richtung. Strom aus der Steckdose wird deshalb **Wechselstrom** genannt. Strom, der seine Richtung nicht ändert, wird **Gleichstrom** genannt.

Um anzugeben, wie rasch die Wechsel erfolgen, muss man angeben, wie viele Wechsel in einer bestimmten Zeitspanne erfolgen. Man zählt dabei jedoch einen hin- und den folgenden her-Wechsel als einen Wechsel.

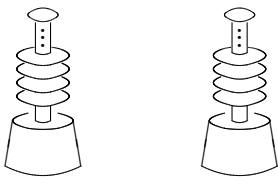
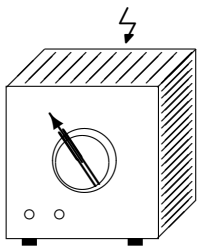
Die Anzahl der Wechsel pro bestimmter Zeitspanne nennt man nach dem lateinischen Wort *frequens*=häufig die **Frequenz**. Als Einheit zur Angabe von Frequenzen benutzt man 1 Hertz, abgekürzt 1 Hz. Ein Hertz bedeutet einen hin- und her-Wechsel pro Sekunde. Also:

Der Wechselstrom aus der Steckdose hat eine Frequenz von 50 Hz.

Heinrich Hertz
dt. Physiker
* 22.2.1857,
Hamburg
† 1.1.1894, Bonn



Wirkungen des elektrischen Stroms



zu ergänzen:

- Verkabelung
- eingespannter Glühdraht

Elektrischer Strom erwärmt i.d.R. die durchflossenen Leiter.

20.1 Wärmewirkung

In unseren bisherigen Versuchen haben wir oft mit Glühlampen gearbeitet. Innerhalb des Glaskolbens einer solchen Lampe erkennen wir einen dünnen, eventuell schraubenförmig gewickelten Draht. Dieser Draht ist einerseits mit der Lampenfassung und andererseits mit der Metallkuppe am Fuß der Lampe verbunden. Fließt Strom durch die Lampe, so fließt dieser Strom also sicher auch durch das Drähtchen. Ist die Lampe eingeschaltet, so glüht dieses Drähtchen und erzeugt so das gewünschte Licht. Da es die Aufgabe dieses Drähtchens ist, zu glühen, nennt man es Glühdraht oder auch Glühwendel, falls es schraubig gewickelt ist. Offensichtlich bewirkt der Stromfluss durch das Drähtchen dessen Erwärmung bis hin zum Glühen.

V 20.1 einen Draht zum Glühen bringen

Aufbau: siehe Rand

Durchführung: Wir erhöhen am Netzgerät langsam die Spannung.

Beobachtung: Der Glühdraht wird wärmer, glüht zunächst schwach rötlich, später intensiver und gelblicher, und er brennt schließlich durch.

Ergebnis: Fließt Strom durch einen Draht, so wird der Draht erwärmt.

Die Kabel in obigem Versuch bestehen natürlich auch aus Drähten. Wir beobachten an diesen Kabeln jedoch keine Erwärmung. Wieso nicht? Die Antwort ist einfach: Die Kabel bestehen aus mehreren und/oder dickeren Drähten. Um diese zum Glühen zu bringen, braucht man einfach mehr Strom. Auch die Drähte im Kabel würden glühen, wenn nur genug Strom flösse. Natürlich sind die Kabel mit Absicht so dick gemacht, dass bei normalem Gebrauch keine nennenswerte Erwärmung auftritt. Ein bißchen wärmer werden die Kabel aber sehr wohl auch bei schwachen Strömen.

Die Wärmewirkung an einem Draht lässt sich verstärken, wenn man den Draht schraubenförmig wendelt. Jede Windung erwärmt dann nämlich ihre Nachbarwindungen, und der Draht erreicht im Ganzen eine höhere Temperatur. Auf dieser Seite ist noch ein wenig Platz für Deine Zeichnung einer Glühlampe!

Die Wärmewirkung des Stromes ist Grundlage vieler weiterer Anwendungen.

- Meistens geht es darum, etwas warm zu machen, wie z.B. in einem Bügeleisen.
- Bei einer Glühbirne geht es aber nicht um eine Erwärmung selbst, sondern um das Licht in Folge einer Erwärmung.
- Bei einer Schmelzsicherung geht es sogar um das Durchbrennen in Folge einer Erwärmung.

Mit einer Schmelzsicherung kann man verhindern, dass ein wertvolles Gerät kaputt geht, wenn aus irgendeinem Grunde der Stromfluss durch das Gerät ansteigt. Die Schmelzsicherung ist dazu eigentlich nur ein Drähtchen im selben Stromkreis, das durchbrennt, bevor das Gerät durchbrennt. Erwärmt der Strom das Drähtchen so sehr, dass es schmilzt, so reißt das Drähtchen. Der Stromkreis ist damit unterbrochen und der Strom hört auf, bevor er weiter steigen und das Gerät beschädigen kann.

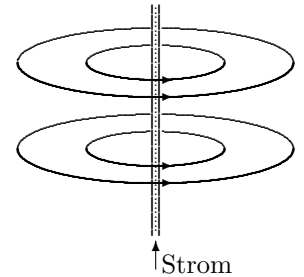
Bei anderen Anwendungen des Stromes ist die Wärmewirkung des Stromes unerwünscht. Ein Elektromotor z.B. soll Dinge bewegen, nicht erwärmen. Trotzdem wird auch ein Elektromotor im Betrieb vom Strom erwärmt, eventuell sogar so sehr, dass der Motor dauernd gekühlt werden muss oder immer nur kurz benutzt werden darf.

20.2 Magnetische Wirkung

Im Jahre 1820 entdeckte Hans Christian Ørsted (wahrscheinlich zufällig), dass Magnetnadeln in der Nähe von stromdurchflossenen Leitern Kräfte erfahren. Das Bild nebenan zeigt die Magnetfeldlinien um ein gerades Stück Leiter herum. Zur Erinnerung: Die Magnetfeldlinien zeigen an jeder Stelle die Richtung an, in die sich der Nordpol einer Magnetnadel dort einstellen würde.

Wickelt man einen Leiter zu einer Spule, so erhält man bei Stromfluss ein Magnetfeld wie bei einem Stabmagneten. Nord- und Südpol sind die runden Öffnungen der Spule. Die folgende Abbildung zeigt zwei Spulen.

Falls Du den „magischen Blick“ beherrschst, kannst Du mit ebendiesem in dieser Abbildung auch eine Spule räumlich sehen. Du musst dazu mit Deinem rechten Auge durch das rechte Spulenbild, mit dem linken Auge durch das linke Spulenbild schauen. Gelingt es Dir, siehst Du drei Spulen, von denen die mittlere räumlich wirkt.

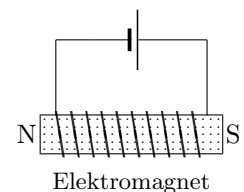


Beachte: Das Magnetfeld einer Spule entsteht nur aufgrund des Stromes. Es entsteht nicht, weil das Spulenmaterial magnetisch wäre. In der Regel sind Spulen sowieso aus Kupferdraht gewickelt, und Kupfer ist nicht magnetisch.

Bringt man in das Innere der Spule jedoch einen Weicheisen-Kern, so wird dieser Kern durch das Magnetfeld der Spule selbst magnetisiert. Er verstärkt dann mit seinem Magnetfeld das Feld der Spule noch erheblich. Aber auch hier verliert die ganze Anordnung ihre Magnetisierung wieder schnell, wenn der Stromfluss aufhört.

Eine solche Anordnung – eine Spule mit einem Eisenkern – nennt man einen **Elektromagneten**. Nebstehende Abbildung zeigt einen solchen Elektromagneten, der an eine Batterie angeschlossen ist. Mit N und S sind der Nord- bzw. der Südpol bezeichnet.

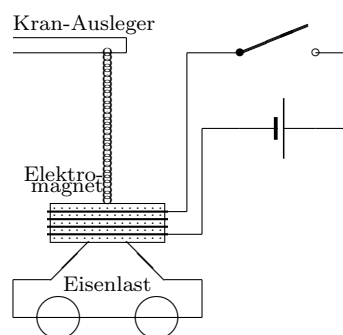
Die Lage der Magnetpole ergibt sich aus dem Umlaufsinn des Stromes um die Spulenachse. Wären entweder die Anschlüsse der Batterie vertauscht, oder wäre die Spule anders herum gewickelt, so lägen die Pole ebenfalls vertauscht.



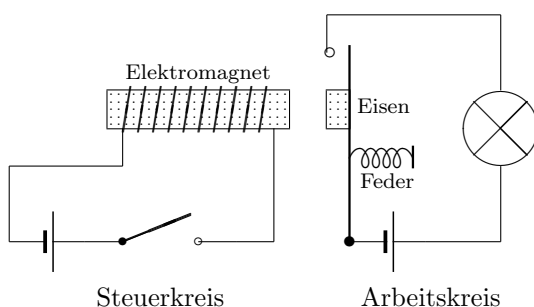
20.2.1 Einige Anwendungen von Elektromagneten

Hebemagnet

Eiserne Lasten (z.B. Autos) lassen sich mit einem Elektromagneten per Knopfdruck halten und wieder loslassen. Dies erspart umständliches Befestigen und Lösen von Haken, Seilen o.ä.



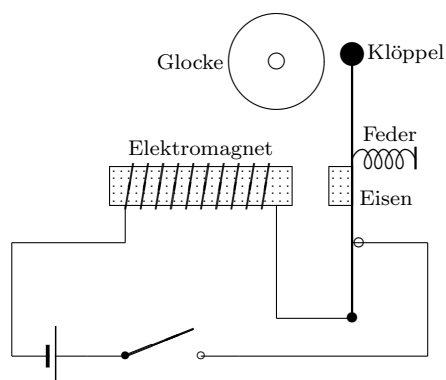
Relais



Ein Elektromagnet im Steuerkreis betätigt einen Schalter im Arbeitskreis. So kann man auch gefährlich starke Ströme im Arbeitskreis mit einem Schalter ein- und ausschalten, der selbst nur vom schwachen Strom für den Elektromagneten durchflossen werden kann.

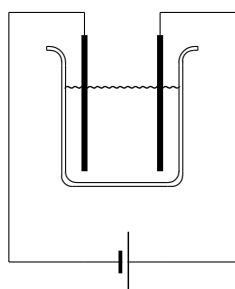
Klingel (Wagner'scher Hammer)

Ein Elektromagnet unterbricht seine eigene Stromversorgung. So kommt es zu einem automatischen Ein- und Ausschalten des Magneten und damit zu einem dauernd wiederholten Anschlagen der Glocke.



Eine weitere – allerdings kompliziertere – Anwendung sind **Elektromotoren**. In ihnen bewegen Elektromagnete sich oder andere Bauteile so, dass sich eine Drehung ergibt. Elektromotoren werden in einem späteren Schuljahr genauer behandelt.

20.3 Chemische Wirkung



Noch diese eine Wirkung sei kurz erwähnt: Elektrischer Strom kann Leiter chemisch verändern, d.h. ihre stoffliche Zusammensetzung ändern. Fließt z.B. im „Hofmann'schen Zersetzungsapparat“ ein Strom durch Schwefelsäurelösung, so bewirkt der Strom letzten Endes eine Aufspaltung von Wasserteilchen. Es bilden sich an den beiden Elektroden getrennt die Gase Wasserstoff und Sauerstoff. Eine Mischung dieser Gase ist explosiv und wird daher „Knallgas“ genannt. Bei einer Knallgasexplosion entsteht wieder Wasser.

Anderes leicht vorführbares Beispiel: Es bilden sich Chlorgas und Wasserstoff, wenn Strom über Kohlelektroden durch eine Kochsalzlösung fließt. (Vorsicht: Das aus der Lösung entweichende, „nach Schwimmbad“ riechende Gas reizt die Atemwege!)

Messgrößen des elektrischen Stroms

21.1 Die Stromstärke

Wenn wir ein Glühlämpchen an ein einstellbares Netzgerät (z.B. an einen Trafo für Modelleisenbahnen) anschließen, können wir das Lämpchen verschieden hell leuchten lassen. Diese unterschiedlichen Helligkeiten können wir natürlich gut wahrnehmen. Ebenso können wir einen Elektromotor verschieden schnell laufen lassen, je nachdem, wie weit wir das Netzgerät „aufdrehen“.

In beiden (und vielen weiteren) Fällen erleben wir *unterschiedlich starke Wirkungen* des elektrischen Stromes. Zur Erklärung dieser unterschiedlichen Stärke in seiner Wirkung schreiben wir dem elektrischen Strom selbst eine Stärke zu. Wir führen daher also eine Größe namens **Stromstärke** ein. In Formeln schreibt man für die Stromstärke das Symbol **I**. Warum I? Nun, man denke an *Intensität* des Stromes.

Jetzt fehlen noch ein Messverfahren und eine Einheit für die Stromstärke. Die vorangegangenen Überlegungen weisen auch schon den Weg dahin: Wir werden Stromstärken über ihre Wirkungen messen. Grundsätzlich kann jede der uns bekannten Wirkungen (Wärme, Magnetfelder, chemische Veränderungen) zur Messung herangezogen werden.

übliches Symbol für die elektrische Stromstärke: I

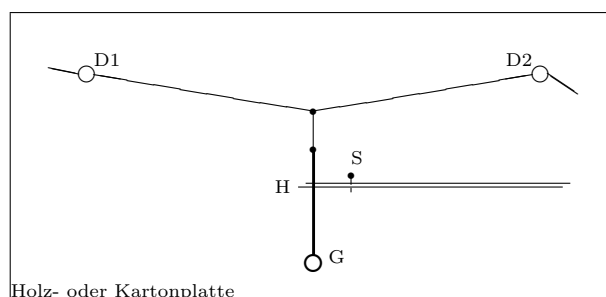
21.1.1 Messverfahren für die Stromstärke

Hitzdrahtinstrument

Stromdurchflossene Leiter erwärmen sich, und die meisten Körper dehnen sich bei Erwärmung ein wenig aus. Ein stromdurchflossener Draht sollte demnach umso länger werden, je stärker ein Strom durch ihn fließt. Allerdings verlängert sich ein Draht bei Erwärmung nur um einen recht kleinen Bruchteil seiner ursprünglichen Länge. Wir müssen daher durch einen geschickten Aufbau diese kleinen Längenänderungen des Drahtes in gut sichtbare Veränderungen des Aufbaus übersetzen.

Bauplan für ein Hitzdrahtinstrument:

Spanne auf einer Holz- oder Kartonplatte einen etwa 20 cm langen, dünnen Draht zwischen zwei Reißzwecken D1 und D2! Dieser Draht ist der Hitzdraht, der später vom Strom durchflossen wird, sich erwärmt und ausdehnt. Befestige einen Gummi-Faden mit der Reißzwecke G am unteren Rand der Platte! Binde mit einem Drähtchen das andere Ende des Gummifadens an den Hitzdraht, so dass das Gummi gespannt ist! (Da der Hitzdraht warm oder sogar heiß wird, wäre es nicht günstig, das Gummi direkt an ihn zu binden. Das Gummi könnte nämlich verschmoren.) Klebe einen Strohhalm an der Stelle H (möglichst weit oben) an den Gummifaden. Durchbohre den Halm mit der Stecknadel S und stecke ihn mit der Nadel fest, so dass der Halm sich um die Stecknadel drehen kann! Fertig!



?

Was passiert, wenn Strom durch den Hitzdraht fließt?

Zunächst wird der Hitzdraht durch den Stromfluss wärmer und dehnt sich ein wenig aus. Die Gummischnur kann die Mitte des Drahtes dann ein deutliches Stück weiter nach unten ziehen. Das festgeklebte Ende des Strohhalmes bewegt sich entsprechend mit nach unten. Wie bei einer Wippe bewegt sich das freie Ende des Halmes nach oben, und zwar um ein noch größeres Stück.

So wird die geringfügige Verlängerung des Drahtes in eine gut sichtbare Bewegung der freien Halm-Spitze übersetzt. Eine große Stromstärke wird dabei die Spitze weiter nach oben bringen als eine kleine Stromstärke. Die Halm-Spitze zeigt damit die Stromstärke an und ist demnach der *Zeiger* unseres Hitzdrahtinstrumentes.

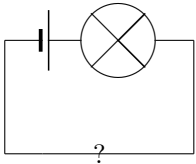
?

Wie benutzt man dieses Instrument nun?

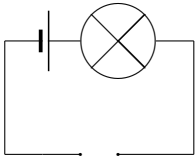
Der Strom, dessen Stärke man messen will, muss durch den Hitzdraht fließen. Wenn Du also die Stromstärke an einer Stelle eines Stromkreises messen willst, musst Du den Stromkreis an dieser Stelle auftrennen und die Lücke mit dem Hitzdraht wieder schließen. Praktisch heißt das, dass Du die beiden losen Enden des aufgetrennten Stromkreises mit den Reißzwecken D1 bzw. D2 in Kontakt bringst.

Die Überlegung, dass der zu messende Strom durch das Messgerät fließen muss, damit dieses die Stromstärke messen kann, ist für jedes Stromstärken-Messgerät richtig. Es gilt daher allgemein:

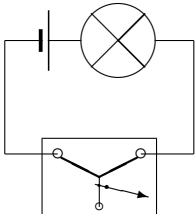
Stromstärken-Messgeräte baut man wie folgt in einen Stromkreis ein: Stromkreis an der interessierenden Stelle auftrennen und das Stromstärken-Messgerät in die entstandene Lücke einschalten.



↓



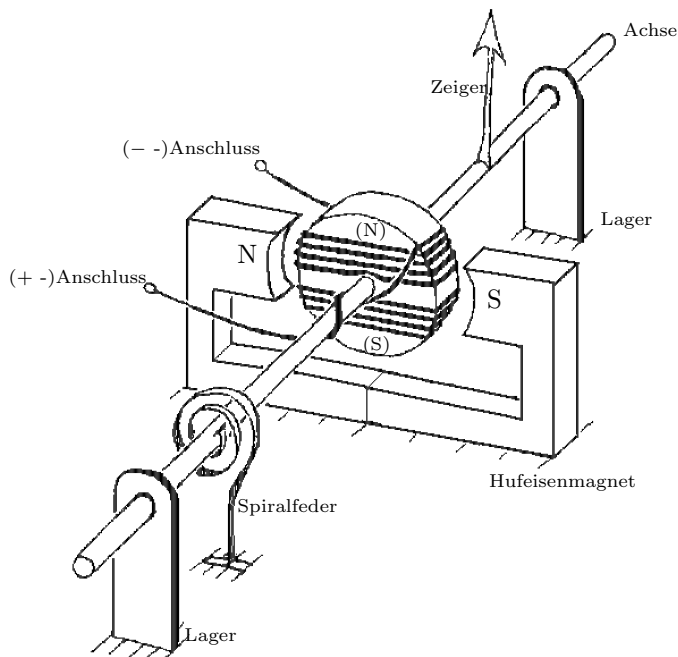
↓



Drehspulinstrument

Ein Drehspulinstrument nutzt die magnetische Wirkung des Stromes aus. Wie der Name bereits andeutet, ist ein wesentlicher Bestandteil dieses Instrumentes eine drehbar gelagerte Spule. Sie ist oft auf einen zylinderförmigen Eisenkern gewickelt. Färbe doch in der folgenden Abbildung die drehbar gelagerte Achse und den fest mit ihr verbundenen Eisenkern gelb! Färbe die Lager braun oder grau!

Spule und Eisenkern bilden einen Elektromagneten. Wir nehmen im Folgenden an, dass die Spule so von Strom durchflossen werde, dass der Nordpol dieses Elektromagneten oben liegt. Diese Spule wird nun zwischen zwei feststehende Magnetpole montiert wie abgebildet. Natürlich fehlt diesen Polen noch die richtige Farbe: Färbe die Nord-Hälfte des Hufeisenmagneten wie üblich rot, die Süd-Hälfte grün!



Bei Stromfluss wird der Nordpol der Spule zum fest montierten Südpol rechts und der Südpol der Spule zum Nordpol links gezogen. Die Spule dreht sich demnach um 90° im Uhrzeigersinn, so dass ihre Pole den jeweils andersnamigen Polen außen gegenüberstehen. Behindert man diese Drehung mit einer Spiralfeder — färbe sie blau! — wie abgebildet, so kann sich die Spule nicht vollständig in diese Stellung drehen. Es gelingt ihr jedoch umso besser, je stärker die Magnetkräfte zwischen ihren Polen und den äußeren Polen sind.

Natürlich sind diese Magnetkräfte umso größer, je stärker der Strom in der Spule ist. Man kann demnach die Stärke des Stromes daran ablesen, wie weit sich die Spule drehen konnte. Zum leichteren Ablesen bringt man noch einen Zeiger an der Spule an. Fertig ist das Drehspulinstrument! Färbe den Zeiger gelb, weil er ja fest mit der Achse verbunden ist!

Knallgas-Produktion

Auch die chemische Wirkung von elektrischem Strom lässt sich zum Messen der Stromstärke ausnutzen. Strom kann Wasser zu Knallgas zersetzen. Ein stärkerer Strom erzeugt dabei in gleicher Zeit mehr Knallgas als ein schwacher Strom. Wenn man also bestimmen kann, wie viel Knallgas in einer bestimmten Zeit gebildet wird, hat man damit wiederum ein Maß für die Stromstärke. Wir werden diese Messmethode jedoch nicht näher kennenlernen, weil sie recht unpraktisch ist.

21.1.2 Die Einheit Ampere

Die Einheit der Stromstärke ist das Ampere.
Ein Ampere wird abgekürzt mit 1 A.

Der Name wurde gewählt zu Ehren des französischen Physikers André Marie Ampère.
(* 22.1.1775 bei Lyon, † 10.6.1836 in Marseille)

Die einfachste Art, ein Ampere zu definieren, ist folgende:

Ein Ampere ist die Stromstärke, die bei einem bestimmten Drehspulinstrument einen bestimmten Zeigerausschlag hervorruft.

Diese Definition hat den Nachteil, dass sie sich auf ein bestimmtes Messgerät bezieht. Jemand, der über dieses spezielle Messgerät nicht verfügt, kann mit dieser Definition überhaupt nichts anfangen. Für unsere weiteren Zwecke wird diese Definition jedoch genügen. Wir gehen also davon aus, dass wir bei einem bestimmten Drehspulinstrument eine bestimmte Zeigerstellung

als Marke für 1 A festgelegt haben. Praktisch haben wir dazu vielleicht einfach hinter den Zeiger ein Stück Papier montiert und darauf einen Strich bei der bestimmten Zeigerstellung gezogen. Einen weiteren Strich haben wir dort gezogen, wo der Zeiger bei stromlosem Gerät steht. Damit haben wir bereits eine primitive Skala angelegt, die die Markierungen für 0 A und für 1 A enthält.

Beachte: Wenn wir die Einheit Ampere tatsächlich neu definieren könnten oder müssten, so könnten wir die 1 A-Marke ganz willkürlich setzen. Natürlich wurde das Ampere aber bereits lange vor unseren Bemühungen zum ersten mal definiert. Dein(e) Lehrer(in) hat deshalb dafür gesorgt, dass das von uns (scheinbar willkürlich) definierte Ampere etwa mit dem Ampere übereinstimmt, das der Rest der Welt auch benutzt.

Als nächstes benötigen wir zum Zeichnen von Schaltplänen noch ein Symbol für ein Stromstärken-Messgerät. Es ist einfach ein Kreis mit einem großem „A“ darin, wie Du es neben diesem Text sehen kannst. Das „A“ kommt natürlich von Ampere. Ein Messgerät für Stromstärken nennt man **Amperemeter**. Konkret ist dies oft ein Drehspulinstrument.



In diesem Skript wirst Du in manchen Schaltplänen zur Unterscheidung von mehreren Amperemetern Indices an den A's finden. Es wird dann die Stromstärke, die z.B. das Amperemeter misst, mit I_3 bezeichnet.



Für manche Untersuchungen muss man gleichzeitig die Stromstärken an mehreren Stellen messen. Dazu benötigt man dann mehrere Amperemeter. Bisher haben wir allerdings nur ein erstes Ur-Amperemeter, mit dem wir ja erst das Ampere „willkürlich“ definieren. Natürlich sollen alle weiteren Amperemeter dieselbe Stromstärke als 1 A anzeigen.

Zur Herstellung dieser weiteren Amperemeter kann man daher nicht wieder an willkürlicher Stelle den 1 A-Strich anbringen. Man muss ihn vielmehr dort anbringen, wo der Zeiger steht, wenn genau ein Ampere durch das Gerät fließt. Bei einem Drehspulinstrument, das dem Ur-Amperemeter „exakt baugleich“ ist, ist dies

leicht: Man kann den 1 A-Strich einfach an gleicher Stelle anbringen wie beim Ur-Amperemeter.

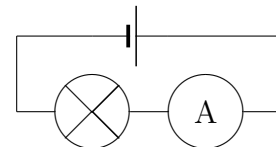
Streng genommen sind allerdings zwei Geräte nie exakt baugleich. Es ist damit streng genommen mit unseren momentanen Kenntnissen auch nicht ganz so einfach, mehrere genau anzeigende Amperemeter herzustellen.

Mit den so hergestellten Messgeräten können wir nun feststellen, ob überhaupt ein Strom fließt, und ob seine Stromstärke größer, kleiner oder gleich einem Ampere ist. Wir können sogar die Richtung des Stromes bestimmen. Wir beginnen nun mit ersten Untersuchungen der Stromstärke.

21.1.3 Eigenschaften der Stromstärke

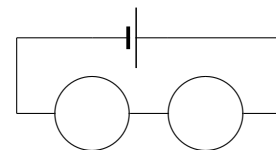
V 21.1 Stromstärken in einem einfachen Stromkreis, I

Aufbau: ein Stromkreis vom Pluspol einer Batterie erst über ein Amperemeter und dann ein Glühlämpchen zum Minuspol der Batterie



Durchführung: Wir merken uns oder markieren die Stellung des Zeigers im Amperemeter.

Wir bauen den Stromkreis nun so um, dass der Stromweg zuerst durch das Glühlämpchen und dann durch das Amperemeter führt:



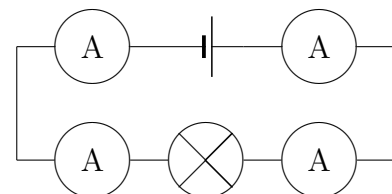
zu ergänzen:

- Inneres der Kreis-Symbole

Beobachtung: /Ergebnis:

V 21.2 Stromstärken in einem einfachen Stromkreis, II

Aufbau: ein unverzweigter Stromkreis mit Batterie, einem Glühlämpchen und mehreren Amperemetern an verschiedenen Stellen



Beobachtung: / Ergebnis:

Diese jüngsten Versuchsergebnisse entsprechen genau unserer bisherigen Vorstellung vom elektrischen Strom: Die Batterie setzt und hält einen Kreislauf elektrischer Ladung in Gang (,indem sie Elektronen durch den Stromkreis treibt).

Die nebenstehende Abbildung stellt einen Leiter mit darin nach oben strömenden Elektronen dar. Mit den Elektronen strömt auch negative Ladung durch den Leiter. Ein Abschnitt des Leiters wird durch zwei Leiterquerschnitte begrenzt. In

diesen Abschnitt fließt durch den unteren Querschnitt negative Ladung hinein, durch den oberen heraus.

Wir wissen bereits, dass es in Stromkreisen nicht zu einem andauernden Anstauen von Ladung kommen kann. In jeden Abschnitt des Stromkreises fließt also von der einen Seite her ebensoviel Ladung hinein, wie in derselben Zeit auf der anderen Seite hinausfließt. Fassen wir zusammen:

Unsere letzten Versuche ergaben, dass die Stromstärke an jeder Stelle eines unverzweigten Stromkreises gleich groß ist.

Unsere Vorstellung beinhaltet, dass in derselben Zeitspanne durch jeden Leiterquerschnitt die gleiche Menge an Ladung hindurchtritt.

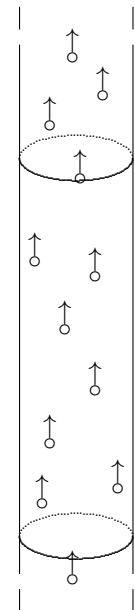
Es ist außerdem gut vorstellbar, dass ein Strom stärker ist,

wenn in derselben Zeit mehr Ladung einen Querschnitt passiert.

Die Stromstärke an einer Stelle ist nichts anderes

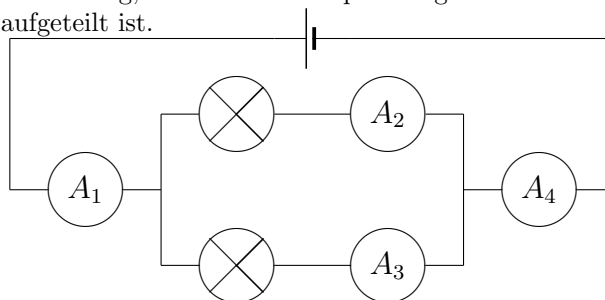
als die dort durchströmende Ladungsmenge pro Zeit.

Wir bestimmen zwar die Stromstärke nicht, indem wir die Ladungsmenge bestimmen, die einen Leiterquerschnitt pro Zeit passiert, aber zum Einbau eines Amperemeters müssen wir einen Stromkreis dennoch „aufschneiden“.



V 21.3 Stromstärken im verzweigten Stromkreis

Aufbau: Eine Schaltung, in der eine Hauptleitung auf einem Abschnitt in zwei Zweigleitungen aufgeteilt ist.



Beobachtung: /Ergebnis: (z.B. in Form von (Un-)Gleichungen)

Auch dieses Ergebnis verträgt sich sehr gut mit unserer bisherigen Vorstellung: Der Gesamtstrom in der Hauptleitung teilt sich am Verzweigungspunkt in zwei Teilströme auf. Es ist klar, dass ein Teilstrom schwächer ist als der Gesamtstrom. Vereinigen sich die Teilströme wieder, erhält man wieder dieselbe Stromstärke wie vor der Verzweigung.

Mit dieser Vorstellung nehmen wir mehr an, als wir zur Zeit messen können: Wir können messen, dass jeder Teilstrom schwächer als der Gesamtstrom ist. Darüber hinaus unterstellen wir, dass die Teilstromstärken addiert genau die Gesamtstromstärke ergeben sollten.

Unsere Amperemeter hatten bisher nur einen 0 A-Strich und einen 1 A-Strich. Nun lassen sie sich vervollständigen: Eine Stromstärke von 2 A sollte vorliegen, wenn sich zwei Teilströme von je 1 A Stärke vereinigen. Wie erhält man zwei Teilströme von je 1 A Stärke? Wenn wir in der Schaltung des vorigen Versuches die Zweigleitungen mit geeigneten baugleichen Glühlampen versehen, ist die Stromstärke in beiden Leitungen stets gleich. (Wie sollte bei baugleichen Zweigen eine Ungleichheit entstehen?)

Wenn wir außerdem als Stromquelle ein einstellbares Netzgerät verwenden, können wir dieses so einstellen, dass gerade 1 A durch die Lämpchen fließt. Dann haben wir in der Hauptleitung einen Strom der Stärke 2 A! Die entsprechende Zeigerstellung wird markiert und mit „2 A“ beschriftet.

Nach gleichem Verfahren lassen sich alle ganzzahligen Vielfachen von 1 A bestimmen. Für z.B. 10 A könnte man fünf Ströme von je 2 A vereinigen. Und Bruchteile? Ein zehntel Ampere sollte vorliegen, wenn sich ein Strom von 1 A auf zehn Teilströme aufteilt. ...

Um hohe und niedrige Stromstärken bequem angeben zu können, kann man – wie bei jeder Einheit – auch bei der Einheit Ampere die üblichen Vorsilben (siehe Seite 6) benutzen.

Beispiele: Statt 1000 A kann man auch 1 kA schreiben und „1 Kiloampere“ lesen. 0,02 A sind 20 mA und 1,3 mA sind 1300 μ A.

Den Zusammenhang zwischen den Stromstärken in den Leitungen an einem Verzweigungspunkt formuliert man ganz allgemein als die sogenannte Knotenregel. Zuvor musst Du natürlich wissen, was die „Knoten“ sind.

Kompliziertere Schaltpläne ähneln oft Netzen, wie z.B. Fischer sie verwenden. Man nennt daher einen Verzweigungspunkt auch einen **Knoten**, und einen geschlossenen Leitungsweg nennt man eine **Masche**. Du wirst später auch noch eine Maschenregel kennenlernen. Nun jedoch erst einmal die

Knotenregel: An jedem Knoten ist die Summe der hinfließenden Stromstärken gleich der Summe der abfließenden Stromstärken.

21.2 Die Spannung einer Quelle

21.2.1 Pumpstärke einer Quelle

Dieselbe Glühlampe kann unterschiedlich hell brennen, je nachdem, an welche Stromquelle wir sie anschließen. Dieselbe Lampe, die an einer Flachbatterie ordentlich hell leuchtet, leuchtet nur ganz schwach oder gar nicht, wenn wir sie an eine Taschenlampenbatterie anschließen. Sie brennt an einer Auto-Batterie sehr hell oder so-

gar durch. An der Steckdose würde sie sicher durchbrennen. Schließen wir zu der Lampe noch ein Amperemeter in den Stromkreis, können wir zu den verschiedenen Helligkeiten auch die entsprechend verschiedenen Stromstärken ablesen. Offensichtlich gilt:

Verschiedene Stromquellen können durch dasselbe Gerät verschieden starke Ströme antreiben.

Dieser Sachverhalt wurde bereits auf Seite 168 anschaulich beschrieben: Stromquellen pumpen (die Ladungsträger) unterschiedlich stark. Mit dieser Formulierung benutzen wir nun das Wort *stark* in zwei Be-

deutungen: Einen *Strom* nennen wir stark, wenn viel Ladung pro Zeit durch einen Leiterquerschnitt fließt. Eine *Stromquelle* nennen wir stark, wenn sie es schafft, durch eine Schaltung einen starken Strom zu treiben.

Die „Pump-Stärke“ einer Stromquelle bezeichnet man als **Spannung** und misst sie in der Einheit Volt.

Die Einheit Volt wurde zu Ehren des italienischen Physikers Alessandro Graf Volta (* 18.2.1745 in Como, † 5.3.1827 ebd.) benannt, der sich mit Elektrizität und insbesondere mit Batterien beschäftigte.

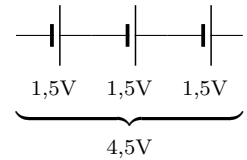
Klären wir noch das Formale: Als Symbol für die Spannung ist der Buchstabe **U** üblich. Warum U? Merkhilfe: Die Spannung einer Stromquelle ist die *Ursache* für den Stromfluss in der daran angeschlossenen Schaltung. Die Abkürzung für 1 Volt ist 1 V. Messgeräte für Spannungen heißen **Voltmeter** und werden mit eingekreistem V dargestellt.

Man kann die Einheit 1 Volt festlegen als die Spannung einer bestimmten Sorte von Batterien. Praktisch wäre es für uns am einfachsten, zu definieren, was 1,5 V ist: 1,5 V ist die *Spannung einer Taschenlampenbatterie*.



21.2.2 Serienschaltung von Quellen

Für viele Batterien wäre 1,5 V tatsächlich eine naheliegende Einheit: Eine Flachbatterie z.B. ist aus drei Zellen zusammengesetzt, von denen jede wie eine Taschenlampenbatterie aufgebaut ist. Durch eine sogenannte **Serienschaltung** dieser drei Zellen ergibt sich als Gesamtspannung die Summe der drei Zellenspannungen: $3 \cdot 1,5 \text{ V} = 4,5 \text{ V}$.



Eine Serienschaltung nennt man auch „Hintereinanderschaltung“ oder „Reihenschaltung“. Das Typische an einer Serienschaltung ist nämlich, dass die „in Serie“ geschalteten Bauteile nacheinander von derselben elektrischen Ladung durchflossen werden. Bei einem verzweigten Stromweg ist dies nicht der Fall. Dort bilden die Zweigleitungen zwischen den Knoten eine sogenannte **Parallelschaltung**, sind „nebeneinander“ geschaltet.

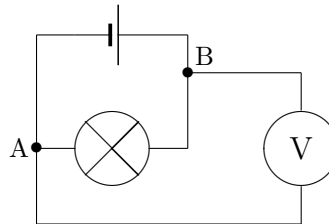
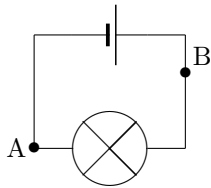
Nun weißt Du auch, warum eine Taschenlampenbatterie auch Monozelle genannt wird: *mono-* kommt aus dem Griechischen und bedeutet *ein-*. Eine Monozelle ist demnach eine Batterie, die aus nur einer 1,5V-Zelle besteht. Du denkst Dir jetzt sicher schon, wie Stabbatterien mit 3 V und Blockbatterien mit 9 V aufgebaut sein könnten!

21.2.3 Spannungsmessung

Eine Spannung besteht *zwischen* den beiden Polen einer Stromquelle. Ein Voltmeter muss deshalb mit seinen beiden Anschlüssen an die beiden Pole angeschlossen werden. Spannungen treten jedoch nicht nur zwischen den beiden Polen einer Quelle auf. Eine Spannung besteht aber immer *zwischen zwei Punkten*. Unterscheide: Eine Stromstärke misst man *an einer Stelle*.

Um die Spannung zwischen den Punkten A und B zu messen ...

...muss man das Voltmeter einfach nur an diese Punkte anschließen:



Die Schaltung muss nicht vorher unterbrochen werden wie zum Einbau eines Amperemeters.

Ein Voltmeter kann als Drehspulinstrument gebaut sein, durch das die zu messende Spannung einen Strom geringer Stärke antreibt. (Gering, damit dieser Test-Strom die übrige Schaltung nicht weiter stört und belastet) Da dieser Strom jedoch umso stärker ausfällt, je größer die anliegende Spannung ist, schlägt der Zeiger bei größerer Spannung weiter aus, – wie man es von einem Voltmeter erwartet. Neben solchen analogen Voltmetern mit Zeiger gibt es auch elektronische Digitalvoltmeter mit Ziffernanzeige.

21.3 Der Widerstand eines Verbrauchers

Unserer Erfahrung nach kommt ohne antreibende Spannungsquelle weder ein Stromfluss z.B. durch ein Lämpchen oder einen anderen Verbraucher zustande noch bleibt ein Stromfluss bestehen, wenn der Antrieb wegfällt. Offensichtlich hemmen Verbraucher den Stromfluss. Da dieselbe Spannungsquelle durch verschiedene Verbraucher unterschiedlich starke Ströme treibt, setzen offenbar verschiedene Verbraucher dem Strom unterschiedlich viel Widerstand entgegen.

Ebenso natürlich scheint es uns, dass man durch denselben Verbraucher gegen dessen Widerstand mit mehr Spannung mehr Strom treibt.

?

21.3.1 U-I-Kennlinien

Wie viel Strom treibt man mit wie viel Spannung durch einen Verbraucher?

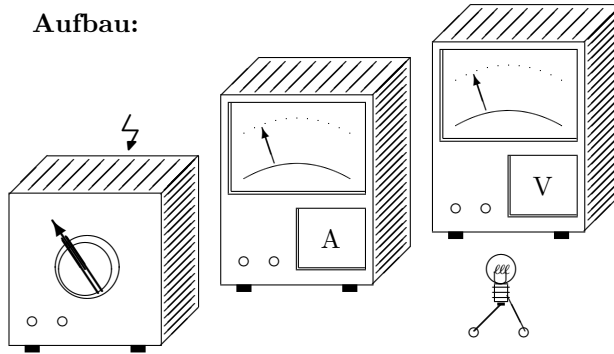
Eine Kurve, die den Zusammenhang zwischen der Stromstärke durch ein Bauteil und der an ihm anliegenden Spannung zeigt, heißt **U-I-Kennlinie** des Bauteils.

V 21.4 Aufnahme einer U-I-Kennlinie

zu ergänzen:

- Verkabelung
- Schaltbild

Aufbau:



Durchführung: Wir messen bei einem Glühlämpchen für verschiedene Spannungen die sich einstellende Stromstärke.

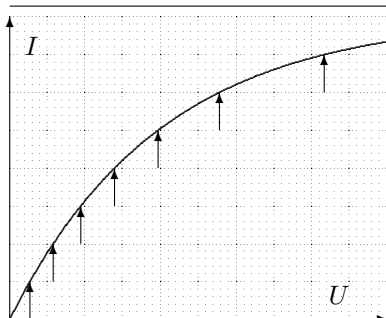
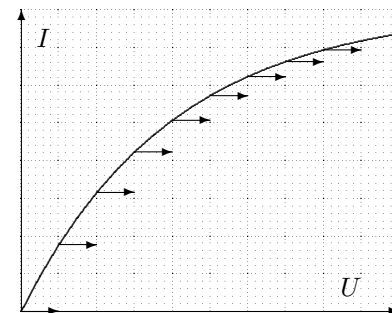
Messwerte: / Auswertung:

Glühlämpchen		

Ein großes Gitternetz für die Auswertung der Messwerte.

21.3.2 Der Widerstand und seine Einheit

Beim Glühlämpchen führt eine Verdoppelung der Spannung nicht immer zu einer Verdoppelung der Stromstärke. Stromstärke und Spannung sind bei ihm also nicht proportional, wie man vielleicht erwartet hätte. Wir schauen uns den Zusammenhang zwischen Spannung und Stromstärke anhand der nebenstehenden Kennlinie eines Lämpchens etwas genauer an: Gleiche Erhöhungen der Spannung ($\rightarrow$) führen zu immer kleineren Zuwächsen der Stromstärke. Zeichne diese Zuwächse der Stromstärke noch ein!

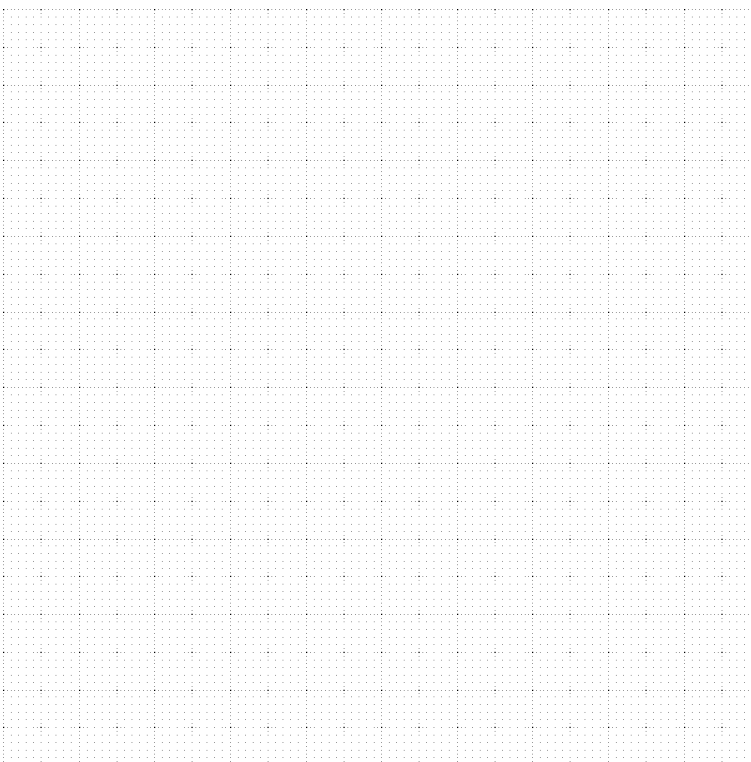


Anders betrachtet: Mit zunehmender Stromstärke werden immer größere – noch einzuzeichnende – Spannungszuwächse erforderlich, um gleiche Zuwächse der Stromstärke ($\uparrow$) zu erreichen. Für jedes zusätzliche Ampe-re braucht man immer mehr zusätzliche Volt. Wenn das Glühlämpchen ein Mensch wäre, würde man sagen: Es sträubt sich gegen jede weitere Erhöhung der Stromstärke immer mehr. Eine Erklärung dieses wach-senden Widerstandes drängt sich im wahrsten Sinne des Wortes offen-sicht-lich auf: Das Abflachen der Kennlinie könnte mit der zunehmenden Temperatur des immer heller glühenden Glühdrahtes zusammenhängen.

V 21.5 Kennlinien gekühlter Drähte

Aufbau: /Durchführung: Wie zuvor, jedoch wird anstelle des Glühlämpchens ein Draht verwendet, der in einem Wasserbad auf gleichbleibender Temperatur gehalten wird. (Für eine spätere Diskussion werden bereits hier Kennlinien von zwei unterschiedlich langen Drähten aufgenommen.)

	kurzer Dr.		langer Dr.	
U	I		I	



Ergebnis: Bei gleichbleibender Temperatur eines Drahtes sind

Stromstärke und Spannung zueinander.

Bei proportionalen Größen hat bekanntlich deren Quotient bei allen Messwertpaaren denselben Wert. Es ließen sich der Quotient I/U oder U/I betrachten. Obwohl I/U als Steigung der Kennlinien in obigen Schaubildern auftritt, wird meist U/I betrachtet. Das wäre die Steigung bei Vertauschung der Achsen.

Welche Bedeutung hat der Quotient U/I ? Er gibt ja an, wie viel Spannung pro bestimmter Stromstärke nötig ist. Beträgt der Quotient z.B. $30 \frac{\text{V}}{\text{A}}$, so heißt das

doch, dass man pro gewünschtem Ampere 30 V anlegen muss, für z.B. 4 A also 120 V.

Wenn nun bei einem Bauteil A (z.B. einem längeren Draht) dieser Quotient größer ist als bei einem anderen Bauteil B (z.B. dem kürzeren Draht), dann heißt dies doch, dass man für dieselbe Stromstärke an A mehr Spannung anlegen muss als an B. Bauteil A hätte demnach einen größeren **Widerstand** als B.

Einheit für Widerstand:

$$1 \text{ Ohm} = 1 \Omega :=$$

Fließt durch ein elektrisches Bauteil ein Strom der Stärke I , wenn an ihm die Spannung U anliegt, so heißt der Quotient aus Spannung U und Stromstärke I der **Widerstand** R des Bauteils. Kurz:

$$R :=$$

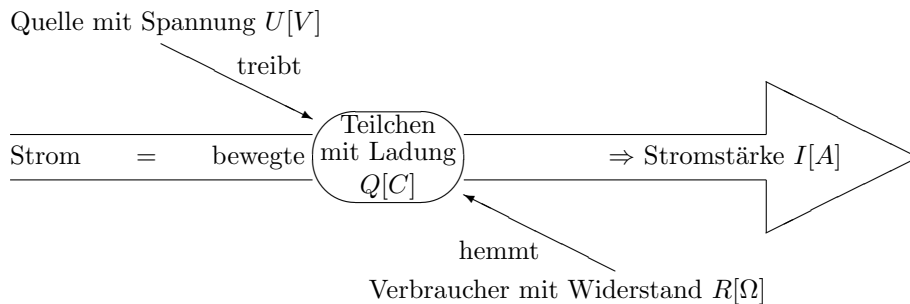
zu ergänzen:

- Formeltermine

Der Buchstabe R wurde als Abkürzung von z.B. *resistance* gewählt. Die Einheit Ohm wurde zu Ehren des Mathematik- und Physik-Lehrers Georg Simon Ohm (* 16.3.1789 in Erlangen, † 6.7.1854 in München) benannt.

Beim Begriff „Widerstand“ kann man das Adjektiv „ohmsch“ hinzufügen, zur Unterscheidung von anderen Widerständen, die wir hier nicht behandeln. Vorläufig haben wir es immer mit ohmschen Widerständen zu tun.

Die Definition der Größe Widerstand gilt unabhängig davon, ob Stromstärke und Spannung zueinander proportional sind.



21.3.3 Ohmsches Gesetz

Die von uns aufgenommene lineare Kennlinie eines gekühlten Drahtes zeigt:

Ohmsches Gesetz: Bei einem metallischen Leiter konstanter Temperatur sind Spannung und Stromstärke proportional zueinander. ($U \sim I$)

Anders gesagt: Der Widerstand eines metallischen Leiters bei konstanter Temperatur ist unabhängig von der Stromstärke oder Spannung.

Wir erkennen aus unseren Kennlinien weiterhin: Bei höherer Temperatur ist der Widerstand eines metallischen Leiters im Allgemeinen größer. Hierzu gibt es jedoch eine Ausnahme: Bei Leitern aus einer Legierung aus 55% Kupfer, 44% Nickel und 1% Mangan hängt der Widerstand praktisch nicht von der Temperatur ab. Diese Legierung nennt man daher **Konstantan**.

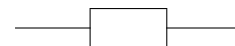
Zuweilen wird mit dem „ohmschen Gesetz“ nur der

Zusammenhang $U \sim I$ gemeint, ohne eine Aussage über die Voraussetzungen, unter denen er gilt. In diesem Sinne ist das ohmsche Gesetz natürlich kein Gesetz mehr. Man formuliert einfach nur gerne, ein Bauteil *erfülle das ohmsche Gesetz*, wenn bei ihm $U \sim I$ gilt. Genau genommen heißt das dann ja nur, dass $U \sim I$ gilt, wenn $U \sim I$ gilt.

Wir werden es in diesem Kapitel nur noch mit Widerständen zu tun haben, für die das ohmsche Gesetz gilt.

21.3.4 Das Bauteil Widerstand

Der Begriff (*ohmscher*) *Widerstand* wird in zwei Bedeutungen gebraucht: Einerseits für die soeben eingeführte Größe U/I , andererseits für ein elektrisches Bauteil, das um seines konstanten Widerstandes willen benutzt wird, z.B. um in einem Stromkreis mit einem empfindlichen Bauteil die mögliche Stromstärke zu begrenzen. Hier am Rand siehst Du das Schaltzeichen für einen Widerstand.

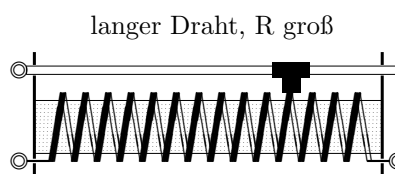
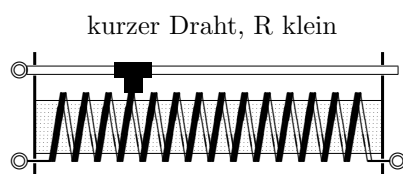


Schiebewiderstände

Nebenbei haben wir im Kennlinien-Versuch auf Seite 184 schon festgestellt, dass der Widerstand eines Drahtes mit seiner Länge zunimmt. Daraus ergibt sich eine recht einfache Möglichkeit, einstellbare Widerstände zu bauen:

Die **Schiebewiderstände** aus der Physik-Sammlung bestehen vor allem aus einem langen Draht. Der Handlichkeit wegen wickelt man ihn schraubenförmig

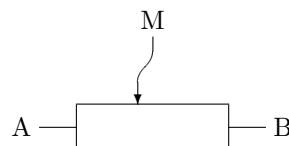
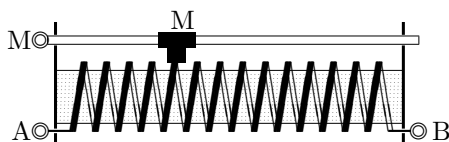
auf. Ein Kontaktplättchen kann quer zu den Windungen verschoben werden, so dass man zwischen ihm und einem Ende des Drahtes unterschiedlich lange Stücke des Drahtes abgreifen kann. Die beiden folgenden Bilder zeigen unterschiedlich lange Drahtstücke zwischen dem linken Ende des Schiebewiderstandes und dem Mittelabgriff.



Es lassen sich natürlich alle drei Anschlüsse eines Schiebewiderstandes benutzen. In seiner ganzen Länge A-B gesehen hat der Schiebewiderstand einen bestimmten, unveränderlichen (Gesamt-)Widerstand. Durch seinen Mittelabgriff M lassen sich an ihm aber auch die beiden Teil-Abschnitte A-M und M-B als veränderliche Widerstände nutzen. In einem Schaltbild zeichnet man den verschiebbaren Mittelabgriff als (geschlängelten) Pfeil an den gesamten Widerstand ein.

zu ergänzen:

- Stromwege
- Beschriftung

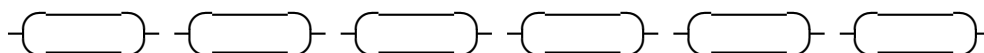


Schichtwiderstände

Widerstände sind oft hergestellt als Keramikzylinderchen, auf denen ein Metall aufgedampft ist. Je nach Metall und den Abmessungen der aufgedampften Schicht ergibt sich der Widerstand eines solchen **Schichtwiderstandes**. Die Größe des Widerstandes (in Ohm) wird in Form von farbigen Ringen aufgetragen. Man zählt die Ringe von dem Ring aus, der am weitesten außen liegt. Dabei steht jeder Ring für eine Ziffer, eine Anzahl Nullen oder eine Toleranzangabe. Eine Toleranz von z.B. 2% bedeutet, dass der tatsächliche Wert des Widerstandes durch Unregelmäßigkeiten bei der Herstellung um bis zu 2% vom angegebenen Wert abweichen kann. Ein heute gebräuchlicher Farben-Code ist dieser →

Farbe	1. Ring	2. Ring	3. Ring	4. Ring	5. Ring
schwarz	0	0	0	,0	
braun	1	1	1	0	±1%
rot	2	2	2	00	±2%
orange	3	3	3	000	
gelb	4	4	4	0 000	
grün	5	5	5	00 000	±0,5%
blau	6	6	6	000 000	
violett	7	7	7	0 000 000	
grau	8	8	8		
weiß	9	9	9		
gold				:10	±5%
silber				:100	

Beispiele:



22.1 Wirkung von Strom im Menschen

Elektrischer Strom kann durch unsere Körper fließen. Ein Mensch besteht nämlich zu etwa 60% aus Wasser, in dem unter anderem Salze gelöst sind. Und solche Lösungen sind bekanntlich Leiter. Fließt Strom durch solche Leiter, können dadurch allerlei chemische Umwandlungen stattfinden. Verändert sich jedoch die chemische Zusammensetzung in unseren Körperzellen, kann dies die Funktionstüchtigkeit der Zellen erheblich beeinträchtigen und sogar Zellen zerstören. Durch die Wärmewirkung des elektrischen Stromes kann es außerdem zu Verbrennungen kommen.

Außerdem arbeitet unser Nervensystem mit (ge-

ringen) elektrischen Spannungen und Strömen: Alle Sinneswahrnehmungen gelangen mit Hilfe elektrischer Spannungen über Nervenbahnen in das zentrale Nervensystem und werden dort auch mit elektrischen Vorgängen verarbeitet. Ebenso wird jedes Zusammenziehen von Muskeln durch elektrische Signale ausgelöst. Zusätzliche Ströme durch unseren Körper können das Nervensystem und damit den ganzen Organismus durcheinander bringen. Es kann z.B. zu starken Krämpfen und Herzrhythmus-Störungen kommen. Man sollte es daher vermeiden, „Strom zu kriegen“.

- Berühre keine Metallteile, (z.B. Kabelenden, Stecker) die möglicherweise an einer Spannungsquelle angeschlossen sind!
- Verwende keine elektrischen Geräte mit defekten Kabeln!
- Behandle Kabel sorgfältig: nicht verschmoren, nicht abscheuern, nicht zerschneiden, möglichst nicht knicken, Stecker nicht am Kabel aus der Steckdose ziehen!
- Wasser kann Strom leiten. Halte deshalb Feuchtigkeit von elektrischen Geräten und Steckdosen möglichst fern!

Du brauchst nun natürlich keine Angst davor zu haben, eine Taschenlampe bei Regen zu benutzen oder ein Fahrrad mit blank gescheuertem Kabel am Scheinwerfer zu benutzen. Die bei diesen Gelegenheiten auftretenden Spannungen unter 25 V sind beim Berühren ungefährlich.

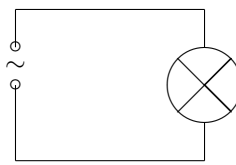
Spannungen ab etwa 25 V können jedoch Schaden an Deinem Körper anrichten! Du solltest Dich jedenfalls vorsehen beim 230 V-Wechselstromnetz im Haushalt und bei Hochspannungsleitungen und Oberleitungen draußen. Auch der nicht sichtbare, fast immer vorhandene, hauchdünne Feuchtigkeitsfilm auf einer Drahtenschnur kann Strom gut genug leiten, um Dich ggf. unter Hochspannung zu setzen!

Hoffentlich fragst Du Dich jetzt, warum es schon gefährlich sein kann, *eine* Hochspannungsleitung zu berühren, oder an *ein* blankes Kabelende zu kommen, wo doch für einen Stromfluss immer *zwei* Anschlüsse nötig sind. Tatsächlich ist es recht ungefährlich, z.B. nur eine Hochspannungsleitung zu berühren, wie jeder darauf sitzende Vogel demonstriert. Wir Menschen berühren aber praktisch immer auch den Erdboden. (oder sind zumindest nahe genug an ihm) Und genau dieser Erdkontakt ist der vermisste zweite elektrische Anschluss! Der Erdboden kann nämlich schon aufgrund seiner Feuchtigkeit elektrisch leitfähig sein. Wieso die Erde in diesem Zusammenhang eine Rolle spielt, erfährst Du im nächsten Abschnitt.

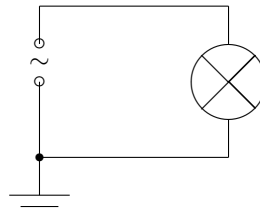
22.2 Das Schuko-System

Wir betrachten stark vereinfacht ein Kraftwerk („~“), das über zwei Leitungen ein elektrisches Gerät mit Strom versorgt.

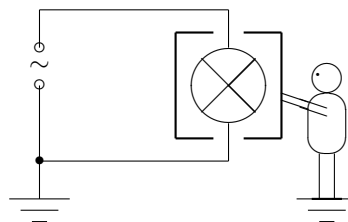
Tatsächlich versorgt auf fast dem gesamten europäischen Festland eher die Gesamtheit aller Kraftwerke die Gesamtheit aller Verbraucher über ein kompliziertes Netzwerk mit drei Wechselstromphasen.



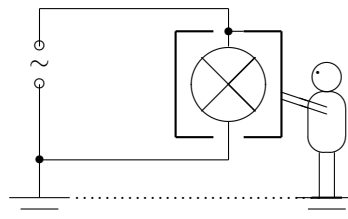
Eine der beiden Leitungen ist mit der Erde (übliches Symbolsymbol links unten) verbunden. Diese *geerdete* Leitung heißt **Null-Leiter** oder **Neutral-Leiter**. Die andere Leitung heißt **Phasen-Leiter**, gerne auch kurz „Phase“.



Nehmen wir nun noch an, dass der Verbraucher ein leitendes Gehäuse besitzt, das der Benutzer berühren kann. Und der Benutzer steht typischerweise auf der Erde – seltener auf anderen Planeten.



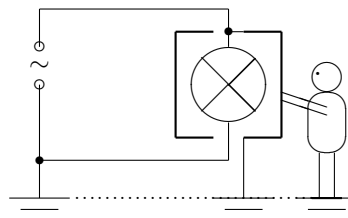
Bekommt nun z.B. durch ein schadhafte Kabel der Phasenleiter Kontakt zu dem Gehäuse, ist der Schaden da: Der Stromkreis ist vom Kraftwerk über den Phasenleiter, das Gehäuse, den Benutzer und die Erde zum Kraftwerk geschlossen. Den Benutzer durchfließt Strom!



zu ergänzen:

- Stromweg durch Verbraucher blau

Damit in solchen Unfällen dem Benutzer nichts passieren kann, sind leitende Gehäuse meist über eine eigene Leitung geerdet. Strom fließt dann fast nur über diesen gut leitenden **Schutzleiter** und nicht über den viel schlechter leitenden Benutzer oder den Verbraucher.



zu ergänzen:

- Stromweg durch Benutzer rot

Haushaltsgeräte werden also über 3 Leitungen angeschlossen: den Phasen- den Null- und den Schutzleiter. Ihre Kabel haben demnach oft drei Adern.

In einer Steckdose befinden sich der Phasen- und der Nulleiter hinter je einem der beiden Löcher. Der Schutzleiter ist mit den federnden Metallbügel am Rand der Dose verbunden. Man nennt diese Bügel deshalb die *Schutz-Kontakte* und das ganze System **Schuko-System**.

zu ergänzen:

- Stromweg durch Schutzleiter rot

Um zu prüfen, welches der beiden Löcher den Phasenleiter enthält, gibt es *Phasenprüfer*. Dies sind Schraubendreher mit eingebaute Glimmlampe. Die Lampe leuchtet, wenn die Spitze des Prüfers Kontakt zum Phasenleiter bekommt und das andere Ende des Prüfers Kontakt zum Benutzer (meist seinem Daumen) hat. Dabei wird zwar ein Stromweg vom Phasenleiter über die Glimmlampe und den Benutzer zur Erde ge-

schlossen, aber der Phasenprüfer lässt nur ungefährlich kleine Stromstärken zu. Man steckt also den Phasenprüfer mit oben aufgelegtem Daumen in ein Loch der Steckdose. Leuchtet dann die eingebaute Lampe, so hat man den Phasenleiter gefunden. Leuchtet die Lampe nicht, sollte man daraus nicht schließen, den Nulleiter gefunden zu haben: Es könnte auch der Prüfer defekt sein!

Phasenleiter:	schwarz oder braun
Schutz-Leiter:	gelb-grün
Null-Leiter:	blau

In einem Kabel der 230 V-Installation werden nebenstehende Farben für die drei benötigten Adern verwendet: (Was man immer so machen sollte, worauf man sich aber nie verlassen sollte!)

22.3 Sicherungen

Entsteht — wie im eben beschriebenem Un-Fall — ein direkter Stromweg von einem Pol der Quelle ohne Verbraucher zum anderen Pol der Quelle, so nennt man dies einen **Kurzschluss**. Bei einem Kurzschluss ist im Stromkreis kein Gerät, das den Stromfluss nennenswert hemmt. Nur die Leitungen selbst hemmen ein wenig. Die Stromstärke würde demnach auf sehr hohe Werte ansteigen. Leitungen würden glühen und Brände verursachen.

Man verhindert daher zu hohe Stromstärken durch Sicherungen. Du kennst bereits Schmelzsicherungen. Jeder Stromanschluss eines Hauses an das 230 V-Netz ist mit einer Schmelzsicherung ausgestattet. Da durchgeschmolzene Sicherungen nicht wieder verwendet werden können, verwendet man daneben auch Sicherungen, die die magnetische Wirkung des Stromes ausnutzen

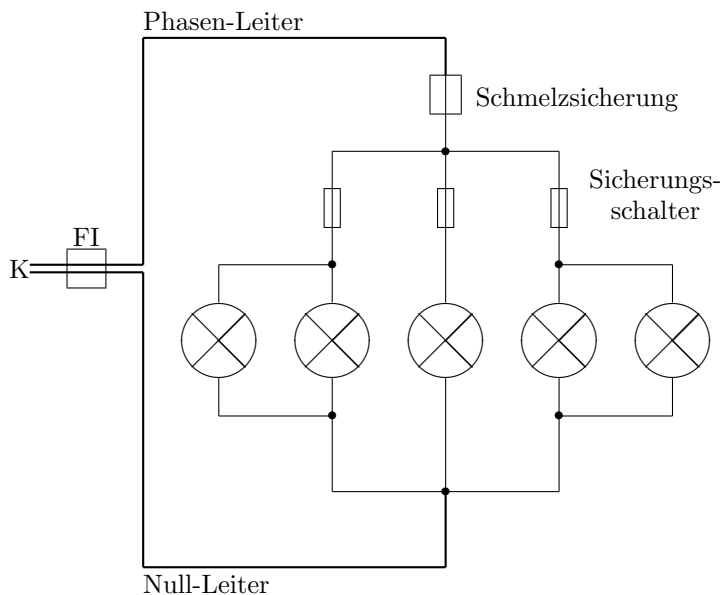
und viele Male benutzt werden können.

Der zu begrenzende Strom durchfließt einen Elektromagneten und einen Schalter. Bei zu großer Stromstärke wird der Elektromagnet stark genug, um den Schalter zu öffnen und damit den Stromfluss zu unterbinden. Der Schalter rastet in dieser „offen“-Stellung ein und kann nur von Hand wieder geschlossen werden.

Mit solchen Sicherungsschaltern werden meist einzelne Zimmer, Gebäudeteile oder starke Geräte bis zu 16 A abgesichert. Ein Kurzschluss z.B. im Wohnzimmer lässt dann nur dort den Strom ausfallen. Natürlich muss diese magnetische Sicherung schon bei geringeren Stromstärken und daher früher ansprechen als die Schmelzsicherung. Andernfalls schaltet die Schmelzsicherung den Strom im ganzen Haus ab.

22.4 FI-Schalter

Zu jedem Haus führt (vereinfacht) vom Kraftwerk K eine Phasen- und eine Null-Leitung. (Tatsächlich sind es drei Phasen.) Diese Haupt-Zuleitungen verzweigen sich zu den einzelnen Verbrauchern des Hauses.



Diese Abbildung stellt das Leitungsnetz eines Hauses vereinfacht (z.B. ohne Schalter) dar. Nach der Knotenregel sind die Stromstärken in dem vom Kraftwerk kommenden Phasen- und Null-Leiter gleich. Und so soll es auch sein. Stellt man einmal fest, dass im Null-Leiter weniger Strom fließt als im Phasenleiter, so kann man daraus schließen, dass ein Teil des Stromes einen anderen Weg als den Null-Leiter genommen hat. In Frage kommen der Schutzleiter und andere Verbindungen zur Erde. In jedem Falle ist dann aber ein Gerät im Haus nicht ordnungsgemäß in Betrieb. Gefahr!

Ein **FI-Schalter** (auch: „Fehlerstromschutzschalter“, FI steht also für Fehler-*I*) vergleicht ständig die Stromstärken in Phasen- und Null-Leiter. Stellt er eine größere Abweichung (größer als z.B. 30 mA) fest, schaltet er den Strom ab. Letzten Endes wird der FI-Schalter betätigt durch das Magnetfeld, das Phasen- und Null-

Leiter nebeneinander verlegt zusammen erzeugen. Im störungsfreien Fall fließen die Ströme in diesen Leitungen gleich stark in entgegengesetzten Richtungen, so dass sich insgesamt kein Magnetfeld ergibt. Erst im Fehlerfall heben sich die Magnetfelder der einzelnen Leiter nicht mehr auf.

Elektrische Ladung

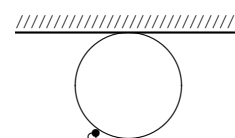
Elektrischen Strom stellen wir uns schon seit dem 7.Schuljahr als einen Strom elektrisch geladener Teilchen vor. Ab heute wollen wir die elektrische Ladung als Eigenschaft von Teilchen genauer kennenlernen. Wir werden zunächst sehen, welche Beobachtungen es überhaupt nahelegen, einem Körper neben Eigenschaften wie z.B. Masse, Volumen oder Farbe auch eine Eigenschaft „elektrische Ladung“ zuzuordnen, mit der sich dann elektrische Phänomene gut deuten lassen.

23.1 Elektrostatische Kräfte

23.1.1 Erste Erfahrungen

...mit elektrischer Ladung hast Du sicher schon gesammelt, wenn auch vielleicht unbewusst. Hier einige Beispiele:

- Ein Kamm zieht nach dem Kämmen Haare und Papierschnipsel an.
- Beim Ausziehen eines Woll-Pullovers über den Kopf hört man ein Knistern und sieht vielleicht auch Fünkchen.
- Eine Plastikfolie haftet ohne Kleber, Wasserfilm o.ä. mit einer anderen Folie zusammen oder auf der Tischoberfläche.
- Beim Aussteigen aus dem Auto „bekommt man einen Schlag“.
- Hänge doch gerade mal einen Luftballon an die Decke, indem Du ihn (mit möglichst viel seiner Oberfläche) an Deinen Haaren reibst und dann an die Decke hältst!



Schon früh hat man auch an Bernstein ähnliche *Kraftwirkungen* bemerkt, wenn man ihn z.B. an einem Fell gerieben hat. Die Bezeichnung „Elektrizität“ kommt daher, dass Bernstein auf Alt-Griechisch „elektron“ heißt.

Bernstein ist versteinertes urzeitliches Harz. Aus Bernstein wird seit der Antike viel Schmuck hergestellt. Bernstein ist auch wissenschaftlich interessant: Manchmal sind in einem Harztropfen z.B. Insekten eingeschlossen und damit bis in unsere Zeit im Bernstein konserviert worden.

23.1.2 Experimentelle Untersuchung

Nach den bisherigen Erfahrungen lässt sich eine elektrische Aufladung

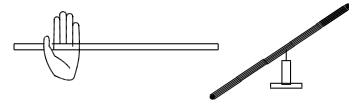
durch Reibung erzeugen

und

an ihren Kraftwirkungen erkennen.

V 23.1 Kräfte zwischen geriebenen Stäben

Aufbau: /Durchführung: Wir untersuchen die Kräfte zwischen zwei Stäben aus Kunststoff oder Glas, nachdem sie mit einem Fell gerieben worden sind. Um die Kraftwirkung zu sehen, wird der eine Stab auf einer Nadelspitze drehbar gelagert, der andere wird ihm mit der Hand angenähert.



Beobachtung:

	1.Stab		2.Stab		beobachtete Kraftwirkung
	Material	gerieben?	Material	gerieben?	
0)	egal	nein	egal	nein	keine
1)	KStoff	nein	KStoff	ja	
2)	KStoff	ja	KStoff	ja	
3)	KStoff	ja	Glas	ja	
4)	Glas	ja	Glas	ja	
5)	Glas	ja	KStoff	nein	

Ergebnis:

0)→1) Offensichtlich ändert das Reiben den Zustand eines Stabes so, dass neue Kräfte auftreten.

Man sagt: Der Stab wird

2)..4) Da wir je nach Material unterschiedlich gerichtete Kraftwirkungen sehen, müssen wir zu ihrer Erklärung (zumindest) 2 Arten elektrischer Aufladung unterscheiden,

die man als und bezeichnet.

2),4) Wir gehen natürlich davon aus, dass Stäbe aus demselben Material, die auf gleiche Weise mit demselben Fell gerieben werden, dieselbe Art der Aufladung erfahren.
Die Beobachtungen 2) und 4) ergeben damit:

Körper, die mit Vorzeichen aufgeladen sind,

3) Wir erkennen hier ganz entsprechend:

Körper, die mit Vorzeichen geladen sind,

1),5) Die Anziehung zwischen einem beliebig geladenen Körper und einem ungeladenen Körper besprechen wir später ausführlich. (Seite 198)

Kräfte, die aufgrund ruhender elektrischer Ladungen wirken, heißen elektrostatische Kräfte oder Coulomb-Kräfte.

Wenn sich Ladungen bewegen, kommen weitere Effekte ins Spiel, die man „elektrodynamisch“ nennt. (von griechisch *dynamis*=Kraft, eine Größe, die ja viel mit Bewegung zu tun hat. Hinter „statisch“ steckt das lateinische *stare*=stehen.)

23.2 Verteilung von Ladung

In weiteren Versuchen stellt man fest, dass elektrische Ladung von einem Körper bei Berührung auf einen anderen Körper übergehen kann. Bei einem metallischen Körper kann sich die Ladung außerdem auf dem ganzen Körper verteilen. Zum Vergleich: Ein geriebener Kunststoff-Stab zeigt nur an den geriebenen Flächen

die oben beschriebenen Kraftwirkungen.

Für weitere Versuche ist es praktisch, statt geriebener Stäbe ein Hochspannungsnetzgerät zu verwenden. Von seinem Plus- und Minuspol kann man über Kabel elektrische Ladung zu anderen metallischen Körpern übertragen.

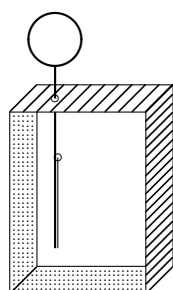
23.2.1 Elektroskop

Die Übertragung und Verteilung elektrischer Ladung erkennt man gut an einem **Elektroskop**, das im folgenden Bild links im ungeladenen und rechts im geladenen Zustand gezeichnet ist.

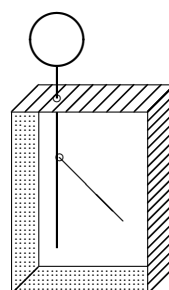
In das Innere eines Blechrahmens ragt ein metallener Stab, ohne metallisch mit dem Rahmen verbunden zu sein. („isoliert“) Am Stab ist unten ein metallischer Zeiger beweglich eingehängt. Oben kann z.B. eine Me-

tallkugel aufgesteckt werden.

Berührt man mit einem geladenen Körper die Kugel, so geht elektrische Ladung auf diese Kugel über und verteilt sich auf ihr, dem Stab und dem Zeiger. Wegen der gleichartigen Ladung von Stab und Zeiger stoßen diese beiden sich ab: der Zeiger schlägt aus. (Keine Angst! ;-)



Metallkugel
Isolation
Blechgehäuse
Stab
Zeiger



zu ergänzen:

- Skala
- Pfeile der Beschriftung

Naheliegenderweise ist der Ausschlag umso größer, je stärker die Aufladung ist. Mit dem Elektroskop haben wir folglich ein primitives Nachweisgerät für elektrische Ladung zur Verfügung.

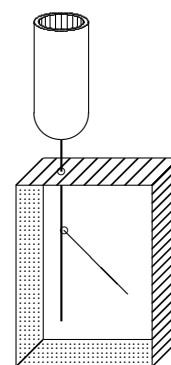
Genauer gesagt: ein Messgerät dafür, wie gut das angeschlossene Objekt Ladung auf das Elektroskop übertragen kann. Tatsächlich zeigt ein Elektroskop nicht an, wie viel Ladung das angeschlossene Objekt trägt, sondern wie viel Spannung es gegenüber der Erde hat.

23.2.2 Ladungsverteilung auf Metallkörpern

Eine wichtige Eigenheit bei der Verteilung von Ladung auf einem Metallkörper zeigt sich an einem aufgeladenen Elektroskop mit aufgestecktem Metallbecher („Faraday-Becher“): Berührt man mit einer ungeladenen Metallkugel die Außenseite des Bechers, geht der Ausschlag des Elektroskopes deutlich zurück. Berührt man den Becher jedoch innen, bleibt der Ausschlag eher gleich. Offensichtlich trägt der Becher nur auf seiner Außenseite Ladung. Dies ist plausibel, wenn man bedenkt, dass gleichnamige Ladungen sich ja gegenseitig abstoßen und sich daher gegenseitig nach außen drücken.

Aus demselben Grund lässt sich auch gut Ladung von innen her auf den Becher bringen: Berührt man die Innenseite mit einer geladenen Metallkugel, geht auch deren Ladung vollständig nach außen auf den Becher über.

Die Innenwände eines metallischen Hohlkörpers können generell nie elektrostatisch geladen sein. In diesem Zusammenhang nennt man einen solchen Hohlkörper einen **Faraday-Käfig**. Sein Inneres ist stets von allen elektrostatischen Einflüssen aus der Umgebung abgeschirmt. Die Abschirmung bleibt sogar weitgehend erhalten, wenn die Wände – nicht allzu große – Löcher haben, wie es bei einem normalen Käfig zwischen den Gitterstäben oder bei einem Auto an den Fenstern der Fall ist.



*Michael Faraday,
Londoner Physiker und
Chemiker,
* 22.9.1791
† 25.8.1867*

Die gegenseitige Abstoßung gleichnamiger Ladungen treibt die Ladung auch bevorzugt an Stellen in der Oberfläche eines Metallkörpers, die nach außen herausragen: in Kanten, Ecken und besonders in Spitzen.

So kann Ladung aus den Spitzen eines stark geladenen Körpers sogar in die Luft übertreten. Diesen Vorgang nennt man daher **Spitzenentladung**.

23.2.3 Mengencharakter

Zur Übertragung von Ladung von Körper zu Körper: Der Begriff „Übertragung“ kann allgemein auf mindestens zwei verschiedene Arten verstanden werden:

- ① Wenn jemand seinen Freunden Pralinen überträgt, hat er selbst nachher weniger Pralinen. Durch das Übertragen wird ja die Gesamtzahl an Pralinen nicht verändert. Es wird also eine in ihrem Gesamtumfang unveränderliche Menge von irgendetwas nur umverteilt. Ein Beispiel aus der Physik ist die Übertragung von Energie nach dem Energieerhaltungssatz.
- ② Bei der Übertragung z.B. einer Krankheit (Ansteckung) gibt es vorher eine kranke Person und nachher eine zweite ebenso kranke Person. Es geht also nicht um eine feststehende Menge von irgendetwas, sondern um eine Eigenschaft, die in beliebigem Ausmaß, zunehmend und abnehmend vorhanden sein kann. Physikalisches Beispiel: Mit einem Magneten lässt sich ein Weicheisenstück magnetisieren, ohne dass der Magnet selbst seine Magnetisierung einbüßt. Man hat nachher zwei Magnete.

(Anderes Beispiel:
*Geteiltes Leid ist
halbes Leid.* ;-)

(Hierzu passend:
*Geteilte Freude ist
doppelte Freude.* ;-)

?

Ist „Ladung“ nun ansteckend oder mengenartig?

V 23.2 Zusammenbringen von Ladung

Aufbau: Zwei Elektroskope mit aufgestecktem Faraday-Becher, Metallkugel am Stiel, Hochspannungsnetzgerät

Durchführung: Die Metallkugel wird am Netzgerät positiv (negativ) geladen und an die Innenseite eines Bechers gehalten.

Beobachtung:

Durchführung: Die Kugel wird nun sofort an das zweite Elektroskop gehalten.

Beobachtung:

Ergebnis: Die Kugel hat nach der Übertragung ihrer Ladung an das erste Elektroskop selbst

Durchführung: Die Kugel wird wiederholt am Netzgerät aufgeladen und an das erste Elektroskop gehalten.

Beobachtung:

Ergebnis: Bringt man zu einer Ladung weitere Ladung desselben Vorzeichens, so nimmt die Aufladung zu.

Durchführung: Das Elektroskop wird zunächst positiv (negativ) aufgeladen. Die Metallkugel wird dann mehrmals negativ (positiv) aufgeladen und an die Innenseite des Bechers gehalten.

Beobachtung:

Ergebnis:

Ladung hat Mengencharakter, sie ist eine Erhaltungsgröße.

Andere Erhaltungsgrößen: Masse und Energie (Dir bereits bekannt) Impuls (kannst Du in der Oberstufe kennenlernen), Drehimpuls (z.Zt. kein Schulstoff)

Wir gehen davon aus, dass die Gesamt-Ladung des ganzen Universums unveränderlich ist. (vermutlich gleich Null) Folglich ist es auch nicht möglich, positive oder negative Ladung einzeln zu erzeugen, weil ja dadurch die Gesamt-Ladung verändert würde. Es können

immer nur beide Ladungsarten gemeinsam in gleichen Maßen hervorgebracht werden. Wenn wir also z.B. einen Stab durch Reiben an einem Fell positiv aufladen, wird dabei das Fell sicher negativ aufgeladen.

23.2.4 Erdung

In der Elektrizitätslehre spielt die Erde eine besondere Rolle. Sie ist durch den Feuchtigkeitsgehalt ihrer Oberfläche und die darin enthaltenen Mineralien ausreichend leitfähig, dass elektrische Ladung sich auch in ihr verteilen kann.

Bringt man nun einen elektrisch geladenen Körper in (leitfähigen) Kontakt zur Erde, so verteilt sich die Ladung des Körpers auf Körper und Erde. Aufgrund des Größenverhältnisses zwischen Erde und Körper ist klar, wie viel Ladung noch auf dem Körper bleibt: praktisch keine mehr! Daher kann man einen einzelnen Körper entladen, indem man ihn „erdet“.

Schaltzeichen
Erdung



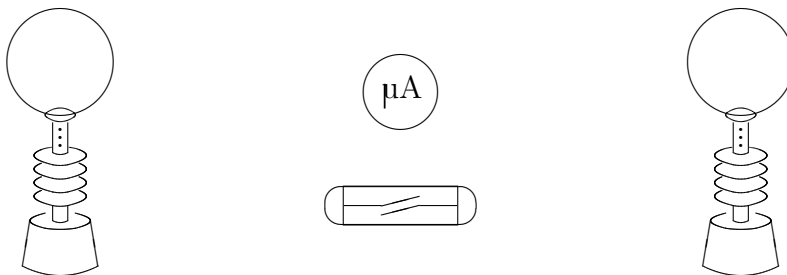
23.3 Elektrischer Strom

23.3.1 Ladungsausgleich

Elektrische Ladungen verschiedener Vorzeichen haben immer das Bestreben, sich zu kompensieren: Die Anziehungskräfte zwischen positiven und negativen Ladungen wirken stets daraufhin, dass diese Ladungen zusammenkommen. Sind sie aber zusammengekommen, so heben sie sich in ihren Wirkungen auf die Außenwelt gegenseitig auf.

V 23.3 Ladungsausgleich

Aufbau: zwei mit unterschiedlichen Vorzeichen geladene Kugeln, isoliert aufgestellt



Durchführung: die beiden Kugeln werden über

Beobachtung:

Ergebnis: Die elektrische Ladung, die sich hier von einer Kugel zur anderen bewegt, stellt einen dar.

zu ergänzen:

- Vorzeichen
- Verkabelung über empfindliches Amperemeter oder Glühlampe

23.3.2 Definition der Größe Ladung

Was im 7.Schuljahr nur Behauptung war, ist jetzt bestätigt: Elektrischer Strom ist bewegte elektrische Ladung. Dabei gehen wir davon aus, dass die Stromstärke umso größer ist, je mehr Ladung pro Zeit durch einen Leiterquerschnitt fließt. Bezeichnet man die Stromstärke mit I , die Ladung mit Q und die Zeit mit t , so gilt die nebenstehende Gleichung.

$$I = \frac{Q}{t}$$

Entsprechend könnte man definieren, dass die Stromstärke 1 Ampere beträgt, wenn *soundsoviel* Ladung pro Sekunde fließt. Um das „soundsoviel“ anzugeben, benötigt man jedoch *vorher* eine Einheit für Ladungen. Es scheint ja auch natürlicher zu sein, zuerst die grundlegende Teilcheneigenschaft „Ladung“ mit einer Einheit zu versehen, bevor man nach der Größe „Stromstärke“ schaut, die sich ja erst einstellt, wenn bereits vorhandene Ladung sich bewegt.

Elektrische Ladungen lassen sich allerdings nicht so gut messen wie Stromstärken. Auch wenn die als Anmerkung folgende Definition des Ampere nicht sonderlich einfach und praktisch umsetzbar klingt, sind die Kräfte zwischen geeigneten stromdurchflossenen Leitern in der Praxis doch recht gut messbar. (→ Drehspulinstrument)

Deshalb geht man heute zur Definition des international gültigen Einheitensystems (*Système International d'Unités*, kurz: SI) anders herum vor: Erst wird das Ampere (neben Meter, Sekunde, Kilogramm, Kelvin, Mol und Candela) als eine Basiseinheit festgelegt, dann aus diesen die Einheit Coulomb für Ladungen abgeleitet.

André Marie Ampère
frz. Physiker
und Mathematiker,
* 22.1.1775
in Polémieux/Lyon,
† 10.6.1836
in Marseille

Hier die Definitionen der sieben Basiseinheiten:

- Die **Sekunde** ist das 9 192 631 770-fache der Periodendauer der Strahlung, die dem Übergang zwischen den beiden Hyperfeinstrukturniveaus des Grundzustandes des Nuklids $^{133}_{55}\text{Cs}$ entspricht.
- Das **Meter** ist jene Distanz, die das Licht in $1/299\,792\,458$ einer Sekunde im Vakuum durchläuft.
- Das **Kilogramm** ist gleich der Masse des internationalen Kilogrammprototyps.
- Das **Ampere** ist die Stärke eines konstanten elektrischen Stromes, der, durch zwei parallele, geradlinige, unendlich lange und im Vakuum im Abstand von 1 Meter voneinander angeordnete Leiter von vernachlässigbar kleinem, kreisförmigen Querschnitt fließend, zwischen diesen Leitern je 1 Meter Leiterlänge die Kraft $2 \cdot 10^{-7}$ Newton hervorrufen würde.
- Das **Kelvin** ist der 273,16te Teil der absoluten Temperatur des Tripelpunktes des Wassers.
- Das **Mol** ist die Stoffmenge eines Systems, das aus ebensoviel spezifizierten Einzelteilchen besteht, wie Atome in 12 Gramm des Kohlenstoffnuklids $^{12}_6\text{C}$ enthalten sind.
- Die **Candela** ist die Lichtstärke einer Strahlungsquelle, welche monochromatische Strahlung der Frequenz $540 \cdot 10^{12}$ Hertz aussendet, und deren Strahlstärke in dieser Richtung $1/683$ Watt durch Steradian beträgt.

Charles Augustin
de Coulomb,
frz. Physiker,
* 14.6.1736
in Angoulême,
† 23.8.1806
in Paris

Bei einem elektrischen Strom der Stärke I fließt in der Zeit t durch einen Leiterquerschnitt die Ladung Q mit

$$Q =$$

Entsprechend definierte Einheit:

$$1 \text{ Coulomb} = 1 \text{ C} :=$$

zu ergänzen:

- Formeltermine

23.4 Die Hauptakteure: Elektronen

23.4.1 Glühelektrischer Effekt

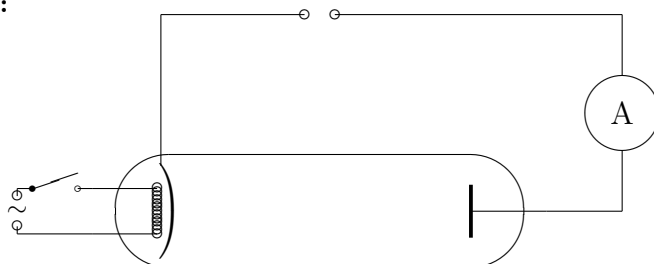
Der vorige Versuch auf Seite 195 lässt offen, ob

- negative Ladung von der negativen Kugel zur positiven fließt, oder
- positive Ladung von der positiven Kugel zur negativen, oder
- beide Ladungsarten sich bewegen.

Dies zu klären hilft folgender Versuch.

V 23.4 Röhren-Diode

Aufbau:



zu ergänzen:

- Beschriftung
- Polung(en)

Durchführung: Wir messen die Stromstärke in obigem großen Stromkreis jeweils mit und ohne Heizung, sowie mit verschiedenen Polungen der Spannungsquelle oben.

Beobachtung:

a) ohne Heizung, $\oplus$ an Katode:

b) ohne Heizung, $\ominus$ an Katode:

c) mit Heizung, $\oplus$ an Katode:

d) mit Heizung, $\ominus$ an Katode:

Die Beobachtungen a) und b) bei ausgeschalteter Heizung sind klar, weil wir ja einfach nur einen offenen Stromkreis vor uns haben. Für einen Stromfluss durch die Leiter-Lücke in der Diode müssten Ladungsträger aus einer der beiden Elektroden austreten und in die andere eintreten. Dies ist jedoch offensichtlich nicht so einfach möglich.

Offensichtlich ist dies aber im Fall d) möglich. Es könnten bei der dabei vorliegenden Polung negative Ladungsträger von der Katode zur Anode fliegen, oder positive in umgekehrter Richtung. Entsprechend müsste die stark erhöhte Temperatur der Katode den Austritt negativer Teilchen bzw. den Eintritt positiver Teilchen begünstigen.

Wir stellen uns unter erhöhter Temperatur vor, dass die Teilchen eines Körpers sich heftiger bewegen. Dies begünstigt sicherlich, dass einzelne Teilchen aus dem Körper ausreißen. Wir kennen dies ja vom Verdampfen einer Flüssigkeit.

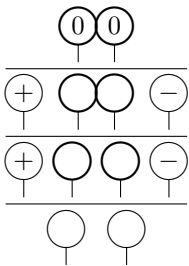
Folglich erklären wir den Stromfluss durch die Diode wie folgt:

Bei Erhitzung treten aus der Katode negative Ladungsträger aus. Ist nun die Katode negativ und die Anode positiv gepolt, so werden die ausgetretenen Ladungsträger zur Anode hin beschleunigt. Dort angelangt treten sie in die Anode ein und setzen ihren Weg durch den Stromkreis in den metallischen Leitern fort.

- In jedem Körper gibt es positive und negative Ladungsträger, nämlich die stets mit einer „Elementarladung“ (das sind $1,602 \cdot 10^{-19} \text{ C}$) positiv geladenen Protonen und die stets mit dem gleichen Betrag negativ geladenen Elektronen.
- Ein elektrisch neutraler Körper enthält gleiche Anzahlen von Protonen und Elektronen.
- Die Protonen können sich in einem festen Körper nicht ausreichend frei bewegen, um zu einem elektrischen Strom beizutragen.
- Metalle enthalten leicht bewegliche Elektronen. Elektrische Ströme – und insbesondere auch alle Ladungs- oder Entladungsvorgänge – in Metallen sind stets Bewegungen von Elektronen.
- Negative Aufladung eines festen Körpers bedeutet daher Elektronenüberschuss, positive Aufladung bedeutet Elektronenmangel gegenüber der unveränderten Anzahl an Protonen.
- Bei höherer Temperatur können Elektronen ein Metall besser verlassen. Dies nennt man den „glüh elektrischen Effekt“.

23.4.2 Influenz

Ladungstrennung in einem Körper durch Annäherung eines geladenen anderen Körpers nennt man **Influenz**.



Influenz zeigt sich in nebenstehender Bildergeschichte. Erste Bildzeile: Zwei anfangs ungeladene Metallkugeln berühren sich. Zweite Zeile: Nähert man dem Paar von links (und/oder rechts) einen positiv (bzw. negativ) geladenen Körper an, so treten Elektronen aus der einen in die andere Kugel über. Trage ab der zweiten Bildzeile noch die Vorzeichen der so entstehenden Aufladungen in die Kugeln ein! Dritte Zeile: Trennt man die beiden Kugeln voneinander, können die übergetretenen Elektronen nicht wieder zurückfließen, wenn zur vierten Zeile die seitlichen Ladungen wieder entfernt werden.

Mit Influenz lässt sich auch die Anziehung zwischen einem geladenen und einem ungeladenen Körper (wie zu Beginn des Kapitels beschrieben) erklären. Selbst wenn es sich dabei nicht um metallische Körper mit leicht beweglichen Elektronen handelt, so können sich die gebundenen Elektronen doch ein wenig gegenüber den Protonen des Körpers verschieben.

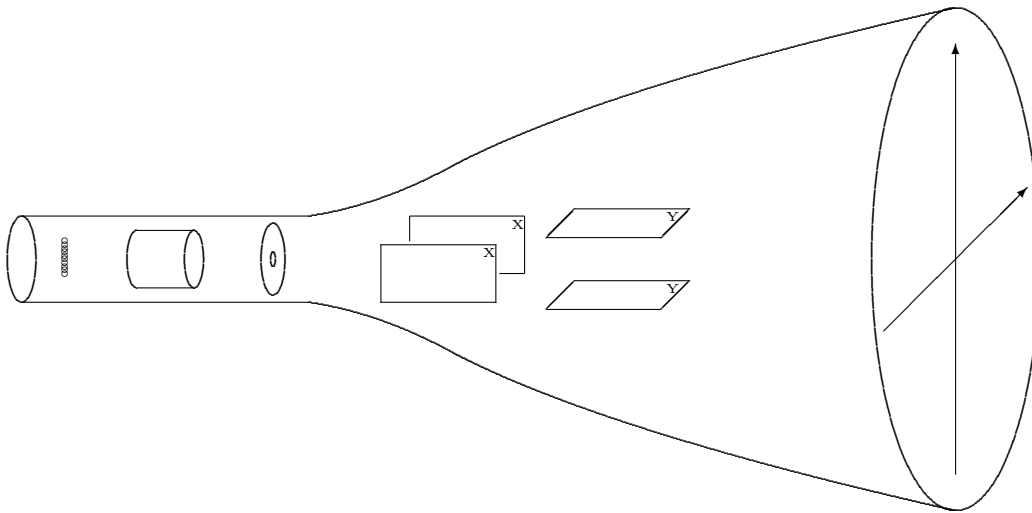
Nähert man einem ungeladenen Körper U einen negativ geladenen Körper N von rechts an, so werden die Elektronen des ungeladenen Körpers U ein wenig nach links gedrückt. Damit sind aber diese negativen Ladungen in U durchschnittlich etwas weiter vom negativ geladenen Körpers N entfernt als U's Protonen. Die Abstoßung zwischen den negativen Ladungen ist damit schwächer als die Anziehung zwischen U's positiven und N's negativen Ladungen.



23.4.3 Braunsche Röhren

- ... wurden 1897 vom deutschen Physiker Karl Ferdinand Braun (* 6.6.1850 in Fulda, † 20.4.1918 in New York) erfunden
- ... dienen dazu, elektrisch gesteuert auf einem Leuchtschirm ein Bild zu zeichnen, das von außerhalb der Röhre betrachtet wird

- ...nutzen wie eine Röhren-Diode den glühelektrischen Effekt zur Freisetzung von Elektronen aus der Glüh-Katode
- ...bilden aus diesen Elektronen einen Elektronenstrahl (durch Beschleunigen zwischen Katode und Anode, Bündeln mit Hilfe des Wehneltzylinders und Ausblenden an der Loch-Anode)
- ...lenken diesen Strahl mit Hilfe elektrisch aufladbarer Metallplatten zu einem gewünschten Punkt auf dem Schirm, der dank seiner chemischen Beschichtung an der Auftreffstelle aufleuchtet



Liegt an den Plattenpaaren keine Spannung an, so passiert der Elektronenstrahl die gesamte Anordnung geradlinig und trifft auf die Mitte des Schirmes, wo wir dann einen Leuchtfleck erkennen.

Legt man an das Plattenpaar X eine Spannung an, z.B. mit Pluspol an der vom Betrachter aus linken Platte (in der Zeichnung im Vordergrund), so wandert der Leuchtfleck nach links, und zwar umso weiter, je größer die Spannung ist. Bei umgekehrter Polung wandert der Leuchtfleck nach rechts.

Analog bewirkt eine Spannung am Plattenpaar Y je nach Polung eine Ablenkung nach oben oder unten.

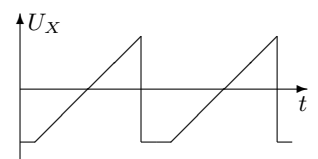
Dabei sind sowohl die horizontale wie auch die vertikale Auslenkung des Leuchtfleckes aus der Schirm-Mitte proportional zu der jeweiligen Ablenk-Spannung. Da unterschiedliche Polungen einer Spannung unterschiedlichen Vorzeichen des Spannungswertes entsprechen, kann man mit einem geeigneten Koordinatensystem auf dem Schirm sofort die angelegten Spannungen ablesen: Zum Beispiel könnte mit $U_X = -4\text{ V}$ der Leuchtfleck an den linken Schirmrand abgelenkt werden und mit $U_X = +4\text{ V}$ an den rechten Rand.

zu ergänzen:

- Beschriftung
- Spannungsquellen
- Verkabelung
- Elektronenstrahl
- Leuchtfleck
- Achsen- beschriftungen U_X, U_Y

Aus jeweils einer solchen Röhre bestehen **Oszilloskope**. Oszilloskope dienen dazu, die zeitliche Entwicklung elektrischer Spannungen oder Zusammenhänge zwischen Spannungen darzustellen. Indirekt lassen sich damit auch alle Größen darstellen, die in Form einer Spannung gemessen werden können. (Dies sind im Zeitalter elektronischer Messgeräte praktisch alle Größen!) Im Groben unterscheidet man zwei Betriebsarten:

Y-t-Modus Im Oszilloskop selbst wird eine sogenannte Sägezahnspannung erzeugt und an das Plattenpaar X für die horizontale Ablenkung gelegt, so dass der Elektronenstrahl immer wieder neu mit konstanter Geschwindigkeit von links nach rechts über den Schirm geführt wird. So hat der Leuchtfleck jederzeit eine X-Koordinate, die der aktuellen Zeit t entspricht. Die X-Achse wird daher zu einer t -Achse.



Der Benutzer legt von außen eine Spannung an das Plattenpaar Y für vertikale Ablenkung an. So hat der Leuchtfleck jederzeit eine Y-Koordinate gemäß der momentan von außen angelegten Spannung. Folglich durchläuft der Leuchtfleck den $U_Y(t)$ -Graphen.

Interessant wird es nun, wenn $U_Y(t)$ hochfrequent periodisch ist: „Periodisch“ heißt, dass $U_Y(t)$ nur aus der endlosen Wiederholung eines gleichbleibenden Verlaufs-Musters besteht. Startet jeder Sägezahn von U_X bei der gleichen Stelle innerhalb dieses Verlaufs-Musters, so durchläuft der Leuchtpunkt bei jedem Überqueren des Schirmes dieselbe Bahn. Das Abwarten des richtigen Startzeitpunktes für einen Sägezahn nennt man „Triggern“. Man kann das Oszilloskop z.B. auf die

Zeitpunkte triggern lassen, an denen U_Y einen Vorzeichenwechsel von Minus nach Plus macht. Oft macht man auf diese Art den Verlauf einer Schwingung (=Oszillation) sichtbar (griech. *skopein*=sehen).

„Hochfrequent“ heißt, dass das Muster nur von kurzer Dauer ist und daher viele Wiederholungen pro Zeiteinheit stattfinden. Um das Muster vernünftig aufzuzeichnen, muss der Leuchtfleck dann entsprechend schnell von links nach rechts geführt werden. Bewegt sich der Leuchtfleck aber recht schnell, nehmen wir nicht mehr den bewegten einzelnen Leuchtfleck wahr, sondern eine durchgehende Linie, quasi die Spur des Leuchtfleckes. Es ergibt sich dann ein stehendes Bild des $U_Y(t)$ -Graphen.

X-Y-Modus Der Benutzer legt sowohl die Spannung $U_X(t)$ für die horizontale Ablenkung als auch $U_Y(t)$ für die vertikale von außen an. Damit stellt man die Beziehung zwischen diesen beiden Spannungen dar: Jedes vorkommende Wertepaar ($U_X|U_Y$) wird durch den Leuchtfleck angezeigt, bei ausreichend rascher Wiederholung sogar in einem stehenden Bild.

23.5 Elektrisches Feld

23.5.1 Definition

Zum Abschluss dieses Kapitels über elektrische Ladung werfen wir noch einen Blick darauf, wie Physiker die Kräfte auf elektrische Ladungen für eine weitergehende Behandlung des Themas zweckmäßig beschreiben können. Bei weitergehender Behandlung stellt sich z.B. heraus, dass elektrisch geladene Körper nicht nur in der Nähe anderer Ladungen Kräfte erfahren können. Bei-

spiel: Die Elektronen in der Antenne Deines Radios erfahren Kräfte, weil die Antenne von Radiowellen getroffen wird. Es erweist sich schon von daher als günstig, die Kräfte auf geladene Körper losgelöst von der Existenz und Nähe anderer Ladungen allgemeiner zu beschreiben. Und genau dazu dient der Begriff vom elektrischen Feld:

Würde ein elektrisch geladener Körper an einer Stelle des Raumes aufgrund seiner elektrischen Ladung eine Kraft erfahren, so sagt man, es herrsche an dieser Stelle ein **elektrisches Feld**.

Da in der Umgebung einer Ladung Q_1 ein mit Q_2 geladener Körper eine Kraft erfährt, kann man sagen: Q_1 verursacht in seiner Umgebung ein elektrisches Feld, das auf Q_2 eine Kraft ausübt. Natürlich gilt das gleiche auch mit vertauschten Indices. So vermitteln elektrische Felder die Wechselwirkung zwischen zwei geladenen Körpern.

Ein elektrisches Feld hat an jeder Stelle eine bestimmte **Feldstärke**. Diese gibt an, wie viel Kraft pro dort befindlicher Ladung ausgeübt wird. Man sagt also z.B.: „Die elektrische Feldstärke in 3 cm Entfernung von Q_1 beträgt 8,15 Newton pro Coulomb.“ Die elektrische Feldstärke ist damit vergleichbar mit dem Ortsfaktor, der z.B. für das Schwerfeld an der Erdoberfläche angibt, dass 9,81 Newton Schwerkraft pro Kilogramm Masse ausgeübt werden.

Da zur vollständigen Angabe einer Kraft auch die Richtungsangabe gehört, hat das elektrische Feld an jeder Stelle auch eine Richtung, und zwar die Richtung der Kraft auf eine positive Ladung.

Zur Untersuchung eines elektrischen Feldes könnte man also einfach einen positiv geladenen Körper teilweise an verschiedene Stellen des Feldes halten und jeweils Betrag und Richtung der elektrischen Kraft auf diesen Körper messen. Die Ladung des Test-Körpers dürfte allerdings nicht allzu groß sein, damit sie nicht Ladungen in ihrer Umgebung durch Influenz nennenswert verschieben kann. Dadurch würde ja das zu messende Feld verändert. Eine solche, nicht allzu große positive Ladung zu Messzwecken nennt man eine „Probekladung“.

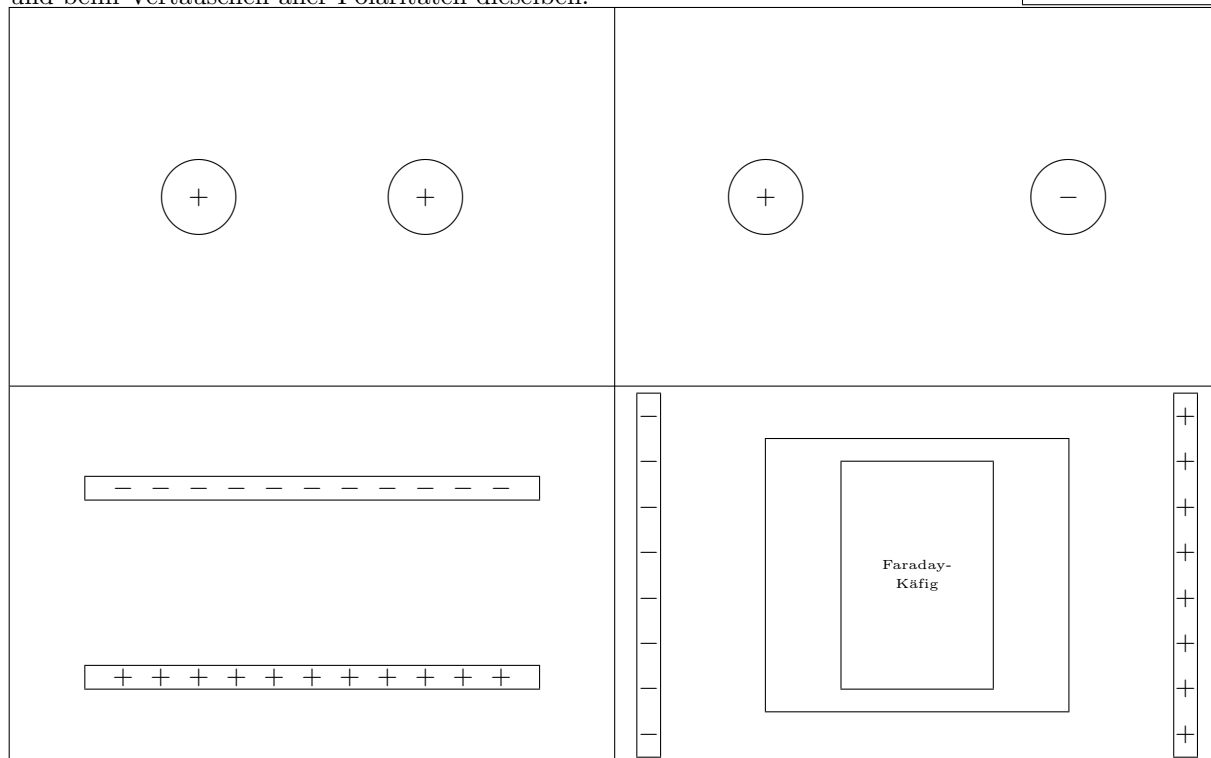
23.5.2 Darstellung

V 23.5 Feldlinienbilder

Aufbau: In einer flachen Schale befindet sich ein Öl mit Grieskörnern. In das Öl ragen verschieden geformte Metallteile, die zwecks elektrischer Aufladung an ein Hochspannungsnetzgerät angeschlossen werden können.

Durchführung: Die Grieskörner werden aufgerührt und die Metallteile elektrisch aufgeladen.

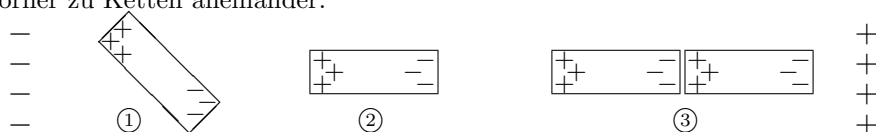
Beobachtung: Es bilden sich die am Rand und weiter unten gezeichneten Muster, und beim Vertauschen aller Polaritäten dieselben.



Erklärung: ① Das elektrische Feld influenziert in jedem Grieskorn positive und negative Aufladungen an entgegengesetzten Enden des Korns.

② Längliche Körner richten sich dementsprechend im Feld aus.

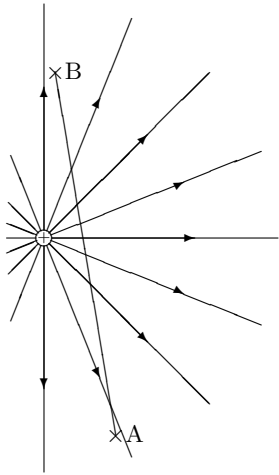
③ Wegen der Anziehung zwischen den Enden benachbarter Körner reihen sich nun Körner zu Ketten aneinander.



Die Grieskörner deuten ein Muster aus Linien an, den Feldlinien. Eine Feldlinie ist eine gedachte orientierte Linie, die in jedem ihrer Punkte dieselbe Richtung hat wie das Feld. Eine positive Ladung erfährt demnach eine Kraft in Richtung der Feldlinie, eine negative Ladung entgegengesetzt. Feldlinien zeigen uns die Struktur des Feldes. Um mit Feldlinien auch die Stärke des Feldes darstellen zu können, vereinbart man, in Gebieten höherer Feldstärke dichter liegende Feldlinien zu zeichnen. Darüber hinaus gilt noch:

- Feldlinien kreuzen sich niemals. Sonst gäbe es keine eindeutige Krafrichtung im Schnittpunkt.
- Feldlinien können sich noch nicht einmal berühren. Sonst hätte das Feld im Berührungspunkt entsprechend der unendlich hohen Liniendichte eine unendliche Feldstärke.
- Elektrische Feldlinien können nur an positiven und negativen Ladungen beginnen bzw. enden. Sie können allerdings auch in sich geschlossene Linien sein, die ein sich änderndes Magnetfeld umschließen.

- Aus Leiteroberflächen treten Feldlinien elektrostatischer (d.h. zeitlich unveränderlicher elektrischer) Felder immer senkrecht aus. Bei schrägem Verlauf hätten die Feldkräfte auch Komponenten parallel zur Oberfläche, die die beweglichen Ladungsträger unter der Oberfläche verschieben würden. Dies wäre keine statische Situation.
- Innerhalb von Leitern gibt es gar keine elektrostatischen Feldlinien. Ein Feld darin baut sich selbst sofort ab, indem es Ladungsträger entgegen seiner eigenen Polung verschiebt. Ein Feld von außen kann also nicht eindringen, es beeinflusst am Außenrand des Leiters genau die Ladungen, die das Feld im Inneren ausgleichen. (→ Faraday-Käfig)



23.5.3 Spannung

Bringt man innerhalb eines elektrischen Feldes eine Probeladung von einem Punkt A zu einem Punkt B, so wird man i.A. Teile des Weges in Richtung von Feldlinien, entgegengesetzt oder quer dazu zurücklegen. Dementsprechend wird das Feld auf die Ladungen unterschiedliche Kräfte ausüben: darunter Kräfte mit Komponenten in Richtung des Weges und Kräfte mit entgegengesetzten Komponenten.

Entlang des Weges wird man also manchmal gegen die Feldkräfte Arbeit verrichten müssen, andere Male wird das Feld an der Ladung Arbeit verrichten. Wenn wir wollen, verrichtet es dann mittelbar über die Ladung auch Arbeit an uns. In dem einen Fall stecken wir Energie in das System hinein, im anderen Fall gewinnen wir Energie aus dem System. Insgesamt kann es uns also eine gewisse Energie kosten, die Ladung von A nach B zu bringen, oder es kann Energie frei werden. Rechnerisch ist dies nur eine Frage des Vorzeichens der verrichteten Arbeit.

Natürlich ist auf demselben Weg die umgesetzte Energie größer, wenn die Ladung größer ist: Größere Ladungen erfahren ja größere Kräfte. Will man nun den Energieumsatz entlang des Weges unabhängig von einer bestimmten Probeladung beschreiben, kommt dafür nur die umgesetzte Energie *pro Ladung* in Frage. Die Größe *Energie pro Ladung* bekommt einen eigenen Namen:

zu ergänzen:

- Formeltermine

Erfordert es die Arbeit W , eine Probeladung Q vom Punkt A zum Punkt B zu bringen, dann ist die Spannung U über dem Weg von A nach B der Quotient aus Arbeit und Ladung.

$$U :=$$

Passend definierte Einheit: 1 Volt = 1 V :=

Natürlich gilt diese Definition auch für Wege zwischen den Polen einer Stromquelle. Bisher waren dies die einzigen Stellen, zwischen denen wir Spannungen betrachtet haben, ohne uns eines elektrischen Feldes bewusst gewesen zu sein. Da aber in Kabeln zum Antrieb eines Stromes elektrische Kräfte auf die Elektronen wirken, gibt es dort definitionsgemäß elektrische Felder. Trotz obiger Ausführungen zum Faraday-Käfig jetzt doch dauerhaft elektrische Felder in einem Leiter? Ja: Der im Kabel fließende Strom ist der permanent von der Stromquelle vereitelte Versuch des Feldes, sich selbst (und die zugehörige Spannung) durch einen Strom abzubauen. Gegen die Tendenz zum Ladungsausgleich hält die Quelle durch dauerndes Ersetzen abgeflossener Ladungen das Feld zwischen ihren Polen „aufgespannt“.

Für unsere weiteren Betrachtungen können wir noch manches außer Acht lassen: Solange das elektrische Feld keine in sich geschlossenen Feldlinien enthält, spielt es auch keine Rolle, *welchen* Weg von A nach B man betrachtet: Über jedem Weg fällt dieselbe Spannung ab! Statt von einer Spannung *über einem Weg zwischen zwei Punkten* können wir von einer wegunabhängigen Spannung *zwischen zwei Punkten* sprechen.

Um die Sprechweisen noch weiter zu vereinfachen, kann man dann mit einer weiteren Vereinbarung sogar über *einzelne* Punkte Aussagen machen: Legt man zunächst einen Punkt O als Bezugspunkt fest, so kann man dann die Spannung zwischen dem Bezugspunkt O und einem Punkt P als **Potenzial** des Punktes P bezeichnen. Als Bezugspunkt für elektrische Schaltkreise eignen sich z.B. geerdete Stellen hervorragend.

Elektrische Größen

24.1 QUIRP – Steckbriefe

24.1.1 Elektrische Ladung

- Elektrische Ladung ist eine Eigenschaft, die manche Teilchen von Natur aus in einem unveränderlichen Ausmaß haben. Es gibt sie in zwei Ausprägungen, die sich durch ein positives bzw. negatives Vorzeichen beschreiben lassen.
- Beim Zusammenfügen von Teilchen ergibt sich die Ladung des Ganzen als Summe der einzelnen Ladungen.
- Ladung macht sich durch Kräfte zwischen geladenen Körpern bemerkbar. Gleichartig geladene Körper stoßen sich gegenseitig ab, verschiedenartig geladene ziehen sich an.
- Die für uns wichtigsten geladenen Teilchen sind die positiv geladenen Protonen und die ebenso stark negativ geladenen Elektronen. Protonen bilden zusammen mit den elektrisch neutralen Neutronen die Atomkerne, die von drumherum schwirrenden Elektronen eingehüllt werden.

24.1.2 Stromstärke

- Elektrischer Strom ist bewegte elektrische Ladung.
- Allermeist haben wir es dabei zu tun mit Elektronen, die sich innerhalb von Metallen relativ leicht bewegen können.
- Die Stärke des Stroms an einer Stelle ist die dort durchfließende Ladungsmenge pro Zeit: $I = \frac{Q}{t}$
- Ein Amperemeter wird *in die Leitung* eingebaut, in der es die Stromstärke messen soll. Ein ideales Amperemeter verändert (insbesondere: hemmt) den zu messenden Strom nicht.
- Vorsicht: Die Bezeichnung „Stärke“ kann irreführen: Stromstärke beschreibt keine Kräfte, insbesondere auch nicht die Stärke des Antriebes der bewegten Ladungen. Man denke bei der Stromstärke I eher an *Intensität*.
- Da die Anzahl frei beweglicher Elektronen in einer metallischen Leitung praktisch unveränderlich ist, sind bei größerer Stromstärke nicht *mehr* Elektronen unterwegs, sondern es driften alle Elektronen im Durchschnitt *schneller* durch den Leiter.

Q

Einheit: Coulomb

1 C = Ladung
bestimmter Anzahl
Protonen

$$I = \frac{Q}{t}$$

Einheit: Ampere

$$1 \text{ A} = \frac{1 \text{ C}}{1 \text{ s}}$$

$$U = \frac{W}{Q}$$

Einheit: Volt

$$1 \text{ V} = \frac{1 \text{ J}}{1 \text{ C}}$$

24.1.3 Spannung

- Im Allgemeinen gibt es immer Effekte, die die Bewegung von Ladungsträgern hemmen. Es ist daher in der Regel Arbeit erforderlich, um Ladungsträger in Bewegung zu halten. Diese Arbeit verrichten Batterien, Netzgeräte und andere Stromquellen.
- Um die Stärke des Antriebes zu beschreiben, gibt man an, wie viel Arbeit pro angetriebener Ladung verrichtet wird. Die Spannung ist demnach definiert über die Gleichung $U = \frac{W}{Q}$
- Eine Spannung herrscht immer zwischen zwei Punkten. Sie beschreibt, wie viel Energie umgesetzt wird, wenn Ladung von dem einen zum anderen Punkt bewegt wird.
- Ein Voltmeter wird *zwischen die beiden Punkte* der Schaltung angeschlossen, zwischen denen es die Spannung messen soll. Ein ideales Voltmeter verändert die zu messende Spannung nicht. Insbesondere stellt es keinen neuen Stromweg zwischen den Anschlusspunkten dar, wird also idealerweise nicht von Strom durchflossen.
- Vorsicht: Spannung als *Stärke des Antriebes* eines Stromes ist nicht die *Stärke des Stromes*. Da Spannung nicht den Strom, sondern seinen Antrieb beschreibt, gibt es auch das Wort „Stromspannung“ nicht.

$$P = U \cdot I$$

Einheit: Watt

$$1 \text{ W} = \frac{1 \text{ J}}{1 \text{ s}}$$

24.1.4 Elektrische Leistung

- Leistung allgemein ist umgesetzte Energie pro Zeit: $P = \frac{W}{t}$
- Fließt ein Strom der Stärke I von einem Punkt zu einem anderen, wobei zwischen den Punkten die Spannung U herrscht, dann bewegt sich in der Zeit t die Ladungsmenge $Q = I \cdot t$, wobei die Energie $W = U \cdot Q = U \cdot I \cdot t$ umgesetzt wird. Die Leistung $P = \frac{W}{t} = \frac{UIt}{t}$ beträgt also $P = U \cdot I$

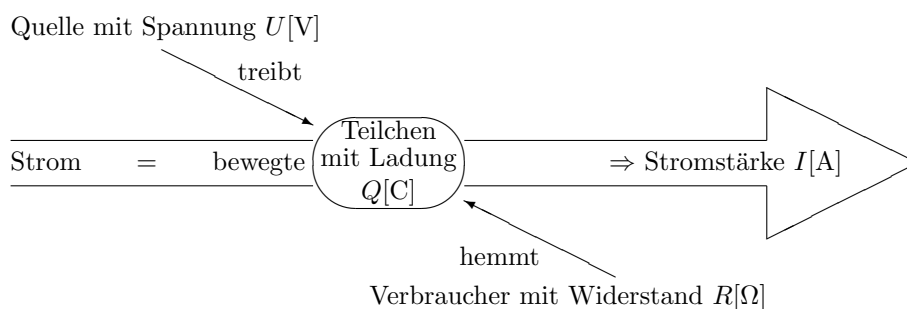
$$R = \frac{U}{I}$$

Einheit: Ohm

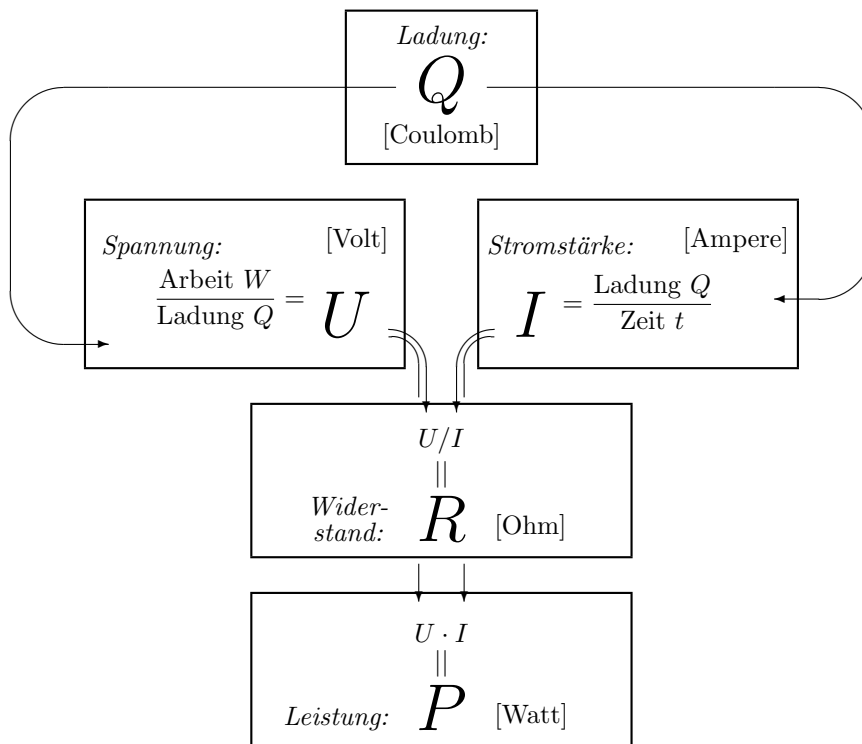
$$1 \Omega = \frac{1 \text{ V}}{1 \text{ A}}$$

24.1.5 Ohmscher Widerstand

- (Fast) Jeder Leiter hemmt den Stromfluss dadurch, dass die bewegten Ladungsträger mit den Atomen des Leiters unter Energieverlust zusammenstoßen.
- Stärkere Hemmung bedeutet mehr erforderliche Spannung für selbe Stromstärke. Daher definiert man als (ohmschen) Widerstand eines elektrischen Bauteils: $R = \frac{U}{I}$
- Die Energieverluste der strömenden Ladungsträger führen zu heftigeren Bewegungen der Atome des Leiters, also zu einer Erwärmung. Aus dem Energieerhaltungssatz folgt für diese „Joule’sche“ Wärme: $W = U \cdot I \cdot t$



24.1.6 Kurz-Übersicht



24.2 Spannung vs. Stromstärke

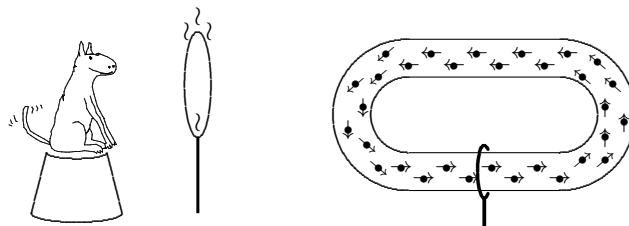
Wie kann man sich Stromstärke vorstellen?

Wenn Du durch einen Gartenschlauch pustest, erzeugst Du darin eine Luftströmung. Mikroskopisch betrachtet heißt dies, dass sich Teilchen der Luft – dies sind v.a. Stickstoff- und Sauerstoffmoleküle – längs durch den Schlauch bewegen.

Beim elektrischen Strom durch ein Kabel bewegen sich Elektronen durch das Kabel. Kabel sind zwar keine hohlen Schläuche, aber zwischen ihren Atomkernen ist für die super-super-winzigsten Elektronen reichlich Platz zum Durchströmen.

Wir markieren nun eine Stelle des Schlauches, indem wir dort einen Draht zu einem Ring um den Schlauch wickeln. So strömen die Luftteilchen jetzt auch durch den Drahtring. (ähnlich wie im Zirkus Löwen durch brennende Ringe springen) Der Luftstrom wird „stärker“ genannt, wenn einfach mehr Teilchen pro Zeit durch den Ring fliegen. (Egal, wie schnell sie fliegen, egal, wie stark man pusten muss, damit sie so fliegen: Sie werden nur gezählt.)

Gedanklich könnte man in einem elektrischen Stromkreis ebenso an einer markierten Stelle zählen, wie viele Elektronen dort pro Sekunde hindurchtreten. Diese Anzahl pro Zeit entspricht der elektrischen Stromstärke I . Allerdings kann man die strömenden Elektronen wesentlich schlechter zählen als Löwen, so dass I ein bisschen anders definiert wird.



Schon die im Vergleich mit Elektronen viel größeren strömenden Luftmoleküle sind praktisch kaum zu zählen. Zur Angabe der Stärke eines Luftstromes zählt man deshalb nicht die Teilchen. Statt dessen wiegt man die durch den Ring strömende Luftmenge, oder man bestimmt ihr Volumen. Die Luft-Stromstärke wird daher eher in Gramm pro Sekunde oder in Liter pro Sekunde angegeben.

Bei den strömenden Elektronen interessiert natürlich weder ihre Masse noch ihr Volumen: Für elektrische Phänomene ist ja die elektrische Ladung entscheidend. Die elektrische Stromstärke I an einer Leiterstelle ist demnach definiert als die dort durchströmende Ladungsmenge Q pro Zeitspanne t . (Diese Definition passt auch auf Fälle, bei denen andere Teilchen als Elektronen den Strom bilden.)

Wie kann man sich Spannung vorstellen?

Pustet jemand von der anderen Seite des Schlauches her genauso stark wie Du dagegen, strömt natürlich nichts. Mit physikalischen Größen ausgedrückt bedeutet „genauso stark pusten“ so viel wie „gleich großen Druck erzeugen“. Für den Antrieb des Luftstromes muss ein Druckunterschied zwischen den Enden des Schlauches herrschen.

Auch ein Elektron bewegt sich nicht durch ein Kabel, wenn es von beiden Seiten her gleich starken Abstoßungskräften der Elektronen in seiner Nachbarschaft ausgesetzt ist. (Von zufällig gerichteten thermisch bedingten Bewegungen sehen wir hier ab.) Die Rolle der zum Antrieb nötigen „Druckdifferenz“ spielt die Spannung.

Erfährt ein Teilchen von links und rechts unterschiedlich große Druckkräfte F_L bzw. F_R , hebt die kleinere Kraft die größere teilweise auf und es bleibt eine Differenz $F = |F_L - F_R|$ in Richtung der größeren Kraft wirksam übrig. Werden Teilchen von dieser Kraft entlang einer Strecke s geführt, so wird dabei pro Teilchen die Arbeit $W = F \cdot s$ verrichtet. Die Stärke des Antriebes lässt sich daher statt mit Drücken oder Kräften auch mit einem Energieumsatz pro angetriebenem Etwas beschreiben.

Man kann den Antrieb von Elektronen zwar nicht immer mit einem „Elektronendruck“ im „Elektronengas“ erklären, aber die Überlegung zum Energieumsatz einer antreibenden Kraft entlang einer Strecke bleibt richtig. Was auch immer für den Antrieb verantwortlich ist: Auf der Strecke zwischen zwei Stellen eines Stromweges wird pro Elektron – und damit auch pro bestimmter Ladung Q – eine bestimmte Energie W umgesetzt. Der Antrieb wird daher durch die Spannung $U = W/Q$ beschrieben.

Wie hängen Stromstärke und Spannung zusammen?

Um in einem geschlossenen Schlauchkreislauf eine höhere Stromstärke zu erhalten, muss man die im Kreislauf befindliche Luft schneller strömen lassen. Dazu muss eine Pumpe sie stärker antreiben. Schneller strömende Luftteilchen erfahren nämlich stärkere Reibungskräfte, die zu kompensieren sind. (Nebenbei zu dieser Reibung: Luftteilchen an der Schlauchwand haften an der Wand. Schnellere Luftschichten weiter innen im Schlauch reiben daher an langsameren Schichten weiter außen.) Folglich wird entlang des Schlauches mehr Energie pro strömender Luftmenge umgesetzt.

Um in einem geschlossenen Stromkreis eine höhere Stromstärke zu erhalten, muss man die im Stromkreis befindlichen Elektronen schneller strömen lassen. Dazu muss man sie stärker antreiben, denn auch Elektronen unterliegen bei höherer Geschwindigkeit stärkeren hemmenden Einflüssen. (Man kann sich ja vorstellen, dass schnellere Elektronen bei Zusammenstößen mit anderen Teilchen mehr Energie verlieren können als langsamere.) Folglich wird entlang des Leiters mehr Energie pro strömender Ladung umgesetzt, es ist mit anderen Worten eine höhere Spannung erforderlich.

$$U = R \cdot I$$

Kurz: Höhere Stromstärke bedeutet höhere Strömungsgeschwindigkeit, die mehr Energieumsatz pro Ladung erfordert, was höhere Spannung bedeutet. Es gilt (zumindest in Metallen bei konstanter Temperatur): $I \sim v \sim \frac{W}{Q} = U$, wobei der Proportionalitätsfaktor zwischen I und U als Widerstand R bezeichnet wird.

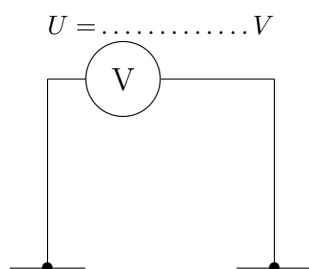
24.3 Bestätigende Experimente

24.3.1 Kombination von Spannungsquellen

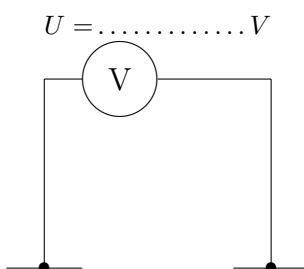
V 24.1 Kombination von Spannungsquellen

Aufbau: /Durchführung: Wir schalten zwei Spannungsquellen auf verschiedene Weisen zusammen und messen jeweils die Spannung der Kombination.

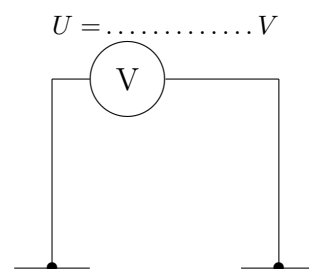
seriell, gleichsinnig:



seriell, gegensinnig:



parallel, gleichsinnig:



(Die fehlende Kombination „parallel, gegensinnig“ wäre ein für die Spannungsquellen ungesunder Kurzschluss der Spannungsquellen.)

zu ergänzen:

- Spannungsquellen
- Messwerte

- Bei gleichsinniger Serienschaltung verrichten an jedem Ladungsträger nacheinander mehrere Quellen eine gewisse Antriebsarbeit mit selber Richtung. Von daher ist klar, dass sich die Spannungen (=Arbeit pro Ladung) addieren. So lassen sich aus Quellen kleinerer Spannung Quellen mit größerer Spannung zusammenstellen. Man baut z.B. 4,5 V-Flachbatterien aus drei 1,5 V-Zellen auf.
- Gegensinnige Serienschaltung scheint zunächst kaum vernünftig anwendbar zu sein. Mit ihr ließe sich jedoch direkt prüfen, welche von zwei Spannungsquellen die stärkere ist. Die stärkere Quelle bestimmt die Richtung, in die diese Kombination einen Strom treiben würde. Sind die Quellen gleich stark, ergibt sich insgesamt kein Antrieb, keine Spannung, kein Strom.
- Gleichsinnige Parallelschaltung gleich starker Quellen scheint ebenfalls keinen Vorteil zu bringen, weil ihre Spannung ja dieselbe ist wie bei den einzelnen Quellen. Jeder Ladungsträger wird ja nur von einer einzigen Batterie angetrieben. Doch liegt ein Vorteil darin, dass jede Quelle nur einen Teil des Gesamtstromes antreiben muss, so dass die Quellen geschont werden. Betreibt man z.B. ein 1,5 V-Lämpchen statt mit einer einzigen Monozelle mit der Parallelschaltung aus 3 Monozellen, so brennt das Lämpchen zwar nicht heller, aber ohne Batteriewechsel dreimal so lange.

24.3.2 Elektrische Leistung

Die Gleichung $P = U \cdot I$ ist noch experimentell zu bestätigen. Stimmt die von einem Strom der Stärke I bei einer Spannung U pro Zeit t verrichtete Arbeit W tatsächlich mit $U \cdot I$ überein?

Zum Vergleich von W/t mit $U \cdot I$ wäre möglich, einen Elektromotor eine mechanische Arbeit verrichten zu lassen, ihn z.B. etwas hochheben zu lassen. Wegen dabei auftretenden Reibungseffekten und ähnlichem ist aber die vom Strom insgesamt gelieferte Energie dabei nicht so leicht zu messen.

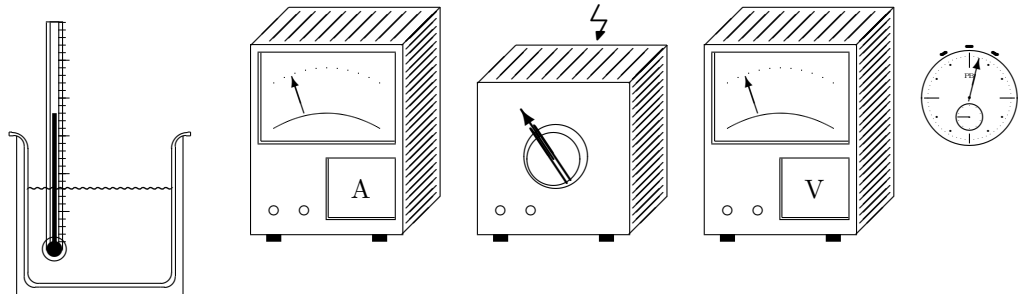
Günstiger ist es, zu bestimmen, welche Leistung ein Strom beim Erwärmen eines Stoffes bringt. Die zum Erwärmen der Stoffportion nötige *Energie* können wir aus der Masse des erwärmten Stoffes, seiner Wärmekapazität und der Temperaturzunahme berechnen. Um die *Erwärm-Leistung* zu berechnen, muss die umgesetzte Energie nur noch durch die benötigte Zeit dividiert zu werden.

V 24.2 Tauchsieder

zu ergänzen:

- Tauchsieder
- Verkabelung
- (Deckel)

Aufbau: Wassermenge im Kalorimeter: g



Messwerte:

Auswertung:

Elektrische Schaltungen

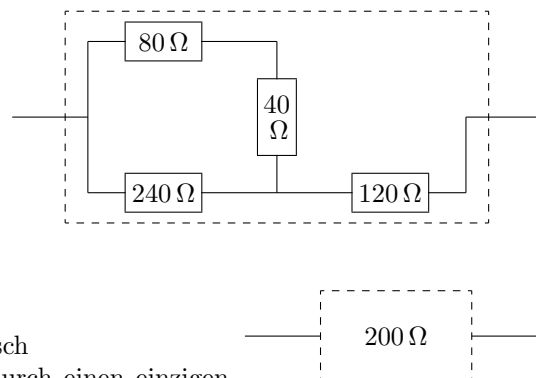
25.1 Serien- und Parallelschaltung

25.1.1 Ersatzwiderstand

Im Folgenden untersuchen wir die Zusammenhänge in komplizierteren Schaltungen. Man kann ja elektrische Bauteile zu beliebig komplizierten Leiternetzen zusammenschalten. Man kann z.B. einige ohmsche Widerstände wie hier rechts gezeichnet verbinden. Legt man an diese Schaltung z.B. 180 V an, so fließen durch die Zuleitungen 0,9 A. Bei doppelter Spannung fließt ein doppelt starker Strom. Die gesamte, gestrichelt eingerahmte Schaltung verhält sich also wie ein einziger ohmscher Widerstand, und zwar einer mit $R = \frac{U}{I} = \frac{180 \text{ V}}{0,9 \text{ A}} = 200 \Omega$.

Man würde also außerhalb des gestrichelten Rahmens elektrisch keinen Unterschied bemerken, wenn man die Schaltung darin durch einen einzigen 200 Ω -Widerstand ersetzen würde. Man sagt daher, die gezeichnete Schaltung habe einen **Ersatzwiderstand** (oder auch „Gesamtwiderstand“) von 200 Ω .

Innerhalb des Rahmens könnte zum Beispiel interessieren, wie viel Strom durch den 40 Ω -Widerstand in der Mitte fließt, wenn man von außen 180 V anlegt. Um darauf die Antwort 0,6 A berechnen zu können, benötigt man alle Regeln, die wir nun kennenlernen werden.



25.1.2 Serienschaltung

V 25.1 Stromstärken bei Serienschaltung

Aufbau:

Beobachtung: /Ergebnis:

V 25.2 Spannungen bei Serienschaltung

Aufbau:

Beobachtung: /Ergebnis:

Berechnung des Ersatzwiderstandes:

Serienschaltung:

$$R_S =$$

$$U_{\text{ges}} =$$

$$=$$

$$=$$

$$=$$

$$=$$

$$| U_x = R_x \cdot I_x$$

$$| I_x = I_{\text{ges}}$$

$$| I_{\text{ges}} \text{ auskla.}$$

$$| : I_{\text{ges}}$$

25.1.3 Parallelschaltung

V 25.3 Stromstärken bei Parallelschaltung

Aufbau:

Beobachtung: /Ergebnis:

V 25.4 Spannungen bei Parallelschaltung

Aufbau:

Beobachtung: /Ergebnis:

Berechnung des Ersatzwiderstandes:

$I_{\text{ges}} =$	$I_x = U_x / R_x$
$=$	$U_x = U_{\text{ges}}$
$=$	U_{ges} auskla.
$=$	$I : U_{\text{ges}}$
$=$	

Parallelschaltung:

=

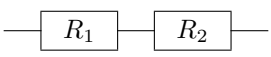
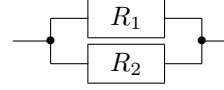
25.1.4 Zusammenstellung der Regeln

Die Namen *Knoten*- und *Maschen*-regel in folgender Übersicht stammen aus dem Vergleich einer Schaltung mit einem Fischer-Netz: Verzweigungspunkte, an denen mehr als zwei Leitungen miteinander verbunden sind, entsprechen den Knoten eines Netzes. Zwei verschiedene Leiterwege zwischen zwei Punkten bilden eine Masche des Netzes. Unterscheidet man die zu einem Knoten hin- und die wegfließenden Stromstärken durch die Vorzeichen + und –, so lässt sich als Knotenregel formulieren:

An jedem Knoten ist die Summe aller Stromstärken Null.

Auch bei Spannungen sind Vorzeichen sinnvoll, z.B. so: Fließt Strom von A nach B, so ist die Spannung von A nach B positiv und die von B nach A negativ. Sieht man nun von zwei Wegen zwischen A und B den einen als Hinweg, den anderen als Rückweg an, so bilden beide Wege einen Rundweg entlang einer Masche. Die Maschenregel lautet in dieser Sichtweise:

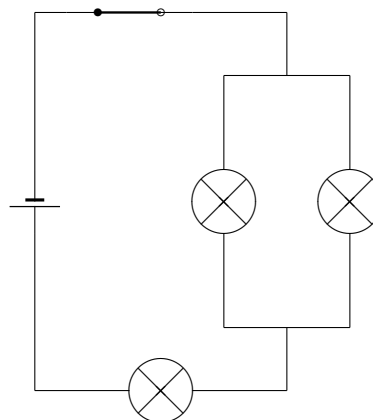
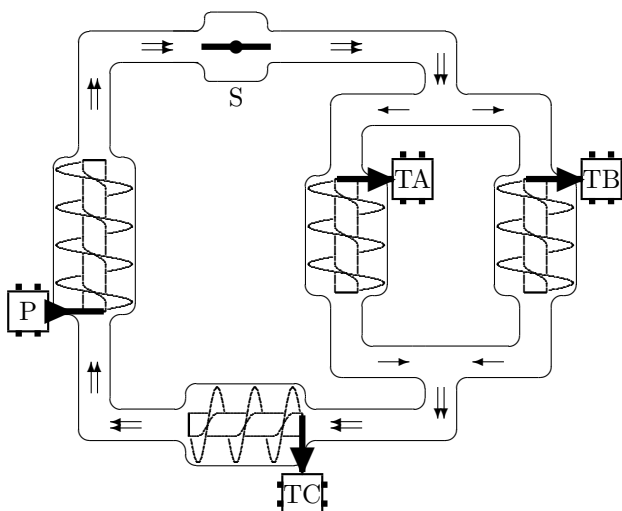
Entlang jeder Masche ist die Summe aller Spannungen Null.

Serienschaltung 	Parallel-schaltung 
$I_{ges} = I_1 = I_2$ In einer unverzweigten Leitung ist die Stromstärke überall gleich. Elektrische Ladung bleibt erhalten und sammelt sich nirgendwo immer weiter an.	$I_{ges} = I_1 + I_2$ Knotenregel: An jedem Knoten eines Leiternetzes ist die Summe der hinfließenden Stromstärken gleich der Summe der abfließenden.
$U_{ges} = U_1 + U_2$ Entlang eines Stromweges addieren sich die Spannungen.	$U_{ges} = U_1 = U_2$ Maschenregel: Über allen Wegen zwischen zwei Punkten fällt dieselbe Spannung ab.
$R_{ges} = R_1 + R_2$ Es addieren sich die Widerstände.	$\frac{1}{R_{ges}} = \frac{1}{R_1} + \frac{1}{R_2}$ Es addieren sich die Leitwerte.
Die Teil-Spannungen verhalten sich wie die Teil-Widerstände.	Die Teil-Ströme verhalten sich umgekehrt wie die Teil-Widerstände.

Bei einer **Serienschaltung** addieren sich die Widerstände, was ja sofort einleuchtet, weil ein Elektron beim Durchfließen einer Serienschaltung alle Widerstände nacheinander überwinden muss.

Bei **Parallelschaltung** addieren sich *Kehrwerte* von Widerständen, was sicher nicht ganz so schnell einleuchtet. Ein Widerstand gibt als erforderliche Spannung pro gewünschter Stromstärke an, wie stark ein Bauteil den Stromfluss hemmt, denn bei *größerer* Hemmung sind pro Ampere *mehr* Volt nötig. Der Kehrwert gibt als erreichte Stromstärke pro angelegter Spannung eher an, wie gut ein Bauteil den Strom leitet, denn bei *größerer* Leitfähigkeit werden pro Volt *mehr* Ampere bewirkt. Deshalb heißen die Kehrwerte von Widerständen auch **Leitwerte**. Und dass zwei parallel nebeneinander verlaufende Leitungen so gut leiten wie die beide Einzelleitungen zusammen, leuchtet dann doch ein.

25.1.5 Vergleich: Stromkreis - Luftstrom



Links treibt eine Pumpe P einen Luftstrom durch ein Rohrleitungsnetz, wie es die Pfeilchen andeuten. Durch Drehen z.B. einer schraubenförmigen Turbine (ähnlich wie bei einem Fleischwolf) befördert sie Luft-Teilchen von ihrem unteren Ende an ihr oberes. (Du kannst Dir anstelle dieser Turbinenform auch einen Flugzeug-Propeller oder ein Windrädchen vorstellen.)

Die Geräte TA, TB, und TC sollen über Turbinen vom Luftstrom angetrieben werden. Da diese Turbinen den Luftstrom natürlich hemmen, entsteht vor den Turbinen TA und TB ein Luft-Stau mit erhöhtem Druck. Da die Pumpe an ihrem unteren Ende Luft einsaugt, entsteht außerdem hinter der Turbine TC ein Unterdruck. Färbe dem Druck entsprechend den oberen Rohrabschnitt zwischen P und TA/TB in kräftigerem Rot, das Stück unten links zwischen TC und P in schwächerem Rot! Das mittlere Rohrstück unten rechts zwischen TA, TB und TC bekommt natürlich ein mittelstarkes Rot. Nun solltest Du in diesem Rohrsystem drei Bereiche mit jeweils einem anderen Druck gut erkennen können.

An jeder Turbine herrscht auf der einen Seite ein höherer Druck als auf der anderen Seite. Diese Druck*differenz* ist es nun eigentlich, die die Luft durch die Turbine treibt. Würde an beiden Enden der Turbine gleicher Druck herrschen, gäbe es keinen Luftstrom, egal wie groß der Druck wäre! Je größer die Druckdifferenz, desto stärker der Luftstrom.

Kommen wir nun zu obigen Gleichungen, zunächst für die Parallelschaltung. Betrachte also die Parallelschaltung aus TA und TB!

$I_{ges} = I_1 + I_2$ Am oberen Verzweigungspunkt vor TA und TB teilt sich der Luftstrom in zwei Teilströme auf, die natürlich zusammen genauso viel Gramm Luft pro Sekunde vom Verzweigungspunkt wegführen, wie der ankommende Strom herbeiführt.

$U_1 = U_2$ Beide Turbinen liegen zwischen denselben Druckgebieten, sind also derselben Druckdifferenz ausgesetzt.

Um eine Serienschaltung zu sehen, betrachte zunächst die Kombination von TA und TB als ein einziges Bauteil, quasi als die „Doppel-Turbine“ TAB (wenn Du willst, gestrichelt umrahmt), betrachte dann die Serienschaltung aus dieser Doppelturbine und TC:

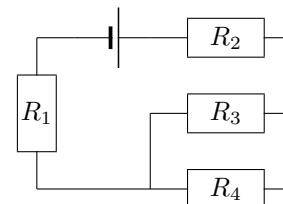
$I_1 = I_2$ Wenn aus der Doppelturbine pro Sekunde *mehr* Luft austräte als in TC einströmt, dann würde sich im Bereich dazwischen Luft anstauen. Der Druck stiege dort folglich an. Dadurch nähme aber die Druckdifferenz über TAB ab, die über TC zu. Die Stromstärke durch TAB nähme daher ebenfalls ab, die durch TC zu. Die beiden Stromstärken würden sich demnach einander angleichen. Eine analoge Überlegung zeigt auch, dass aus TAB nicht dauerhaft *weniger* Luft austreten kann als in TC einströmt. Auf Dauer können die beiden Stromstärken also nur gleich sein.

$U_{ges} = U_1 + U_2$ Herrscht z.B. vor der Doppelturbine ein Druck von 4 bar, zwischen TAB und TC 3 bar und hinter TC nur 1 bar, dann betragen die Druckdifferenzen über TAB 1 bar, über TC 2 bar und über der gesamten Turbinen-Anordnung 3 bar. Die Additivität ist kaum zu übersehen.

25.2 Komplexe Schaltungen

...lassen sich rechnerisch oft dadurch in den Griff kriegen, dass man sich immer wieder geeignete Ausschnitte der Schaltungen durch ihren Ersatzwiderstand ersetzt denkt und auf geeignete Ausschnitte die passenden Grundregeln anwendet. Beispiel:

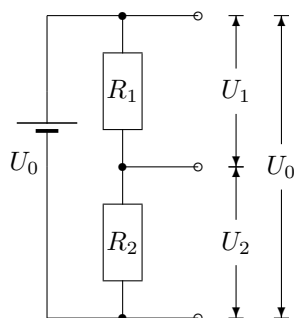
In der nebenstehenden Schaltung lässt sich zunächst die Parallelschaltung aus R_3 und R_4 durch ihren Gesamtwiderstand ersetzen, den wir (– hoffentlich selbsterklärend? –) R_{34} nennen wollen. Dieser bildet mit R_1 eine Serienschaltung mit Gesamtwiderstand R_{134} . Schließlich erhält man mit R_2 den Gesamtwiderstand $R_{ges} = R_{1..4}$. Damit lässt sich I_{ges} ausrechnen, das gleichzeitig I_1 , I_2 und I_{34} ist, womit sich U_1 , U_2 und U_{34} (per $U = R \cdot I$) errechnen lassen. Bekanntlich ist noch in der Parallelschaltung $U_{34} = U_3 = U_4$, woraus schließlich I_3 und I_4 gewonnen werden können.



25.3 Anwendungen

25.3.1 Spannungsteiler

Bei einer Serienschaltung teilt sich die Gesamtspannung auf die einzelnen Widerstände auf. Die Teilspannungen verhalten sich dabei wie die Widerstände.

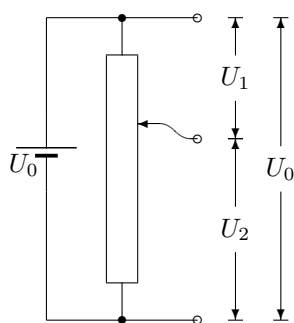


$$\begin{aligned} I_1 &= I_2 & | \quad I &= \frac{U}{R} \\ \frac{U_1}{R_1} &= \frac{U_2}{R_2} & | \cdot R_1 & : U_2 \\ \frac{U_1}{U_2} &= \frac{R_1}{R_2} & | \end{aligned}$$

Ebenso gilt natürlich auch $\frac{U_1}{U_0} = \frac{R_1}{R_1 + R_2}$.

Durch geeignete Wahl der Widerstände lässt sich also U_0 in beliebigem Verhältnis aufteilen, es lassen sich beliebige (echte) Bruchteile von U_0 abgreifen. Wir wollen als Beispiel $\frac{3}{5}$ von U_0 als U_1 abgreifen. Der Rest von U_0 – also U_2 – ist dann $\frac{2}{5}$ von U_0 . Wir müssen also U_0 im Verhältnis 3 : 2 teilen. Dazu müssen wir zwei Widerstände benutzen, die in diesem Verhältnis stehen, zum Beispiel $30\,\Omega$ als R_1 und $20\,\Omega$ als R_2 .

Einen sehr flexiblen Spannungsteiler („Potenziometer“) erhält man, wenn man statt zweier fester Widerstände die beiden Abschnitte eines Schiebewiderstandes verwendet. Durch Verschiebung des Mittelabgriffs lassen sich der Gesamtwiderstand und damit auch die Gesamtspannung recht fein abgestuft in verschiedenen Verhältnissen zu R_1 und R_2 bzw. U_1 und U_2 zerlegen.



?

Wozu braucht man Spannungsteiler?

Zunächst zum Naheliegenden: Spannungsteiler sind in der Praxis *nicht* gut als Stromquellen mit einstellbarer Spannung einsetzbar. Bei kleinem Widerstand des Spannungsteilers wird nämlich in ihm durch die hohe Stromstärke viel Energie unnütz verbraucht. Bei großem Widerstand ändert ein angeschlossener Verbraucher mit seinem relativ kleinen Widerstand die Spannungsaufteilung deutlich zu seinen Ungunsten: Beim Anschluss des Verbrauchers an eine der beiden Teilspannungen „bricht diese Spannung ein.“

Spannungsteiler werden verwandt, wenn man nur einstellbare Spannungen benötigt, aber (fast) keine Ströme fließen. Du fragst, wozu man sonst Spannungen braucht, wenn nicht um Ströme anzutreiben? Dazu drei Beispiele:

- Gelegentlich braucht man elektrische Felder, die man ohne Ströme mit einer Spannung aufrecht erhält.
- Manchmal will man eine Spannung messen, indem man ihr eine so große Gegenspannung entgegenwirkt, dass irgendwo gerade kein Strom mehr fließt.
- Viele elektronische Bauteile werden durch Spannungen gesteuert, ohne nennenswerte Ströme.

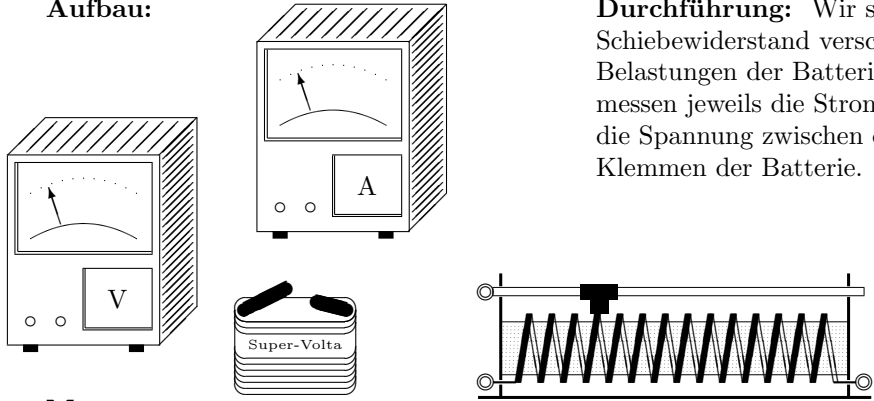
25.3.2 Innenwiderstände

Stromquellen

Nicht nur beim Spannungsteiler kann man die Erfahrung machen, dass die Spannung einer Stromquelle abnimmt („zusammenbricht“, „in die Knie geht“), sobald man der Quelle mehr Strom entnimmt.

V 25.5 Batteriespannung bei Belastung

Aufbau:



Durchführung: Wir stellen am Schiebewiderstand verschiedene Belastungen der Batterie ein und messen jeweils die Stromstärke und die Spannung zwischen den Klemmen der Batterie.

zu ergänzen:

- Verkabelung

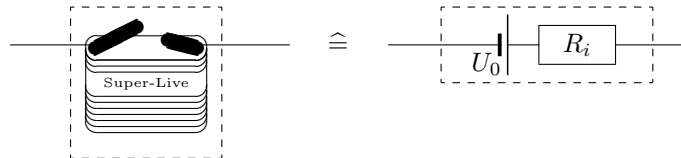
Messwerte:

Auswertung:

Grid area for evaluation and calculations.

Erklärung: Der Strom, den die Batterie antreibt, durchfließt auch die Batterie selbst. Dabei setzen die durchflossenen Teile der Batterie (Anschlussklemmen, Elektroden, Elektrolyt) wie (fast) jeder Leiter dem Strom auch Widerstand entgegen.

In einer realen Batterie wird daher nicht nur eine Spannung erzeugt, sondern es fällt bei Stromfluss auch schon ein Teil der erzeugten Spannung über dem Widerstand der Batterie ab. Eine reale Batterie verhält sich daher so wie eine Serienschaltung aus einer idealen Batterie, die unabhängig von der Belastung immer die **Leerlauf-Spannung** U_0 erzeugt, und einem ohmschen **Innen-Widerstand** R_i :

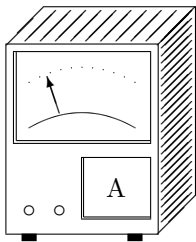


Von außen misst man zwischen den Klemmen der Batterie die **Klemmenspannung** U_K . Sie beträgt U_0 abzüglich der Spannung, die bereits über dem Innenwiderstand abfällt:

$$U_K = U_0 - R_i \cdot I$$

Dies ist natürlich gerade die Gleichung der im vorigen Versuch erhaltenen Geraden. Ermittle aus dem Schaubild die Leerlaufspannung und den Innenwiderstand unserer Batterie! Was versteht man wohl unter der **Kurzschluss-Stromstärke** I_K ? Trage alle markanten Größen im Diagramm ein!

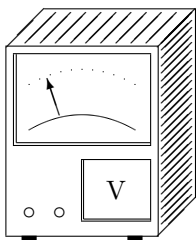
Natürlich gibt es einen Innenwiderstand nicht nur bei Batterien, sondern bei jeder Stromquelle. (Stromquellen, die ausschließlich aus Supraleitern bestehen, brauchen wir beim momentanen Stand der Entwicklung noch nicht zu betrachten.)



Amperemeter

Einen Innenwiderstand hat unvermeidlich auch jedes Amperemeter, denn es muss ja der zu messende Strom durch Leitungen des Messgerätes fließen, zum Beispiel durch den Hitzdraht eines Hitzdrahtinstrumentes oder durch die Windungen der Spule eines Drehspulinstrumentes.

Auch hier benimmt sich ein reales Amperemeter wie die Serienschaltung aus einem idealen Amperemeter ohne Widerstand und dem Innenwiderstand. Insbesondere verringert ein Amperemeter durch seinen Innenwiderstand den Strom, den es messen soll! Von einem guten Amperemeter erwartet man daher einen möglichst geringen Innenwiderstand.



Voltmeter

Du ahnst es bereits: Auch Voltmeter haben einen Innenwiderstand. Anders als bei Stromquellen und Amperemetern sind wir bei Voltmetern aber froh um den Widerstand: Hätte ein Voltmeter keinen Widerstand, so würde es seine beiden Anschlussstellen kurzschließen! Von einem guten Voltmeter erwarten wir daher sogar einen möglichst großen Widerstand, damit es praktisch keinen zusätzlichen Stromweg darstellt.

25.3.3 Spezifischer Widerstand

Bisher haben wir betrachtet, wie sich ein Bauteil aufgrund seines elektrischen Widerstandes in einem Schaltkreis verhält, welchen Einfluss der Widerstand auf Ströme und Spannungen hat. Diese Frage, wie sich Ströme und Spannungen aus Widerständen ergeben, wird mit Hilfe der Gleichung $R = \frac{U}{I}$ beantwortet.

Diese Gleichung sagt aber nichts darüber aus, wie

sich der Widerstand des Bauteils aus dem Aufbau des Bauteils ergibt. In der Praxis lassen wir Ströme meist nur durch metallene Leitungen (v.a. Kabel, Spulen, Leiterbahnen auf Platinen) fließen. Die ohmschen Widerstände sind dann letzten Endes immer die Widerstände von metallenen Leitungen. Wir fragen in diesem Abschnitt daher:

Wie ergibt sich der Widerstand einer Leitung aus ihrem Aufbau?

?

Wir wissen bereits, dass der Widerstand einer Leitung mit ihrer Länge l zunimmt. Wir wissen sogar schon genaueres: Ein 7 m langer Draht kann als Serienschaltung aus 7 Drahtstücken von jeweils 1 m Länge betrachtet werden. Folglich hat eine 7 m lange Leitung den 7-fachen Widerstand eines 1 m langen Stückes. Allgemeiner erwarten wir, dass der Widerstand einer Leitung proportional zu ihrer Länge ist.

$$\rightarrow R \sim l$$

Weiteren Aufschluss gibt uns die Parallelschaltung von Drahtstücken. Sieben parallel geschaltete Drahtstücke von 1 m Länge entsprechen einer Leitung mit 7-facher Querschnittsfläche. Nach den Regeln der Parallelschaltung hat diese Leitung nur ein Siebtel des Widerstandes der einfachen Leitung. Wir erwarten daher, dass der Widerstand einer Leitung umgekehrt proportional zu seiner Querschnittsfläche A ist.

$$\rightarrow R \sim \frac{1}{A}$$

Ob sich die 7 Drahtstücke der Länge nach berühren oder nicht, spielt keine Rolle: An einer Berührstelle würde mangels Spannung zwischen den betroffenen Drahtstellen sowieso kein Strom von einem Draht zum anderen fließen.

Wir erwarten also insgesamt $R \sim \frac{l}{A}$.

Wegen $R \sim l$ muss die gesuchte Formel für R die Form $R = \dots \cdot l$ haben. Andererseits trifft wegen $R \sim \frac{1}{A}$ auch $R = \dots \cdot \frac{1}{A}$ zu. Insgesamt muss daher $R = \dots \cdot l \cdot \frac{1}{A} = \dots \cdot \frac{l}{A}$.

Wovon kann R noch abhängen? Wir kennen noch zwei Einflüsse auf den Widerstand: Die Temperatur und das Leitermaterial. Wir haben zwar gesehen, dass der Widerstand eines Metalles i.d.R. mit der Temperatur zunimmt, aber erst bei großen Temperaturänderungen ($\rightarrow$ Glüh-Lämpchen) für uns nennenswert. Den Einfluss der Temperatur vernachlässigen wir daher einfach.

Es bleibt also nur noch der Einfluss des Leitermaterials. Bekanntlich gibt es ja verschieden gute Leiter. Der Proportionalitätsfaktor zwischen $\frac{l}{A}$ und R ist demnach eine Materialkonstante, nämlich der spezifische Widerstand ρ .

Einige Messungen bestätigen unsere Überlegungen:

Ein Leiter aus dem Material X mit der Länge l und der Querschnittsfläche A hat den ohmschen Widerstand

$$R =$$

Dabei ist ρ der spezifische Widerstand des Materials X.

Den spezifischen Widerstand gibt man üblicherweise in der Einheit $\frac{\Omega \cdot \text{mm}^2}{\text{m}}$. Querschnitte werden nämlich meist in mm^2 und Längen in m angegeben. Natürlich könnte man auch nach Kürzen der Einheiten $\Omega \cdot \text{m}$ als Einheit benutzen.

$$\left(\frac{\Omega \cdot \text{mm}^2}{\text{m}} = \frac{\Omega \cdot (0,001 \text{ m})^2}{\text{m}} = \frac{\Omega \cdot 0,001^2 \text{ m}^2}{\text{m}} = 10^{-6} \Omega \cdot \text{m} \right)$$

Diese neue Formel ist in größerem Zusammenhang so einzuordnen:



26.1 Motoren

26.1.1 Magnetfelder von Strömen

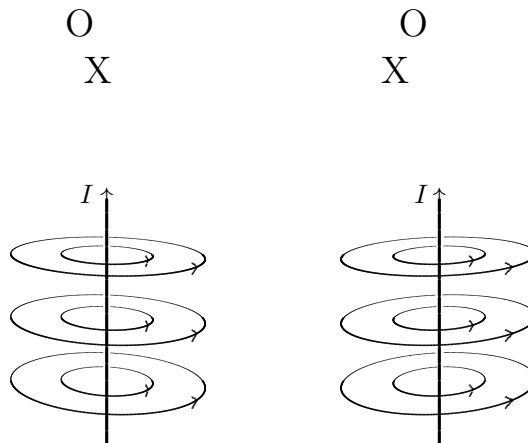
Elektrische Ströme sind bekanntlich ($\rightarrow$ 7. Schuljahr) von Magnetfeldern umgeben. Zur Auffrischung hier noch einmal einige Fakten über Magnete und Magnetfelder:

- Magnete kommen immer als Dipole, mit Nord- und Südpol vor.
- Ungleichnamige Pole ziehen sich an, gleichnamige stoßen sich ab.
- Man sagt, in einem Gebiet herrsche ein Magnetfeld, wenn dort Kräfte auf Magnete wirken.
- Feldlinien sind gedachte Linien, die nach Definition in der Richtung verlaufen, in die sich der Nordpol einer frei beweglichen Kompassnadel einstellen würde.
- Außerhalb eines Magneten verlaufen die Feldlinien seines Magnetfeldes daher von seinem Nordpol zum Südpol.
- Feldlinien kreuzen und berühren sich nie.
- Eisen, Nickel und Kobalt können magnetisiert werden.

Gerader Leiter

Man hat festgestellt, dass eine Kompassnadel auch in der Nähe von stromdurchflossenen Leitern ausgerichtet wird, dass also dort Magnetfelder herrschen. Indem man genauer beobachtet, wie sich Kompassnadeln oder Eisenfeilspäne ausrichten, erhält man für einen geraden Leiter das folgende Feldlinienbild.

Damit Du den räumlichen Verlauf der Feldlinien besser erkennen kannst, ist das Bild als Stereogramm ausgeführt. Wenn Du mit dem linken Auge auf das linke Bild und mit dem rechten Auge auf das rechte Bild schaust, kannst Du ein räumliches Bild sehen. Du musst die Augen also so stellen, als ob Du durch das Papier auf einen weiter entfernten Gegenstand schauen wolltest. Es kann helfen, zunächst wirklich auf einen entfernten Gegenstand zu schauen, und dann das Bild vor die Augen zu bringen, ohne jedoch die Augenstellung zu verändern. Falls Du die gedruckten zwei Bilder jeweils doppelt siehst, hast Du schon fast gewonnen: Verändere dann Deine Augenstellung nur noch so, dass sich die beiden mittleren von den insgesamt vier Bildern genau überdecken. Das dabei entstehende mittlere Bild ist das räumliche. Falls Du das ganze an einem einfacheren Motiv üben willst, versuche es zunächst mit diesen O's und X'en:



Die Feldlinien sind also Kreise um den Leiter in Ebenen senkrecht zum Leiter. Bei Umkehr der Stromrichtung kehrt sich auch die Richtung der Magnetfeldlinien um. Ihre Form bleibt ansonsten gleich. Natürlich ist das Magnetfeld bei größerer Stromstärke stärker. Von

einem Nord- oder Südpol kann man hier nicht sprechen, denn es gibt keine markanten Stellen, an denen Feldlinien aus- bzw. eintreten. Über die Richtung der Feldlinien gibt es aber die

Rechte-Faust-Regel: Umschließt man den Leiter mit der rechten Faust so, dass der Daumen in (technische) Stromrichtung zeigt, dann zeigen die gekrümmten Finger in Richtung des Magnetfeldes.

Leiterschleife

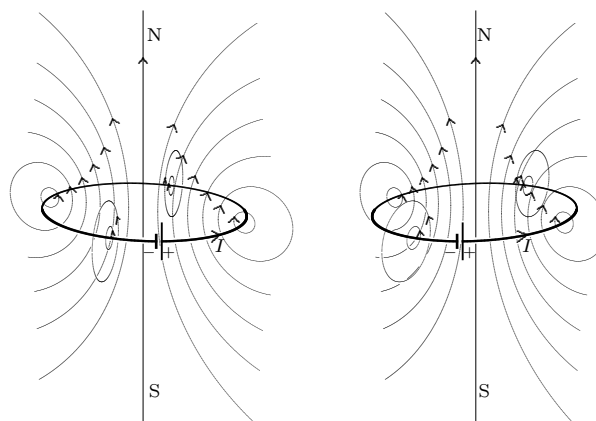
Auch, wenn man einen geraden Leiter zu einer Schleife biegt, behält jedes Abschnittchen des Leiters sein Magnetfeld um sich herum. In der folgenden Abbildung ist dies dadurch angedeutet, dass an vier Stellen jeweils zwei kreisförmige Feldlinien um ein Leiterstück herum gezeichnet sind, und zwar in dem Umlaufsinn, der sich nach der Rechten-Faust-Regel ergibt.

Die Fläche, die von der Schleife in ihrem Inneren aufgespannt wird, wird demnach offensichtlich von allen Feldlinien in derselben Richtung durchsetzt. (in der Abbildung von unten nach oben) Die Feld-Beiträge je-

des Leiter-Abschnittchens addieren sich im Inneren der Schleife zu einem recht starken Feld durch die Schleife.

Daher ähnelt eine stromdurchflossene Schleife einem herkömmlichen Magneten. In der folgend abgebildeten Situation hat die Schleife auf ihrer Oberseite – wo ja ein Bündel Feldlinien austritt – einen Nordpol, an der Unterseite analog einen Südpol. Außerhalb der Schleife verlaufen die Feldlinien in großen Bögen zur anderen Seite der Schleife.

Zur leichteren 3D-Übersicht entspricht die gezeichnete Feldliniendichte nicht der Feldstärke.



Man kann sich vorstellen, dass der Magnetismus mancher Stoffe (v.a. Eisen) auf die Magnetfelder inneratomarer Kreisströme als Elementarmagnete zurückgeführt werden kann.

Spule

Mehrere Schleifen hintereinander bilden eine Spule. Die Felder der einzelnen Schleifen sind im Inneren der Spule gleichgerichtet und verstärken sich demnach zu einem noch stärkeren Feld. Es ergibt sich im Inneren langer Spulen ein Magnetfeld, dessen Feldlinien parallel zur Spulenachse und gleichmäßig über den Spulenquerschnitt verteilt verlaufen.

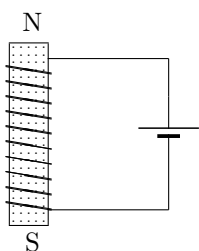
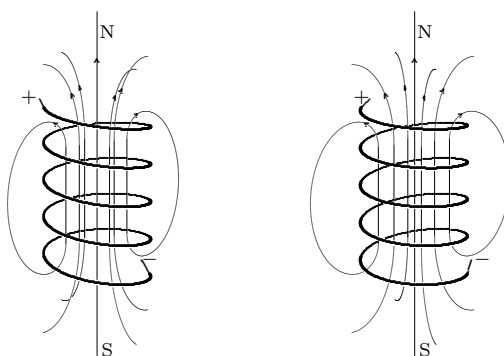
Ein solches Feld mit parallelen, gleichmäßig dichten Feldlinien nennt man **homogen**.

In Bezug auf das äußere Magnetfeld ist eine Spule nicht von einem Stabmagneten zu unterscheiden: An einem Ende (im Bild oben) treten Feldlinien aus und

verlaufen weitschweifig zum anderen Ende (unten). Die beiden Öffnungen der Spule entsprechen als Nord- und Südpol der Spule genau den Enden eines Stabmagneten.

Welches Ende der Spule welcher Pol ist, hängt natürlich über die Rechte-Faust-Regel davon ab, in welchem Sinn der Strom die Spulenachse umläuft. (Von welchem Ende der Strom her kommt, ist egal! Bei mehrlagigen Spulen ist der Draht sowieso mehrfach zwischen den Enden hin- und hergewickelt. Nur der Wicklungssinn zählt!) Für Spulen gibt es folgende Variante der Rechten-Faust-Regel:

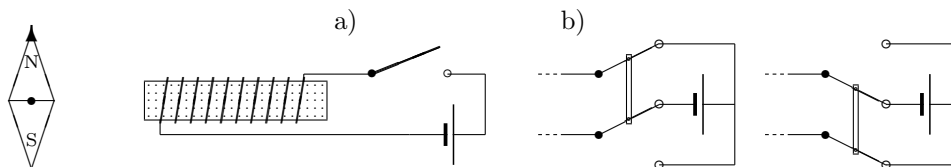
Umschließt man eine Spule mit der rechten Faust so, dass die gekrümmten Finger in Richtung des Stromes zeigen, dann zeigt der Daumen den Nordpol der Spule.



Füllt man eine Spule mit einem Eisenkern, so wird dieses Eisenstück durch das Magnetfeld des Spulenstromes gleichsinnig zu diesem Feld magnetisiert. Insgesamt ergibt sich so im Vergleich mit der eisenlosen Spule ein erheblich (z.B. tausendfach) verstärktes Magnetfeld. Die gesamte Anordnung aus Spule und Eisenkern nennt man einen **Elektromagneten**.

26.1.2 Gleichstrom-Motor

Es soll nun mit Hilfe elektrisch erzeugter Magnetfelder eine Drehbewegung hervorgerufen werden. Die ersten Versuche dazu könnten so aussehen, wie folgend von oben gesehen dargestellt. Eine drehbar gelagerte Magnetnadel steht vor einem Elektromagneten, der über Schaltung a) ein- und ausgeschaltet werden kann.



Beim Einschalten bildet sich am linken Ende des Elektromagneten ein Südpol. Die Magnetnadel wird folglich mit ihrem Nordpol zu diesem Ende hingezogen, dreht sich also um 90° im Uhrzeigersinn. Bleibt der Magnet eingeschaltet, so pendelt sich die Nadel in diese Stellung parallel zum Elektromagneten ein. Der Elektromagnet hält sie in dieser Ausrichtung fest.

Soll die Nadel sich weiterdrehen, muss der Magnet rechtzeitig ausgeschaltet werden, damit die Nadel sich mit ihrem Schwung über die Parallelstellung hinaus drehen kann. Überschreitet sie dann nach einer halben Drehung die Stellung, in der ihr Nordpol ganz links liegt, so kann man den Magneten wieder einschalten, um die Nadel erneut im Uhrzeigersinn anzutreiben. Durch geschicktes Ein- und Ausschalten des Elektromagneten lässt sich also eine dauernde Drehbewegung der Nadel herbeiführen.

Den Mangel, dass bei jeder Umdrehung eine halbe Umdrehung lang keine Kraft wirkt, könnte man noch mit Schaltung b) beheben. In ihr wird durch einen doppelten Wechselschalter der Magnet nicht ein- und ausgeschaltet, sondern umgepolt. So kann er für jede Stellung der Nadel die richtige Polung erhalten, mit der er die Nadel weiter im Uhrzeigersinn antreibt. Mit der Nadel könnte natürlich noch irgendetwas verbunden sein, was wir mit unserem primitiven Motor drehen wollen, z.B. ein Ventilator.

In einem Elektromotor muss also der Strom für einen Elektromagneten in den richtigen Momenten umgepolt werden. Von Hand wie in dem gerade beschriebenen Schulversuch kann man dies natürlich nicht immer machen. Es ist auch nur in Spezialfällen eine Lösung, einfach Wechselspannung zu verwenden: die Wechsel der Spannung müssten sich ja der Drehzahl des Motors anpassen können, die meist zumindest von der Belastung abhängt.

Da die Umpolungen bei gewissen Stellungen des

rotierenden Motorteils („Rotor“) erfolgen müssen, benutzt man einfach das rotierende Teil als Schalter für den Motorstrom. Natürlich dürfen an den Rotor keine Kabel von außen befestigt werden, die sich dann ja beim Drehen aufwickeln würden. Die elektrischen Verbindungen zwischen rotierendem und stehendem Teil („Stator“) werden über **Schleifkontakte** hergestellt. Praktisch sind dies z.B. Kupferstücke, Kohlestücke oder Drahtbürsten, die mit Federkraft vom feststehenden Geräteteil gegen ringförmige Kontaktflächen auf dem Rotor gedrückt werden.

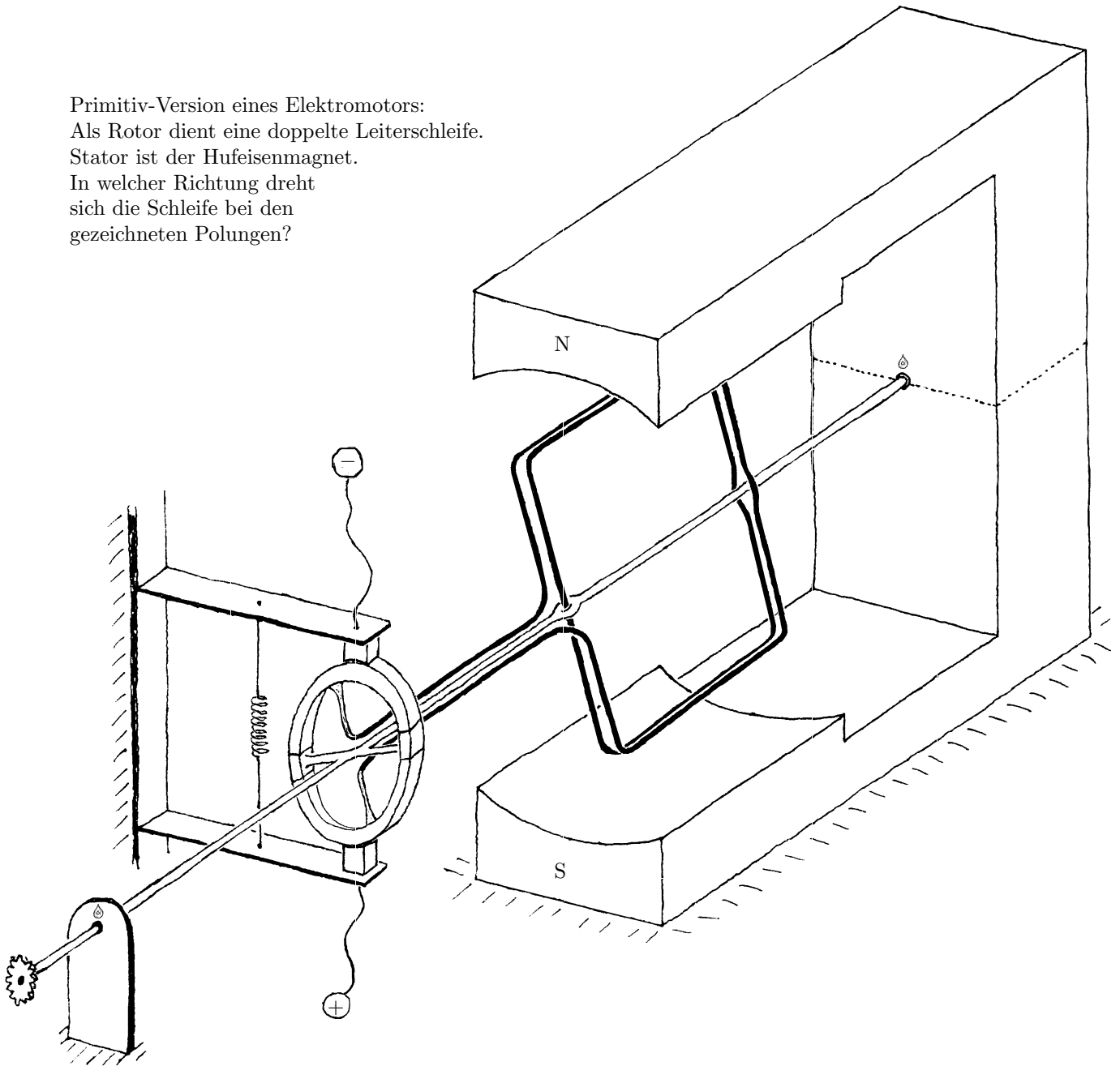
Um Strom von der Batterie über den Rotor zum Stator-Elektromagneten und über den Rotor zurück zu leiten, bräuchte man insgesamt vier Schleifkontakte. (Batterie $\rightarrow$ Rotor $\rightarrow$ Stator $\rightarrow$ Rotor $\rightarrow$ Batterie) Da Schleifkontakte aufgrund der ständigen Reibung hohen Verschleiß aufweisen und daher viel Wartungsaufwand verursachen, möchte man jedoch möglichst wenig Schleifkontakte haben. Man baut deshalb statt des Stators lieber den Rotor als unpolbaren Elektromagneten. Als Stator kann ein Elektromagnet oder ein Dauermagnet verwendet werden.

So kommt man mit nur zwei Schleifkontakten aus. (Batterie $\rightarrow$ Rotor $\rightarrow$ Batterie) Die Kontaktflächen auf dem Rotor sind zwei elektrisch isolierte Halb-Ringe, die mit jeweils einem Ende des Elektromagneten verbunden sind. Sie schleifen an zwei gegenüberliegenden festen Kontakten vorbei, die jeweils mit einem Pol der Stromquelle verbunden sind. Bei der Drehung schleifen dann die beiden Halbringe abwechselnd an den Batterieanschlüssen vorbei, so dass auch die Stromrichtung im Rotor wechselt. Bei richtiger Anordnung der Halbringe und der festen Kontakte erfolgen die Umpolungen sogar im richtigen Moment ;-) Da die beiden Halbringe für das Vertauschen der Pol-Anschlüsse verantwortlich sind, nennt man sie zusammen einen **Kommutator**.

Vorschlag zur Färbung des folgenden Bildes:

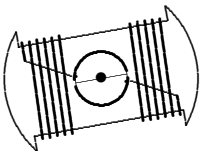
- Achse mit Halterung des Kommutators, der Isolierschicht zwischen den Halbringern sowie Zahnradchen: gelb
- oberer / unterer Halbring mit anliegendem Leiterstück: grün / orange
- oberer / unterer Teil des Hufeisenmagnetes: rot / grün
- oberes / unteres Kontaktstück: blau / rot
- Lager: grau oder braun

Primitiv-Version eines Elektromotors:
 Als Rotor dient eine doppelte Leiterschleife.
 Stator ist der Hufeisenmagnet.
 In welcher Richtung dreht
 sich die Schleife bei den
 gezeichneten Polungen?



Zu echtem Gebrauch gebaute Elektromotoren haben als Rotor (mindestens) eine Spule mit Eisenkern. Dieser Eisenkern wird bei Stromfluss magnetisiert und erfährt folglich Kräfte im Magnetfeld des Stator-Magneten, wie das bei Magneten bekanntlich so ist. Üblicherweise nennt man bei richtigen Motoren diesen Rotor einen **Anker**.

Bei der einfachsten, hier am Rand abgebildeten Ausführung spricht man von einem „Doppel-T-Anker“ – der Form des Ankers entsprechend.

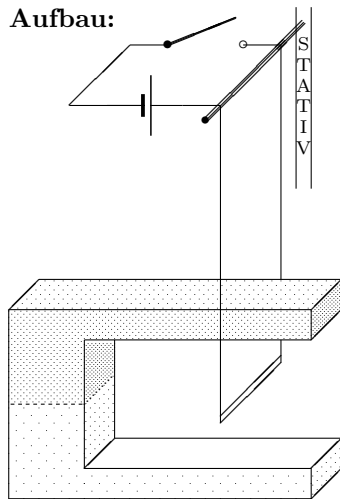


26.1.3 Lorentzkraft

Erstaunen könnte jedoch, dass auch der gezeichnete Motor *ohne* Eisenkern grundsätzlich funktioniert. Es gibt in der Leiterschleife kein einziges Stück Material, das magnetisiert wird! Es gibt zwar ein Magnetfeld, aber Kräfte können eigentlich nur an etwas Materiellem angreifen. Die Kräfte, die die Schleife drehen, können daher nur an dem zur Schleife gebogenen Leiter selbst angreifen. Ob tatsächlich ein einfaches stromdurchflossenes Leiterstück aufgrund eines Magnetfeldes eine Kraft erfahren kann, zeigt der folgende Versuch.

V 26.1 Leiterschaukel /I

Aufbau:



Ein gerades Leiterstück hängt an zwei Kabeln in einen Hufeisenmagnet hinein.

Durchführung: Der Schalter wird geschlossen. (andersfarbig dazuzichnen!)

Beobachtung:

(andersfarbig dazuzichnen!)

zu ergänzen:

- Polung des Magneten
- Magnetfeldlinien
- Stromrichtung

Verändert man die Richtung von Strom oder Magnetfeld (Polungen, Winkel), so ändert sich auch die Lorentzkraft. Beispiel: Jede Umpolung kehrt die Kraftrichtung um. Insgesamt erhält man:

Ein stromdurchflossener Leiter kann aufgrund eines Magnetfeldes eine Kraft erfahren, die man **Lorentzkraft** nennt, und die folgende Eigenschaften hat:

- Die Lorentzkraft steht immer senkrecht auf den Richtungen sowohl des Stromes als auch des Magnetfeldes.
- Es gilt die 3-Finger-Regel der rechten Hand: Zeigt der Daumen in Richtung des Stromes, der Zeigefinger in Richtung des Magnetfeldes, so zeigt der Mittelfinger in Richtung der Lorentzkraft.
- Die Lorentzkraft hat maximalen Betrag, wenn Strom und Magnetfeld senkrecht aufeinander stehen. Sie verschwindet, wenn Strom und Feld parallel verlaufen. (Anm.: Der Betrag ist proportional zum Sinus des Winkels zwischen Strom und Feld.)

Die 3-Finger-Regel nennt man auch UVW-Regel. Die drei Buchstaben sind dabei den drei Fingern zugeordnet und stehen

- für den **Strom** als die **Ursache** des Phänomens,
- der über das **Magnetfeld** als **Vermittlung**
- eine **Kraft** als seine **Wirkung** hervorruft.

Nun solltest Du über die Lorentzkraft erklären können, warum sich der Motor auf Seite 222 dreht.

Hendrik Antoon
Lorentz, * 18.7.1853
in Arnheim, Nederl.
† 4.2.1928, Haarlem

Da wir erst dann eine Lorentzkraft auf einen Leiter im Magnetfeld wirken sehen, wenn der Leiter von Strom durchflossen wird, scheint die Lorentzkraft doch eher auf den Strom als auf den Leiter zu wirken. Da der elektrische Strom im Leiter letzten Endes nur eine Bewegung von Elektronen ist, sind es wohl ursprünglich die bewegten Elektronen, die Lorentzkräfte erfahren. Die strömenden Elektronen werden dadurch seitlich abgelenkt und ziehen über die elektrostatische Anziehung zwischen sich (– immer negativ –) und dem restlichen

(– daher positiven –) Atomverband des Leitermaterials den Leiter hinter sich mit.

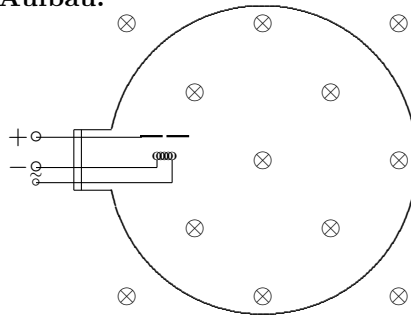
Der folgende Versuch zeigt, dass einzelne Elektronen bei Bewegung quer zu einem Magnetfeld tatsächlich so abgelenkt werden, wie es nach der UVW-Regel zu erwarten ist. Ein Strahl von Elektronen aus einer Elektronenkanone wird dadurch sichtbar, dass die Elektronen in einem geeigneten Gas mit Gasatomen zusammenstoßen und diese zum Leuchten anregen.

V 26.2 Fadenstrahlrohr

zu ergänzen:

- Elektronenbahn

Aufbau:



Durchführung: Das ganze Rohr wird von einem Magnetfeld durchsetzt.

Beobachtung: wie gezeichnet

⊗ = Magnetfeldlinie, die in die Zeichenebene hinein verläuft
 ⊙ = Magnetfeldlinie, die aus der Zeichenebene heraus verläuft

26.2 Generatoren

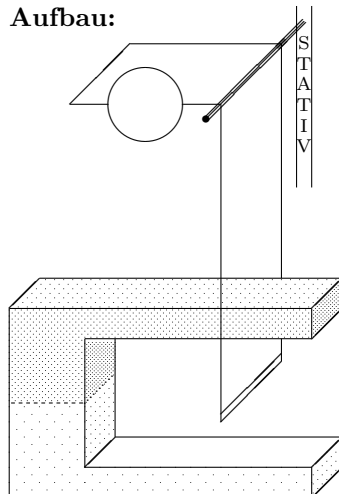
26.2.1 Induktion im bewegten Leiter

V 26.3 Leiterschaukel /II

zu ergänzen:

- Polung des Magneten
- Magnetfeldlinien
- Messgerät-Symbol

Aufbau:



Ein gerades Leiterstück hängt an zwei Kabeln in einen Hufeisenmagnet hinein.

Durchführung: Die Leiterschaukel wird ausgelenkt (andersfarbig dazuzichnen!)

Beobachtung: Während der Bewegung

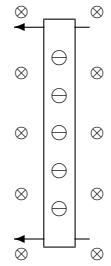
Der Effekt ist noch deutlicher zu sehen, wenn man statt eines einzigen Leiterstückes ein mehrfach um einen Schenkel des Hufeisenmagneten gewundenes Kabel am Schenkel entlang bewegt.

Bei umgedrehtem Magnetfeld oder umgekehrter Bewegungsrichtung

Bei schnellerer Bewegung

Erklärung der Beobachtung: In dem bewegten Leiterstück befinden sich wie in jedem Metall viele frei bewegliche Elektronen. Bewegt man den Leiter, so bewegen sich diese Elektronen natürlich mit. Nun wissen wir aber seit kurzem, dass Ladungsträger, die sich quer zu Magnetfeldlinien bewegen, Lorentzkräfte erfahren.

Wende doch einmal in dem nebenstehenden Bild eines nach links bewegten Leiters die UVW-Regel an, um die Richtung der Lorentzkraft auf ein mitbewegtes Elektron ($\ominus$) zu ermitteln! (Bedenke dabei, dass die Bewegung von negativ geladenen Teilchen nach links einem rechnerischen Ladungsstrom nach rechts entspricht!) Die Magnetfeldlinien ($\otimes$) verlaufen senkrecht zur Zeichenebene von Deiner Seite her in diese hinein.

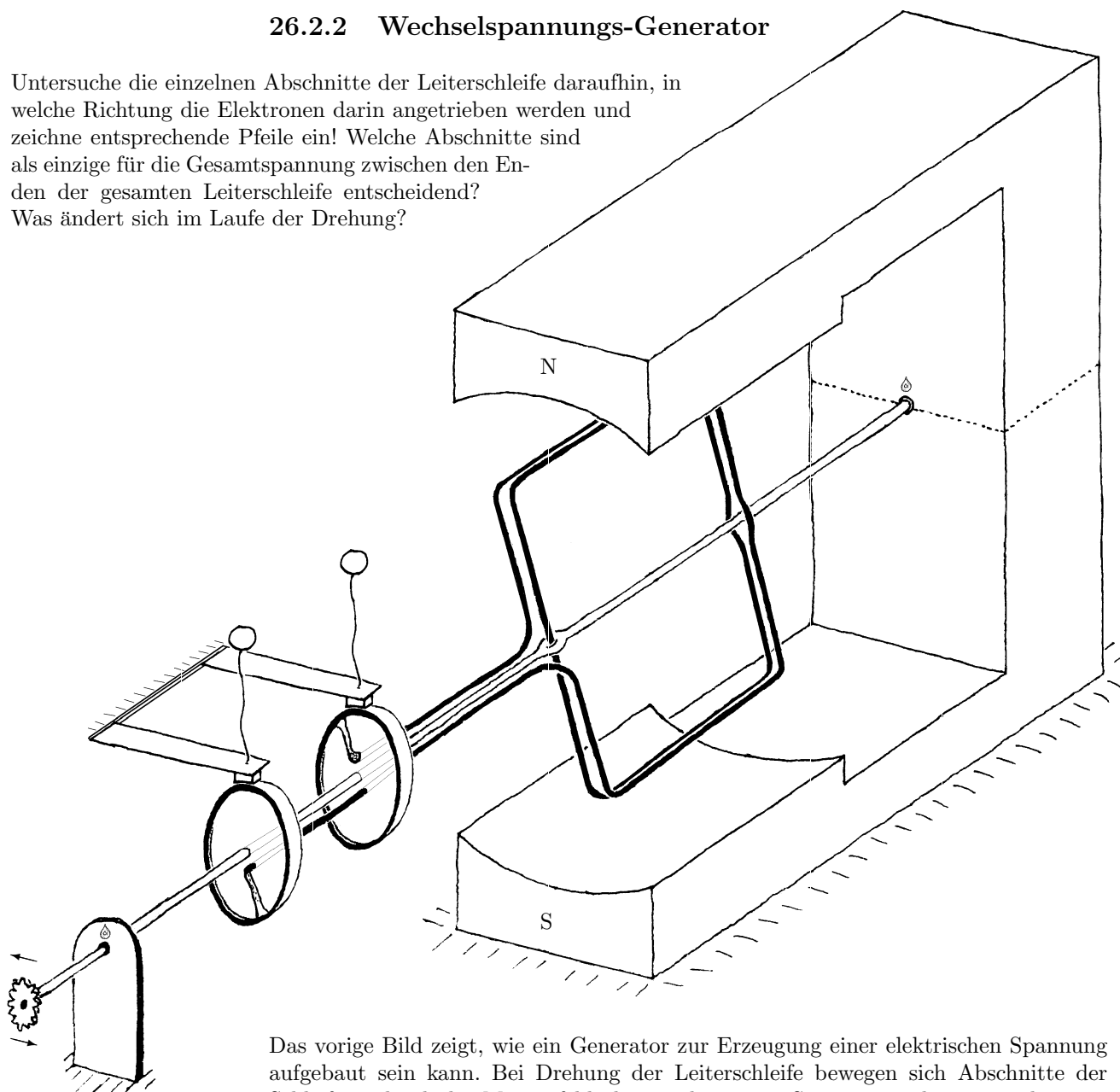


Fertig? Die Lorentzkraft treibt die Elektronen also entlang des Leiters nach oben. Das Leiterstück wird dadurch zu einer Spannungsquelle: das obere Ende mit dem Elektronenüberschuss ist ihr Minuspol, das untere Ende mit dem Elektronenmangel ihr Pluspol. Dieses Phänomen, dass bei der Bewegung eines Leiters durch ein Magnetfeld eine Spannung entsteht, nennt man

Induktion. Man sagt, es wird eine Spannung *induziert*, nämlich die *Induktionsspannung*. Selbstverständlich kann eine Induktionsspannung auch einen *Induktionsstrom* antreiben, wenn das bewegte Leiterstück in einem geschlossenen Leiterkreis eingebaut ist. Im gezeichneten Beispiel flösse dann ein Induktionsstrom von oben nach unten durch das Leiterstück.

26.2.2 Wechselspannungs-Generator

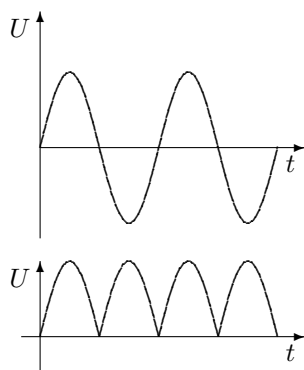
Untersuche die einzelnen Abschnitte der Leiterschleife daraufhin, in welche Richtung die Elektronen darin angetrieben werden und zeichne entsprechende Pfeile ein! Welche Abschnitte sind als einzige für die Gesamtspannung zwischen den Enden der gesamten Leiterschleife entscheidend? Was ändert sich im Laufe der Drehung?



Das vorige Bild zeigt, wie ein Generator zur Erzeugung einer elektrischen Spannung aufgebaut sein kann. Bei Drehung der Leiterschleife bewegen sich Abschnitte der Schleife so durch das Magnetfeld, dass in ihnen eine Spannung induziert wird.

Nur die Abschnitte parallel zur Drehachse tragen etwas zur Gesamtspannung bei, weil nur bei diesen die Lorentzkräfte entlang des Leiters wirken. In den anderen Abschnitten, die senkrecht zur Achse verlaufen, wirken die Lorentzkräfte senkrecht zum Leiter. Die parallelen Abschnitte sind als Spannungsquellen betrachtet entlang der Leiterschleife gleichsinnig in Serie geschaltet, so dass sich ihre Spannungen addieren.

Im Laufe der Drehung ändert sich der Winkel zwischen der Bewegungsrichtung eines Abschnittes und dem Magnetfeld und damit auch der Betrag der Lorentzkraft. Alle halbe Drehung kehren sich die Bewegungsrichtungen der Abschnitte um, so dass auch ihre Spannungen die Polung umkehren. Zwischen den Enden der Leiterschleife greift man also über die Schleifkontakte eine (sinusförmige) **Wechselspannung** ab. Das nebenstehende Diagramm zeigt den zeitlichen Verlauf dieser Spannung während zweier Umdrehungen.



Will man Gleichspannung abgreifen, so hilft derselbe Umpol-Trick wie beim Elektromotor: Abgriff über einen Kommutator statt über zwei Schleifringe! Man erhält allerdings keine konstante Spannung, sondern **pulsierende Gleichspannung**, wie sie das nebenstehende Diagramm zeigt.

Natürlich enthalten tatsächlich genutzte Generatoren nicht nur – wie gezeichnet – zwei Leiterschleifen, sondern Spulen mit Eisenkern. Wie schon beim Elektromotor erhält man auch beim Generator dadurch einen stärkeren Effekt.

26.2.3 Kenngrößen von Wechselspannung/-strom

Legt man an einen ohmschen Widerstand eine Wechselspannung an, so fließt ein Strom, dessen zeitlicher Verlauf dieselbe Form wie der Verlauf der Spannung hat. Die im Folgenden nur anhand der Spannung eingeführten Begriffe sind auf die Stromstärke übertragbar.

Die gesamte (Sinus-)Kurve setzt sich aus der nahtlosen Aneinanderreihung des Wellenberg-&-tal-Musters  zusammen.

Dieses Muster ist durch zwei Parameter festgelegt:

- Die zeitliche Länge T des Musters ist die Zeitspanne, die bei einer vollständigen Umdrehung der Generatorspule vergeht. Diese Länge nennt man je nach Sichtweise
Schwingungsdauer (Ein sich dauernd wiederholender Vorgang wird von Physikern allgemein „Schwingung“ genannt.),
Umlaufzeit (wenn man an die rotierende Spule denkt) oder
Periode (Eine Funktion $D \rightarrow \mathbb{R}$ heißt periodisch mit Periode p , wenn für jedes $x \in D$ auch $x \pm p \in D$ und $f(x) = f(x + p)$ gilt.)
- Die maximale Spannung $\hat{U}$ nennt man
Scheitelspannung (wird erreicht am Scheitel des Berges) oder
Amplitude (lateinisch *amplus*=weit, hier: Schwingungsweite)

Eine bequeme Alternative zur Angabe der Schwingungsdauer ist die Angabe der **Frequenz** f . Dies ist die Anzahl der Schwingungen pro Zeit. Da *eine* Schwingung pro Schwingungsdauer stattfindet, gilt natürlich die nebenstehende Gleichung Als Einheit für Frequenzen ist das Hertz definiert als $1 \text{ Hz} := \frac{1}{\text{s}}$.

$$f = \frac{1}{T}$$

Auch zur Angabe der Scheitelspannung gibt es eine Alternative nach folgenden Überlegungen: Betreibt man z.B. eine Glühlampe oder einen Tauchsieder, so interessiert meist deren Leistung. Bekanntlich errechnet sich die elektrische Leistung nach der Gleichung $P = U \cdot I$, die sich mit $I = U/R$ zu $P = U^2/R$ kombinieren lässt.

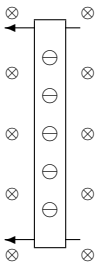
Da jedoch U ständig schwankt, schwankt auch die Leistung. Die Leistung $\hat{P} = \hat{U} \cdot \hat{I} = \hat{U}^2/R$ liegt dabei nur zu denjenigen einzelnen Zeitpunkten vor, an denen die Kurven für U und I maximale Beträge ausweisen. (an den „Bergspitzen“ und den „Talsohlen“) Zu allen anderen Zeiten ist die Leistung kleiner, manchmal sogar Null.

Für den Benutzer eines Gerätes interessanter als die Leistung zu einzelnen Zeitpunkten ist die durchschnittliche Leistung $\bar{P}$. Man kann experimentell feststellen, dass am selben Verbraucher eine Wechselspannung mit Scheitelspannung $\hat{U}$ zur selben mittleren Leistung führt wie eine konstante Gleichspannung $U_{\text{eff}} = \sqrt{1/2} \cdot \hat{U}$. In dieser algebraischen Genauigkeit erhält man diesen Wert jedoch nur durch Berechnung von $\bar{P}$, für die uns hier leider noch ein paar Mathematik-Kenntnisse fehlen. Man nennt U_{eff} die **Effektivspannung**. Es gelten die nebenstehenden Gleichungen. Wenn nichts anderes dazu gesagt wird, sind bei Wechselspannung alle Spannungs- und Stromstärken-Angaben immer Effektivwerte.

Trage im Diagramm auf Seite 226 die Größen T , $\hat{U}$ und U_{eff} ein!

$$U_{\text{eff}} = \sqrt{\frac{1}{2}} \cdot \hat{U}$$

$$I_{\text{eff}} = \sqrt{\frac{1}{2}} \cdot \hat{I}$$



26.2.4 Lenzsche Regel

Auf Seite 225 haben wir bei einem bewegten Leiterstück die Entstehung eines Induktionsstromes mit Lorentzkraften erklärt: Mit dem Leiterstück bewegen sich auch die darin enthaltenen Elektronen durch das Magnetfeld und erfahren eine Lorentzkraft nach oben. Bei geschlossenem Stromkreis wird so ein Strom von oben nach unten (technisch-rechnerische Stromrichtung) induziert. Ergänze das nebenstehende Bild wieder entsprechend!

Dieser Induktionsstrom hat natürlich zur Folge, dass der stromdurchflossene Leiter im Magnetfeld eine Lorentzkraft erfährt! Nach der UVW-Regel wirkt diese nach – halt! erst selbst überlegen, dann in der Fußnote weiterlesen, dann Bild ergänzen!¹

Die Bewegung des Leiters nach links wird durch diese Kraft nach rechts gehemmt. Der Induktionsstrom ist also so gerichtet, dass die von ihm hervorgerufene Lorentzkraft die für den Strom ursächliche Bewegung hemmt.

Dies muss auch schon wegen folgender Überlegung so sein: Mit dem induzierten Strom wird ja auch elektrische Energie erzeugt. Nach dem Energieerhaltungssatz kann diese Energie nicht aus dem Nichts entstehen, sondern muss der Anordnung bei der Bewegung des Leiters zugeführt werden. Das Bewegen des Leiters muss daher physikalische Arbeit sein, also Kraft in Bewegungsrichtung entlang des Weges erfordern. Diese erforderliche Kraft nach links ist natürlich die Kraft, die man gegen die nach rechts wirkende Lorentzkraft ausüben muss, damit das Leiterstück nicht von der Lorentzkraft abgebremst wird. Da diese Energie-Betrachtung bei jedem Induktionsvorgang zutrifft, gilt ganz allgemein die nebenstehende Regel.

Eine Generatorspule lässt sich also leicht drehen, wenn am Generator keine Verbraucher angeschlossen sind und folglich kein Induktionsstrom durch die Spule fließt. Bei angeschlossenen Verbrauchern fällt die Drehung umso schwerer und kostet also umso mehr Leistung, je mehr Strom und damit je mehr Leistung die Verbraucher aufnehmen.

Lenzsche Regel:
Jeder
Induktionsstrom ist
so gerichtet, dass er
seiner Ursache
entgegenwirkt.

¹rechts

26.3 Transformatoren

26.3.1 Induktion in ruhenden Leitern

Relativbewegung

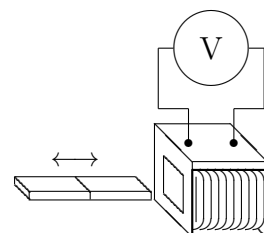
In einem Leiterstück, das sich durch ein Magnetfeld bewegt, wird bei passenden Richtungen eine Spannung induziert. Ein Beispiel dafür ist eine Leiterschaukel wie in früheren Versuchen. Ein anderes Beispiel ist der Metallstab, der in manchen Aufgaben auf Leiterschienen rollt. Ein weiteres Beispiel sind Abschnitte einer rotierenden Generatorspule.

Ob etwas sich bewegt oder nicht, ist jedoch Ansichtssache. Glaubst Du nicht? Hier eine Alltagserfahrung als Beleg: Angenommen, Du sitzt in einem stehenden Zug und betrachtest den ebenfalls stehenden Zug auf dem Nachbargleis. Wenn nun einer der beiden Züge ohne spürbaren Ruck zu rollen beginnt, wirst

Du aufgrund Deiner Beobachtung kaum sagen können, welcher der beiden Züge sich bewegt! Noch anders betrachtet haben sowieso beide Züge nie still gestanden, sondern sich stets zusammen mit der gesamten Erde mit fast 30 km/s um die Sonne bewegt.

Immer ist zur Beschreibung der Bewegung *eines* Körpers ein *zweiter* Körper erforderlich, den man als ruhend annimmt. (Im Beispiel: Bahnsteig oder Sonne) Da man die Bewegung eines einzelnen Körpers nie ohne Bezug zu anderen Körpern beschreiben kann, sind alle Bewegungen sogenannte **Relativbewegungen**. Das Wort *relativ* stammt aus dem Lateinischen und bedeutet etwa *rückbezogen* auf etwas.

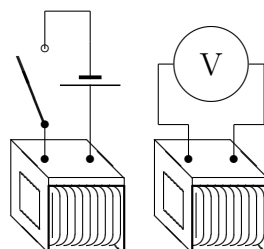
Auch bei allen bisherigen Induktionsversuchen kommt es nur auf die Relativbewegung zwischen Leiter und Magnetfeld an. Wir überzeugen uns davon, indem wir z.B. bei (relativ zum Tisch) ruhender Leiterschaukel den Hufeisenmagneten wegziehen und dabei ebenfalls Induktion feststellen: Hauptsache, es bewegen sich Magnetfeldlinien und Leiterstücke mit geeigneten Richtungen relativ zueinander. Noch deutlicher ist der Effekt, wenn wir in der Nachbarschaft einer Spule einen Magneten bewegen. Auch dabei findet eine Relativbewegung zwischen den magnetischen Feldlinien und den Leiterstücken der Spule statt.



Magnetischer Fluss

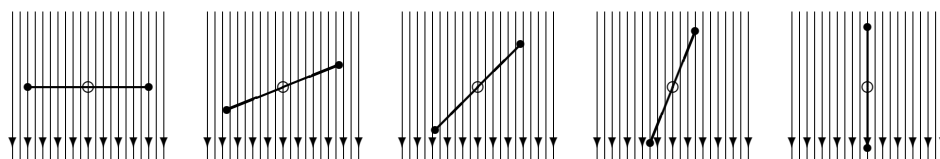
Dieser Versuch mit bewegtem Magneten und ruhender Spule bietet aber auch Anlass zu folgender Überlegung: Da die Spannung *in der Spule* induziert wird, können wir uns fragen, was sich eigentlich *in* der Spule ändert, wenn der Magnet *daneben* bewegt wird. Es kommt als Antwort zunächst die Stärke des Magnetfeldes in Frage. Demnach sollten wir in der Spule auch dadurch eine Spannung induzieren können, dass wir ohne irgendwelche Bewegungen nur die Stärke des Magnetfeldes in ihr ändern.

Dazu schalten wir einen Elektromagneten dicht neben der Spule ein und wieder aus. Tatsächlich wird bei beiden Schaltvorgängen jeweils eine Spannung induziert. Die Vorzeichen der beiden Induktionsspannungen unterscheiden sich dabei – eigentlich erwartungsgemäß.



Anders als alle vorher aufgetretenen Induktionsphänomene lassen sich diese Induktionsspannungen kaum mehr mit Lorentzkraften erklären. Ohne Bewegung gibt es ja keine Lorentzkraft! Statt der Bewegung könnte hier die Veränderung der Stärke des Magnetfeldes eine Rolle spielen. Im Folgenden lernst Du ein hierzu passendes Erklärungsmuster für Induktionsphänomene kennen. Dieses neue Muster beschreibt dabei sogar nicht nur die jüngsten Versuchsergebnisse, sondern auch alle vorigen Induktionsphänomene.

Betrachten wir die Induktion in einer Generatorspule! Auch hier würde sich nichts an der Induktion ändern, wenn sich statt der Spule das äußere Magnetfeld drehen würde. Im Gegensatz zum vorhergehenden Versuch ändert sich hier jedoch die Stärke des Magnetfeldes in der Spule nicht. Was ändert sich hier statt dessen? Betrachte dazu die folgende Bilderserie einer Leiterschleife, die sich in einem Magnetfeld dreht! (im Querschnitt senkrecht zur Drehachse, Feldlinien von oben nach unten, Drehachse o)

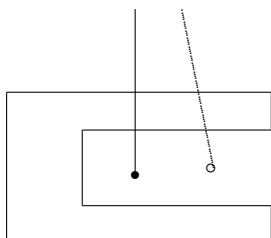


Mit dem Drehwinkel ändert sich ganz offensichtlich die Anzahl der Feldlinien, die die Leiterschleife durchsetzen. (Färbe diese Feldlinien!)

Auch für die anderen Induktionsversuche ist die Anzahl der Feldlinien durch eine Leiterschleife die entscheidende Größe: Bedenkt man, dass die Stärke eines Magnetfeldes durch die Dichte der Feldlinien dargestellt wird, sind Änderungen der Stärke des Feldes auch Änderungen der Feldliniendichte, und damit auch der Linienanzahl durch die von einer Leiterschleife umschlossene Fläche. (Feldliniendichte=Feldlinienanzahl pro durchsetzter Fläche)

Sogar Induktion bei bewegten Leitern lässt sich oft leicht mit der Feldlinienanzahl erklären. Beispielsweise bilden in den Aufgaben mit Leiterschienen und rollendem Metallstab die Schienen, der Stab und der restliche Stromkreis eine Schleife, die je nach Position des Stabes unterschiedlich viele Feldlinien umfasst.

Auch die Leiterschaukel umfasst bei verschiedenen Auslenkungen zusammen mit ihren Aufhängefäden unterschiedlich viele Feldlinien im oberen Schenkel des Hufeisenmagneten, wie nebenstehendes Bild zeigt. (Eingezeichnet ist die Schaukel an zwei verschiedenen Positionen. Ergänze die vollständigen Feldlinien, wie sie durch die Luft und das Eisen verlaufen!)



Nun ist „Anzahl der Feldlinien“ aber keine wohldefinierte physikalische Größe. Um eine solche zu erhalten, geht man wie folgt vor: Die Stärke eines Magnetfeldes an einer Stelle beschreibt man darüber, wie viel Kraft ein stromdurchflossener Leiter dort erfährt. In der Oberstufe kannst Du genauer lernen, wie man als Maß für die Stärke eines Magnetfeldes die „magnetische Flussdichte B “ definiert als Kraft pro Leiterlänge und Stromstärke. Der Name Fluss*dichte* ist schon ein Hinweis darauf, dass diese Größe dasselbe beschreibt wie die Feldliniend*ichte*.

Der Anzahl an Feldlinien, berechnet als Feldliniendichte $\cdot$ durchsetzte Fläche, entspricht dann der **magnetische Fluss** (Φ), berechnet als Flussdichte $B \cdot$ durchsetzte Fläche A . Vorläufig brauchst Du Dir aber nur zu merken:

magnetischer Fluss
entspricht
Anzahl an Feldlinien

Induktionsgesetz: Ändert sich der magnetische Fluss durch eine Leiterschleife, so wird entlang der Schleife eine Spannung proportional zur Änderungsrate des Flusses (Änderung des Flusses pro Zeit) induziert.

Dies beinhaltet, dass die Richtung eines Induktionsstromes davon abhängt, ob der Fluss zunimmt oder abnimmt. Es sei auch darauf hingewiesen, dass auch mit einer kleinen Flussänderung eine hohe Spannung induziert werden kann, wenn die Änderung nur schnell genug erfolgt. Dies tritt insbesondere bei abruptem Ausschalten von Elektromagneten auf.

26.3.2 Transformator

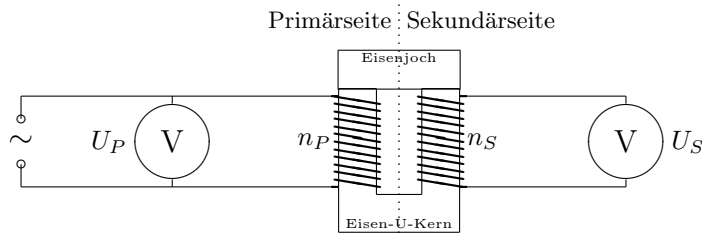
Spannungstransformation

Wechselnd starken magnetischen Fluss erhält man natürlich auch, wenn die felderzeugende Spule von Wechselstrom durchflossen wird. Wir untersuchen im Folgenden, welche Spannung dabei in einer zweiten Spule induziert wird, wenn diese vom selben Feld durchsetzt wird, das die erste Spule erzeugt. Um sicherzustellen, dass beide Spulen – die felderzeugende und die Induktionsspule – vom selben magnetischen Fluss durchsetzt werden, werden beide Spulen auf einen gemeinsamen, geschlossenen Eisenkern montiert. Als Schaltsymbol zeichnet man diese Anordnung wie nebenstehend zu sehen, wobei allgemein der Eisenkern in einer Spule als einfacher Strich neben der Spule gezeichnet wird.



V 26.4 Spannungen am Transformator

Aufbau:



n_P, n_S : Windungszahlen der Primär-/Sekundärspule

U_P, U_S : Effektivwerte der primär-/sekundärseitigen Wechselspannung

Durchführung: Es werden zunächst die zeitlichen Verläufe der Primär- und der Sekundärspannung auf einem Oszilloskop verglichen.

Beobachtung: Bei sinusförmiger Primärspannung

ist die Sekundärspannung

Durchführung: U_P wird variiert, U_S mit einem Voltmeter gemessen.

Messwerte: bei $n_P = \dots$ und $n_S = \dots$

U_P/V					
U_S/V					

Ergebnis: Offensichtlich gilt:

Durchführung: Es werden verschiedene Spulen als Primär- oder Sekundärspule ausprobiert.

Messwerte: /Auswertung:

n_P										
n_S										
U_P/V										
U_S/V										

Ergebnis:

Damit erklärt sich auch der Name **Transformator** für die benutzte Anordnung von zwei Spulen auf gemeinsamem Eisenkern: Es lässt sich durch die Auswahl der Windungszahlen eine Wechselspannung auf einen (weitgehend) beliebigen anderen Effektivwert transformieren. Weitere hierfür gebräuchliche Verben sind „übersetzen“ und „umspannen“. Dem Belieben sind jedoch z.B. durch die unvermeidlichen ohmschen Widerstände des Spulenmaterials Grenzen gesetzt, von denen wir hier gänzlich absehen.

Um das Übersetzungsverhältnis leicht ändern zu können, ohne ganze Spulen auszutauschen, verwendet man als Sekundärspule eine Spule mit einem verschiebbaren Abgriff, ähnlich wie bei einem Schiebewiderstand. Je nachdem, ob man mehr oder weniger Windungen zwischen diesem Abgriff und dem fest angeschlossenen Ende der Spule einschließt, ist die abgegriffene Sekundärspannung größer oder kleiner. Dies ist ein typischer Aufbau von Netzgeräten.

Für ein vertieftes Verständnis der Spannungen am Transformator ist folgende Sichtweise günstig: Der magnetische Fluss verläuft ja (nahezu) vollständig nur durch den Eisenkern. Die Magnetfeldlinien sind in sich geschlossene Linien, die einen Umlauf durch das Eisen machen. Jede einzelne Windung der beiden beteiligten Spulen umfasst daher dieselben Feldlinien, und damit denselben magnetischen Fluss. Bei einer Änderung dieses Flusses wird daher in jeder Windung dieselbe Spannung induziert. Über einer ganzen Spule als Serienschaltung ihrer n einzelnen Windungen herrscht daher die n -fache Spannung wie über einer einzelnen Windung. Daher müssen sich die in den Spulen induzierten Spannungen stets wie deren Windungszahlen verhalten.

Bisher war allerdings nur sekundärseitig von einer Induktionsspannung die Rede. Tatsächlich findet Induktion aber auch in der Primärspule statt. Das Induktionsgesetz von Seite 230 kennt nämlich keinen Unterschied zwischen Primär- und Sekundärspulen: Leiterschleife ist Leiterschleife und Fluss ist Fluss! Auch in den Leiterschleifen, mit denen ein Fluss erzeugt wird, wird bei Flussänderung eine Spannung induziert.

Die in der Primärspule induzierte Spannung ist dabei jederzeit gerade betragsgleich der von außen angelegten Wechselspannung, jedoch umgekehrt gepolt. Warum? Dazu folgende Überlegung:

Treibt eine Spannungsquelle Elektronen des Stromkreises z.B. mit 3 Volt an, so verrichtet sie ja pro Coulomb fließender Ladung 3 Joule an Arbeit. Werden diese 3 Joule nicht vollständig dafür gebraucht, die Widerstände im Stromkreis zu überwinden, so werden die Elektronen weiter beschleunigt, so dass die Stromstärke steigt. Reichen die 3 Joule nicht zur Überwindung der Widerstände, so sinkt die Stromstärke entsprechend.

Mit der Stromstärke ändern sich aber meist auch die zur Überwindung von Widerständen nötigen Spannungen.

Bei ohmschen Widerständen ist bekanntlich $U = R \cdot I$. Höhere Stromstärke führt hier also auch zu mehr benötigten Joule pro fließendem Coulomb. Jede Ungleichheit in der Stärke von Antrieb und Hemmung des Stroms führt also über veränderte Stromstärken zu einer Angleichung von Antrieb und Hemmung: Ist z.B. die Hemmung schwächer als der Antrieb, so steigt die Stromstärke und mit ihr die Hemmung. Daher stellt sich immer (blitzschnell!) eine Stromstärke ein, bei der die Spannungsabfälle über den Strom-hemmenden Bauteilen gleich der Antriebs-Spannung sind. (Dies ist nur eine neue Formulierung der altbekannten Maschenregel.) Ähnliche Überlegungen haben wir auf Seite 213 schon einmal anhand eines Luftstromkreises mit Turbinen angestellt.

Auch eine Spule mit vernachlässigbarem ohmschen Widerstand leistet einem Wechselstrom einen („induktiven“) Widerstand. Steigt z.B. die Stromstärke in ihr an, so wächst auch der magnetische Fluss in ihr. Diese Änderung des Flusses führt zu einer Induktionsspannung, die nach Lenzscher Regel dem Anwachsen der Stromstärke entgegenwirkt. Die Wechselspannungsquelle kann daher die Stromstärke durch die Spule nur so stark steigern, dass die induzierte Gegenspannung der Quellenspannung genau betragsgleich wird.

Bei z.B. 200 V Effektivspannung an einer Primärspule mit 50 Windungen verändert sich die Primär-Stromstärke jederzeit so, dass vom zugehörigen magnetischen Fluss in jeder Windung der Primärspule effektiv 4 V induziert werden. Derselbe magnetische Fluss induziert dann natürlich auch in jeder Sekundärwindung 4 V.

Stromstärken am Transformator

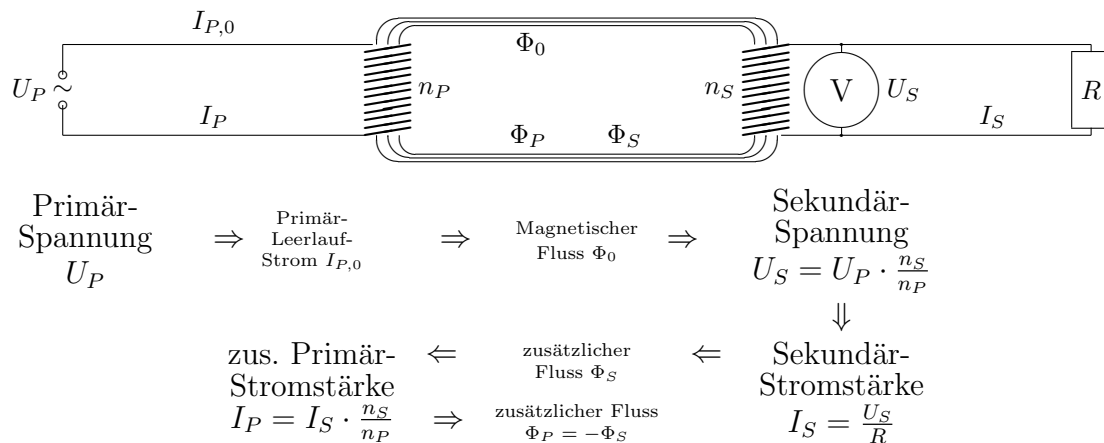
Sind die Anschlüsse der Sekundärspule eines Transformators nicht leitend miteinander verbunden, ist also insbesondere kein Verbraucher angeschlossen, so spricht man von einem unbelasteten Transformator oder einem Transformator „im Leerlauf“. Primärseitig fließt dann ein Leerlaufstrom bestimmter Stärke I_{P0} . Nach obigen vertieften Überlegungen ist dieser Primärstrom gerade so groß, dass die von ihm verursachten Flussänderungen in der Primärspule eine Gegenspannung zu U_P induzieren. Auf der offenen Sekundärseite ist natürlich $I_S = 0$.

Bei Belastung des Transformators, d.h. Anschluss eines Verbrauchers an die Sekundärspule, fließt ein Sekundärstrom, dessen Stärke I_S sich aus dem angeschlossenen Widerstand und U_S ergibt. Da I_S auch die Sekundärspule durchfließt, erzeugt I_S auch einen magnetischen Fluss im Eisenkern. Die Änderung dieses Flusses induziert dabei natürlich auch wieder eine Spannung in der Primärspule, wodurch sich auch der Primärstrom ändert. Und das hat natürlich wieder Auswirkung auf den magnetischen Fluss.

Was sich jetzt vielleicht nach einem schwer berechenbaren Wechselspiel von Hin- und Her-Induktionen anhört, ist aber tatsächlich nicht ganz so schlimm: Nach wie vor muss die in der Primärspule induzierte Spannung (bis auf ihr Vorzeichen) gleich U_P sein. Die für die Induktion verantwortlichen Änderungen des magnetischen Flusses müssen daher im Endeffekt dieselben sein wie im unbelasteten Zustand des Trafos. Aller Fluss, der durch den Sekundärstrom zusätzlich erzeugt wird, wird daher durch Veränderungen des Primärstroms genau kompensiert. Der Primärstrom, der demnach zusätzlich zum Leerlaufstrom I_{P0} auftritt, heiße I_P .

Nun bedenken wir noch, dass ein Strom in *jeder* Windung, die er durchfließt, magnetischen Fluss im Eisenkern erzeugt. Es wirkt also I_S in der Sekundärspule n_S -fach und I_P in der Primärspule n_P -fach. Zum gegenseitigen Ausgleich muss also gelten:

$$I_P \cdot n_P = I_S \cdot n_S$$

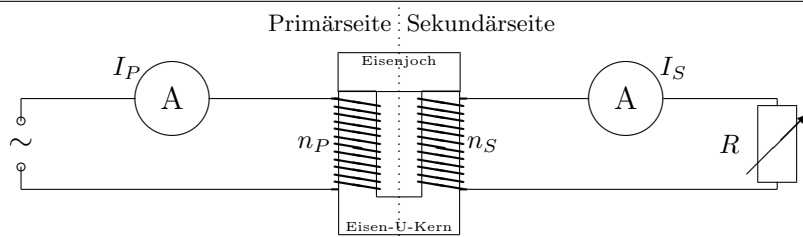


Experimentell werden wir im 10. Schuljahr diesen Zusammenhang jedoch nicht genau bestätigen können. Der Leerlaufstrom I_{P0} und der durch die Belastung verursachte Zusatzstrom I_P sind nämlich Wechselströme, deren $I(t)$ -Kurven zeitlich versetzt verlaufende Wellenlinien sind. Die beiden Ströme erreichen also z.B. ihre Maximalwerte nicht gleichzeitig. Wir können deshalb die effektiven – ebenso wenig wie die maximalen – Stromstärken nicht einfach addieren oder subtrahieren. Insbesondere können wir nicht einfach von einem bei Belastung gemessenen effektiven $I_{P,gesamt}$ ein vorher im Leerlauf gemessenes I_{P0} abziehen um den Zusatzstrom I_P zu erhalten.

Wir können jedoch dafür sorgen, dass I_{P0} keine große Rolle gegenüber I_P spielt, indem wir I_P relativ groß machen. Dazu betreiben wir einen Transformator mit relativ kleinen Widerständen an der Sekundärspule, so dass dort möglichst große Ströme fließen. Im Extremfall schließen wir die Sekundärseite sogar einfach kurz, so dass nur die unvermeidlichen Widerstände der Spulen den Strom begrenzen. Unter diesen Bedingungen betrachten wir im folgenden Versuch I_{P0} als 0 und identifizieren $I_{P,gesamt}$ mit I_P .

V 26.5 Stromstärken am Transformator

Aufbau:



Messwerte: /Auswertung:

$R = \dots\dots\dots$

n_P										
n_S										
I_P/A										
I_S/A										

Ergebnis:

Energiebetrachtung

Aus $\frac{U_P}{U_S} = \frac{n_P}{n_S}$ und $\frac{n_P}{n_S} = \frac{I_S}{I_P}$ folgt $\frac{U_P}{U_S} = \frac{I_S}{I_P}$ und damit $U_P \cdot I_P = U_S \cdot I_S$. In den Produkten aus Spannung und Stromstärke erkennen wir natürlich sofort elektrische Leistungen.² Die gewonnene Gleichung drückt damit den Energieerhaltungssatz aus: Sofern der Transformator verlustfrei arbeitet, gibt er sekundärseitig dieselbe Leistung ab, die er primärseitig aufnimmt.

26.3.3 Anwendungen

²Verlaufen die in Wirklichkeit gemessenen $I(t)$ und $U(t)$ zeitlich versetzt zueinander, muss $\bar{P} = U_{\text{eff}} \cdot I_{\text{eff}}$ zwar um einen Korrekturfaktor ergänzt werden, jedoch können wir unter Beibehaltung der bisherigen Idealisierungen diesen Faktor beiderseits als gleich annehmen und außer Acht lassen.

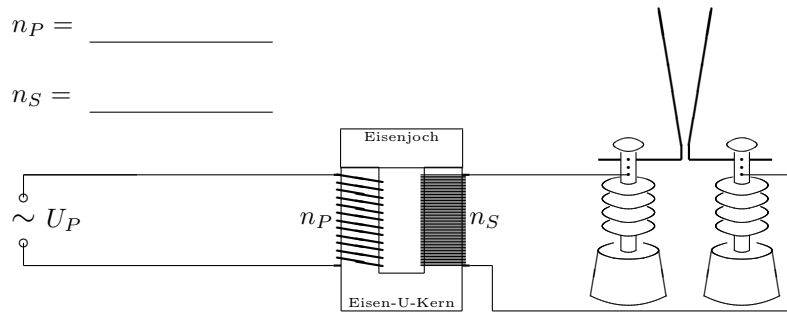
V 26.6 Hochspannung

zu ergänzen:

- Beobachtung

Aufbau: $n_P =$ _____

$n_S =$ _____



Durchführung: Wechselspannung $U_P =$ _____ V einschalten

Beobachtung: siehe Bild

Erklärung:

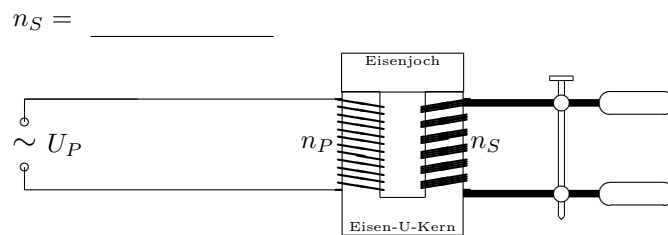
V 26.7 Hochstrom

zu ergänzen:

- Beobachtung

Aufbau: $n_P =$ _____

Sekundärseitig eingespannt:
ein dicker Nagel



Durchführung: Wechselspannung $U_P =$ _____ V einschalten

Beobachtung: siehe Bild

Erklärung:

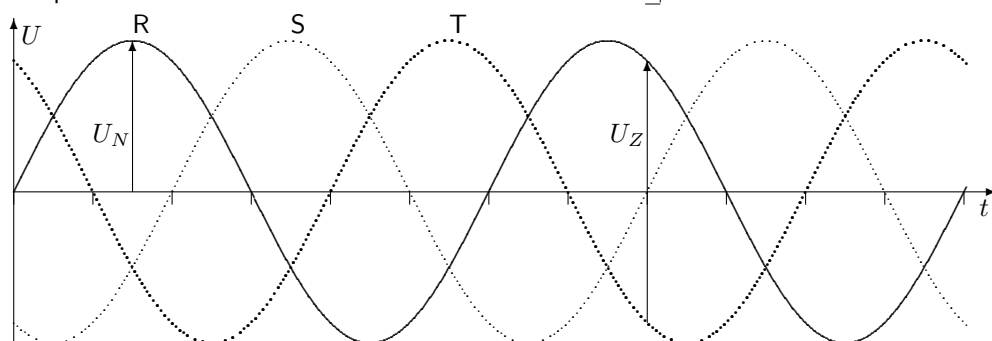
26.4 Unsere Elektrizitätsversorgung

26.4.1 Drehstrom

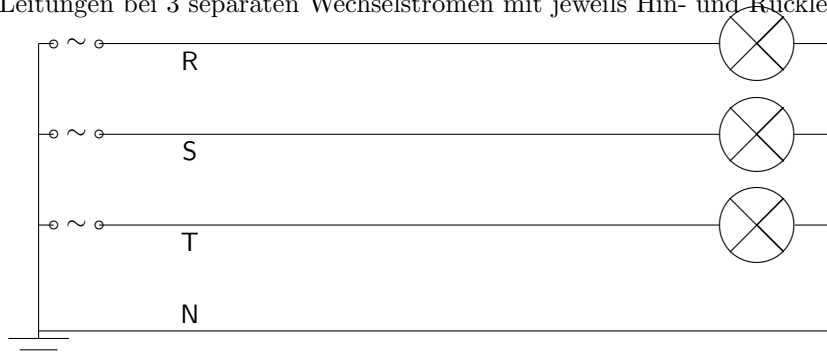
Unser Stromversorgungsnetz ist ein Drehstromnetz: In den (meisten) Kraftwerken besteht ein Generator aus einem Elektromagneten, der zwischen drei fest montierten Spulen rotiert und durch seine Bewegung in jeder dieser Spulen eine Wechselspannung induziert. Die Gesamtübersicht der Stromversorgung auf Seite 237 zeigt oben rechts, wie ein solcher Drehstromgenerator schematisch aufgebaut ist.

Nebenbei: Je nachdem, ob bei einem Generator der Magnet, in dessen Feld die Induktion stattfindet, fest montiert außen steht oder innen rotiert, nennt man den Generator eine Außen- bzw. Innenpolmaschine. Der Generator auf Seite 226 ist also eine Außenpolmaschine, ein Kraftwerksgenerator meist eine Innenpolmaschine.

Da die Spulen um 120° gegeneinander versetzt angeordnet sind, sind auch die drei erzeugten Wechselspannungen zeitlich um eine Drittel Periode gegeneinander versetzt. Das folgende Diagramm zeigt den zeitlichen Verlauf dieser Spannungen. Die drei Generatorspulen sind an einem Ende miteinander verbunden und an dieser gemeinsamen Verbindung (=„Sternpunkt“) auch geerdet. Die dargestellten Spannungen sind also auch die Spannungen (oder „Potenziale“), die die anderen Enden der drei Spulen gegenüber der Erde haben. Die von diesen Enden ausgehenden, gegenüber der Erde spannungsführenden Leitungen werden **Phasenleiter** genannt und im folgenden mit R, S und T bezeichnet. (Alternativ übliche Bezeichnung: L1, L2 und L3)



Wenn man an die drei Wechselspannungen Verbraucher direkt anschließen wollte, könnte man dies nach der folgenden Schaltung mit nur vier Leitungen zwischen den Quellen und den Verbrauchern tun: (statt 6 Leitungen bei 3 separaten Wechselströmen mit jeweils Hin- und Rückleitung)



Ein Vorteil bei der Verwendung von Drehstrom ist nun der, dass bei gleichartigen Verbrauchern die Summe der drei Stromstärken jederzeit Null ergibt. Fließt z.B. durch Leitung R gerade die Scheitelstromstärke nach rechts, so sind die Stromstärken in S und T gerade nach links gerichtet und in der Summe betragsgleich der in R. Am rechten Rand der Schaltung gleichen sich also alle drei Stromstärken aus, so dass durch den **Neutralleiter** N überhaupt kein Strom fließt.

Wären die Wechselströme nicht zeitlich versetzt (sondern „in Phase“), müsste durch N die betragsmäßi-

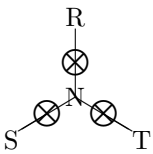
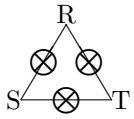
ge Summe der drei Ströme in R, S und T zurückfließen. Um zu starke Erhitzung zu vermeiden bräuchte man also für N eine Leitung mit dreifacher Querschnittsfläche. Für R,S,T und N zusammen bräuchte man dann insgesamt doch 6-fachen Querschnitt, genau wie für drei separate Wechselströme.

Sind bei unterschiedlichen Verbrauchern auch die Stromstärken durch R, S und T unterschiedlich, dann heben sich die Stromstärken am rechten Rand immer noch zumindest teilweise auf, so dass nur ein Rest durch N zurückfließen muss.

Ein weiterer Vorteil von Drehstrom ist, dass ohne zusätzlichen Aufwand Wechselspannungen mit zwei verschiedenen Scheitelwerten zur Verfügung stehen:

- Zwischen jedem der Phasenleiter und dem Neutralleiter besteht eine Wechselspannung mit Scheitelwert U_N .
- Zwischen zwei Phasenleitern kann eine Wechselspannung mit Scheitelwert U_Z abgegriffen werden.

Dabei ist $U_Z = \sqrt{3} \cdot U_N$, also rund 1,7-fach größer. Zeichne im vorigen Schaltplan noch Verbraucher dazu, die mit U_Z versorgt werden!



Es ist günstig, Elektrogeräte so zu bauen, dass sie aus drei gleichartigen Einheiten (z.B. 3 Spulen in einem Motor, 3 Heizelemente in einem Herd) bestehen, die in Stern- oder Dreiecksschaltung angeschlossen werden. (Worin liegt der Unterschied?) Im Haushalt werden v.a. Elektroherde an „Starkstrom“, d.h. an alle 4 Leitungen des Drehstromes (plus Schutz-Erdung als 5. Leitung) angeschlossen.

26.4.2 Übertragung elektrischer Energie

In Wirklichkeit gibt es nicht wie im vorigen Schaltplan durchgehende Leitungen von den Kraftwerken zu den Verbrauchern. Aufgrund der großen Entfernungen und der entsprechend hohen Leitungswiderstände dazwischen wäre dies nicht sinnvoll. Die Leistungsverluste in den Leitungen ergeben sich aus dem Leitungswiderstand R_L und der Stromstärke I bekanntlich zu $P_L = R_L \cdot I^2$. (Herzuleiten aus $R = U/I$ und $P = U \cdot I$)

Zur Verringerung der Verlustleistung kann nicht der Widerstand R_L beliebig verkleinert werden, weil Leitungen nicht beliebig dicker gemacht werden können. Die Leitungen würden viel zu schwer und insbesondere ihre Verlegung viel zu teuer. Mit Transformatoren kann jedoch die Stromstärke verringert werden, was ja die Leistung sogar mit 2 potenziert günstig beeinflusst.

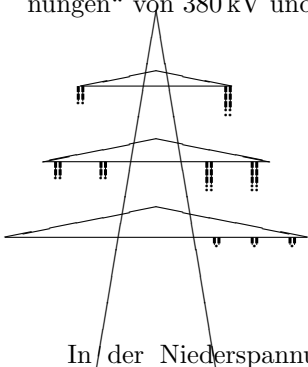
Unsere elektrische Energie wird also über Hochspannungsleitungen übertragen. (Verluste: unter 5%) Je nach Entfernung und zu übertragender Leistung werden verschiedene Spannungen benutzt: „Höchstspannungen“ von 380 kV und 220 kV, „Hochspannung“ von

110 kV und „Mittelspannungen“ von 20 kV und 10 kV.

Je nach ihrer Größe werden Kraftwerke an ein Netz der Hoch- oder Höchstspannungsleitungen angeschlossen, Groß-Verbraucher an ein Mittelspannungsnetz und sonstige Verbraucher an ein Niederspannungsnetz (400/230 V), wie es also auch in unseren Steckdosen endet. Umspannanlagen verbinden die genannten Netzebenen und versorgen in der Regel eine niedrigere Netzebene mit Energie aus der nächsthöheren Ebene. Die folgende Übersicht ist auf nur eine Hochspannungsebene (mittlere „Zeile“ des Bildes) zwischen Kraftwerk (oben rechts) und Haushalt (untere Bildhälfte) vereinfacht.

Wie dieses Bild zeigt, werden zur Übertragung auf der Hochspannungsebene nur drei Leitungen benötigt: Selbst wenn die Verbraucher aus den drei Phasen unterschiedlich viel Strom beziehen, verteilt sich diese Last in den Transformatoren der Umspannanlagen auf die drei Hochspannungsleitungen.

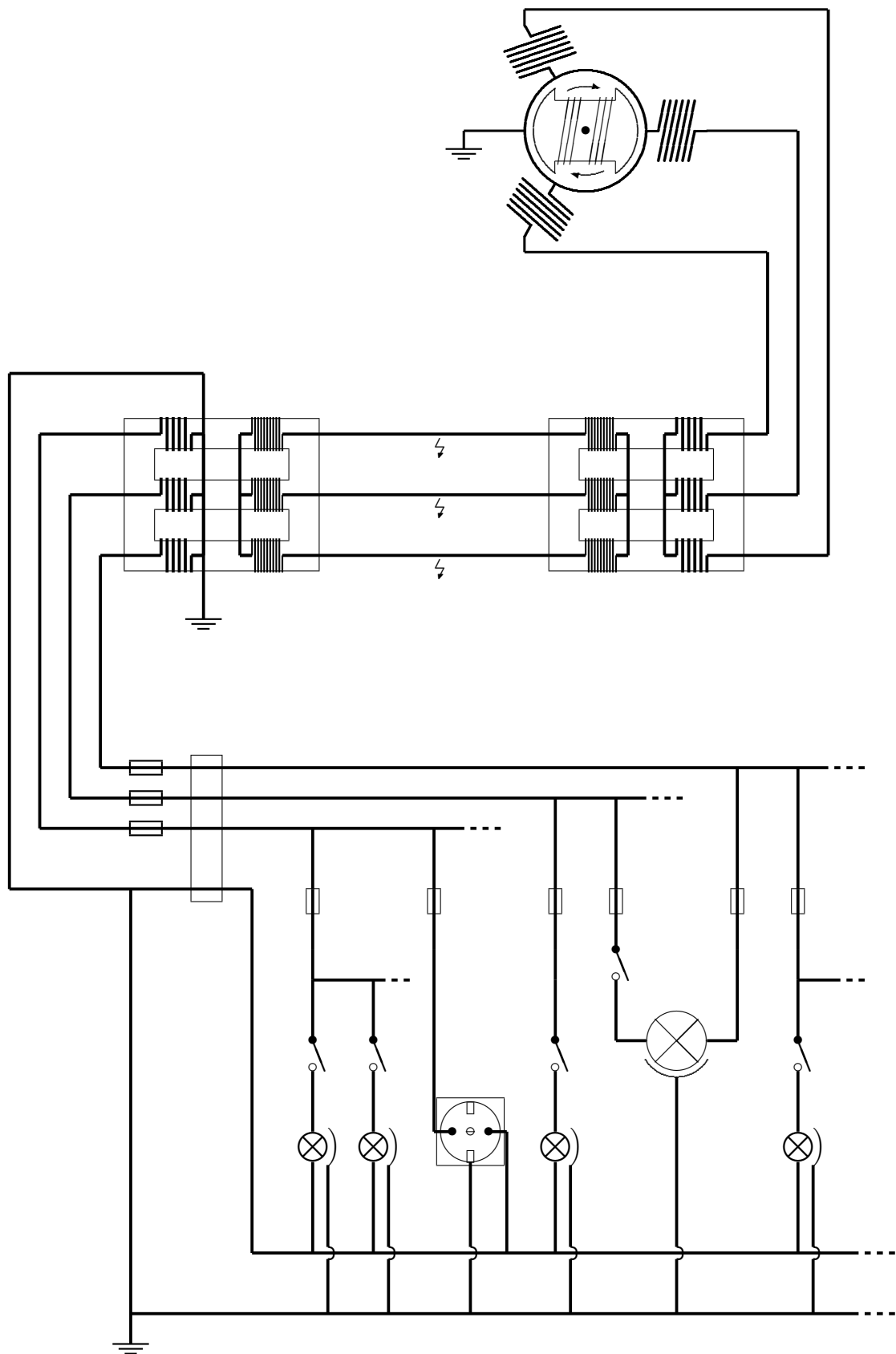
Über die Hochspannungsmaste werden daher immer 3 Freileitungen pro Stromkreis verlegt. Für jeden einzelnen 380 kV-Phasenleiter werden dabei 4 Draht-Seile wie Eckpunkte eines Quadrates angeordnet (///) an eine Kette aus 3 Isolatorstäben aufgehängt. (Durch diese 4-seilige Leiter-Ausführung werden die elektrischen Feldstärken an den Leiteroberflächen verringert, um Ladungsübertritte in die Luft zu erschweren. Die 3 Isolatoren dienen dazu, die Leitungen vom Mastgestänge weit genug weg zu halten.) Bei 220 kV-Leitungen genügen 2 Isolatoren und 2 Seile, bei 110 kV-Leitungen ein Isolator und ein Seil. Der hier skizzierte Mast trägt von jeder dieser Spannungsebenen je einen Stromkreis.



In der Niederspannungsebene wird eine geerdete vierte Leitung als Neutralleiter mitgeführt. So kommen in jeden Haushalt 4 Leitungen. Eine fünfte Leitung, die über das Fundament des Hauses selbst geerdet wird, vervollständigt die Elektro-Installation des Hauses.

Jede der Phasen ist über eine Schmelzsicherung gegen übermäßige Stromstärke geschützt. In ihren Verzweigungen, die die verschiedenen Teile des Hauses ver-

sorgen, sitzt ein weiterer Sicherungsschalter. Ein FI-Schalter vergleicht zudem, ob die Summe der Ströme durch die 3 Phasen und den Neutralleiter Null ergibt. Ist dies nicht der Fall, bedeutet dies, dass die Differenz als „Fehlerstrom“ (= „Fehler-I“) über den Erd-Leiter fließt, was einen Defekt an einem Verbraucher anzeigt. Der FI-Schalter schaltet dann den Strom ab.



26.4.3 Verbundnetz

Die Übersicht auf Seite 237 ist natürlich auch so vereinfacht, dass nur *ein* Kraftwerk und *ein* Haushalt dargestellt sind. Tatsächlich beziehen wir aber zusammen mit einigen Hundert Millionen anderen Menschen jährlich einige Milliarden Kilowattstunden elektrischer Energie aus einem elektrischen Verbundnetz, das sich über weite Teile Europas bis nach Nordafrika und in die Türkei erstreckt. Es wird organisiert durch den Verband *European Network of Transmission System Operators for Electricity*, kurz *ENTSO-E*.



In allen Leitungen des gesamten Verbundes auf dem Festland pulsieren die drei Wechselspannungen des Drehstromnetzes synchron. Elektrisch gesehen bildet dieser Teil Europas und seiner Nachbarn ein „Synchrongebiet“. Wenn in einer unserer Steckdosen z.B. gerade der Scheitelwert der Spannung erreicht wird, wird er im selben Moment auch in sämtlichen an dieselbe Phase angeschlossenen Steckdosen von Portugal bis Polen, von Griechenland bis Jütland erreicht.

Es sind nämlich alle Generatoren quasi parallelgeschaltet und stimmen folglich im zeitlichen Spannungsverlauf überein. Abweichungen von der synchronen Drehung führen über veränderte Induktionsspannungen und veränderte Stromstärken zu Veränderungen der Lorentzkräfte in den Generatoren, die der Abweichung automatisch entgegenwirken.

Der Vorteil eines so großen Verbundes liegt darin, dass Schwankungen in der Produktion und im Verbrauch elektrischer Energie besser ausgeglichen werden können. Elektrische Energie kann ja kaum auf Vorrat hergestellt und gelagert werden, bis sie gebraucht wird, sondern muss genau in dem Moment erzeugt werden, in dem sie gebraucht wird.

Wären nur wenige Verbraucher an ein Netz angeschlossen, so würde jedes Ein- und Ausschalten eines Verbrauchers die Leistungsnachfrage im Netz verhältnismäßig stark verändern. (Wenn man zu einer brennenden Glühbirne eine zweite einschaltet, wird 100% mehr Leistung gebraucht. Bei 100 Glühbirnen macht eine weitere nur 1% Mehrleistung aus.)

Entsprechend gilt: Würden nur wenige Kraftwerke das Netz versorgen, so brächte der Ausfall eines einzigen Kraftwerkes einen verhältnismäßig großen Leistungsmangel mit sich. In einem kleinen Netz müssen bei Ausfall eines Kraftwerkes wenige andere Kraftwerke ihre Leistungsabgabe drastisch steigern. In einem großen Netz brauchen viele Kraftwerke ihre Leistung nur ein wenig zu steigern.

Weiterer Vorteil: In einem großen Verbund können sich Schwankungen der Leistung einzelner Verbraucher auch mit höherer Wahrscheinlichkeit gegenseitig ausgleichen: Wenn Du die Kühlschranktür öffnest und dessen Beleuchtung dadurch angeht, gibt es mit recht hoher Wahrscheinlichkeit irgendwo in Europa auch jemanden, der etwa zu dieser Zeit eine ähnliche Beleuchtung gerade ausschaltet.

Wenn nun die Verbraucher mehr Leistung in Form von mehr Strom aus dem Netz beziehen, steigt die

Stromstärke in allen Generatoren des Verbundes an. Es kommt zu größeren Lorentzkäften, die die Drehung der Generatoren stärker hemmen. Die Drehzahl der Generatoren sinkt und mit ihr der Scheitelwert und die Frequenz der Wechselspannung. Die Verbraucher beziehen die Mehrleistung daher letzten Endes aus der Bewegungsenergie der Generatoren.

Natürlich soll die Spannung nach Frequenz und Betrag möglichst stabil bleiben. Bei uns sind Abweichungen von nur bis zu $\pm 0,01$ Hz die Regel. Bei einem Absinken der Frequenz wird in den Kraftwerken sofort „mehr Gas gegeben“, um die Generatoren wieder auf ihre Sollfrequenz zu bringen. Je nach Kraftwerkstyp wird also z.B. mehr Kohle oder Öl verbrannt. Reicht dies nicht aus, werden zusätzliche Generatoren eingeschaltet.

Zuletzt noch ein paar Worte dazu, was ein Wechsel des Stromlieferanten bedeutet. Alle Erzeuger speisen Energie in das Verbundnetz ein, aus dem die Verbraucher sie entnehmen. Das Netz ist in dieser Hinsicht wie ein See, in den viele Bäche Wasser einspeisen, und aus dem viele Lebewesen trinken, ohne dass sich der Weg einer bestimmten Portion Wasser durch den See verfolgen lässt.

Es speist einfach jeder Erzeuger soviel Energie ein, wie seine Kunden durchschnittlich verbrauchen. Jede Einspeisung eines Erzeugers wird dabei ebenso gemessen wie jede Entnahme der Verbraucher. Anhand dieser Messwerte wird dann geregelt, wer von wem wofür wie viel Geld bekommt. Außer, dass nach einem Wechsel der alte Erzeuger zukünftig vielleicht ein bisschen weniger, der neue ein bisschen mehr Energie ins Netz einspeist, ändern sich nur Finanzströme.

Energienutzung

27.1 Techniken zur Energiegewinnung

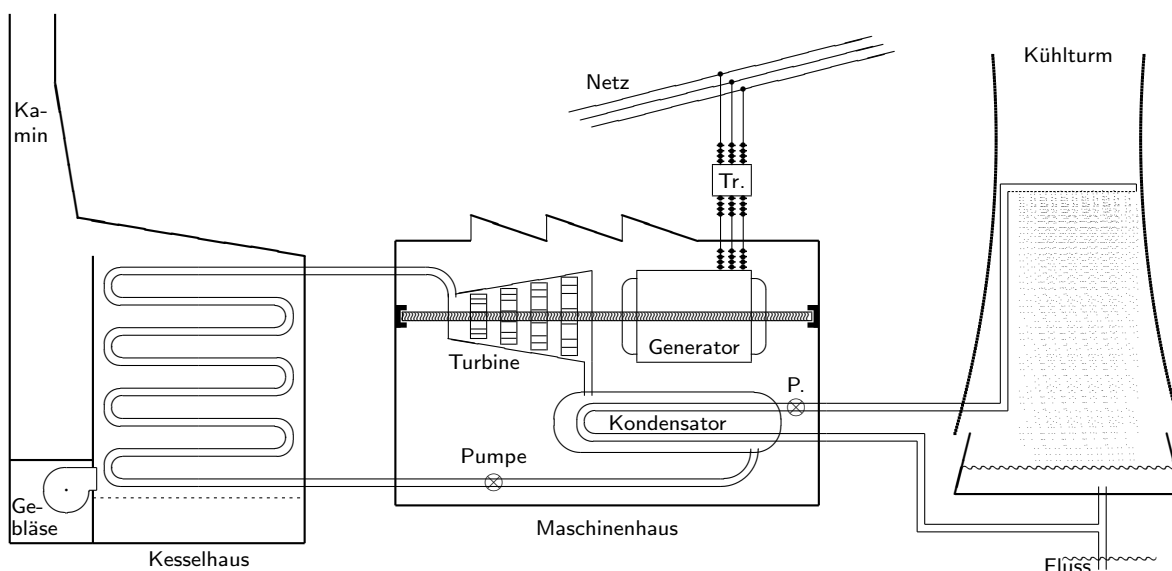
27.1.1 Dampfkraftwerke

Der größte Teil unserer elektrischen Energie kommt aus Kraftwerken, die im Wesentlichen so aufgebaut sind, wie hier skizziert. Ins turmhohe Kesselhaus werden Kohlestaub, Öl oder Gas zusammen mit Luft-Sauerstoff eingeblasen und verbrannt. In der dadurch entstehenden Hitze wird Wasser in einem großflächigen Röhrensystem verdampft und der Dampf weiter erhitzt. Er erreicht Temperaturen von über 500 °C und erzeugt damit in den abgeschlossenen Rohren Drücke von über 200 bar.

(Nicht eingezeichnet: Die Abgase werden gereinigt (ent-

staubt, entstickt, entschwefelt), bevor sie in die Atmosphäre gelassen werden.)

Mit solchem Hochdruck lässt man den Dampf nun durch Turbinen strömen. Die Turbinen sind im Prinzip Schaufelräder, die wie ein Windrad vom strömenden Dampf in Drehung versetzt werden. Bei der Arbeit an den nachgebenden Schaufelrädern entspannt sich der Dampf: Sein Volumen nimmt zu, (was sich auch in der Form der abgebildeten Turbine niederschlägt,) sein Druck und seine Temperatur fallen.



Damit die Turbinen noch stärker angetrieben werden, lässt man den abgekühlten Dampf im Kondensator an kalten Wasserrohren kondensieren. Beim Übergang vom gasförmigen in den flüssigen Zustand nimmt das Volumen des Wassers nämlich drastisch ab, wodurch

im Kondensator ein viel niedrigerer Druck entsteht. Mit der so vergrößerten Druckdifferenz zwischen den Enden der Turbinen nehmen auch die Antriebskräfte der Turbinen zu. Das kondensierte Wasser wird wieder in den Dampferzeuger gepumpt.

Das beim Kondensieren erwärmte Kühlwasser aus dem Kondensator wird in einem Kühlturm fein zer­sprüht niederrieseln gelassen. Dabei gibt es seine Wärme der Luft im Turm ab, die infolge der Erwärmung aufsteigt und wieder kühle Luft von unten nachströmen lässt. Das abgekühlte Wasser kann nun direkt wiederverwendet oder gegen frisches Wasser aus einem Fluss ausgetauscht werden. Man darf erwärmtes Kühlwasser aus dem Kondensator nicht ungekühlt in einen Fluss leiten, weil das Flusswasser bei Erwärmung weniger Sauerstoff bindet und folglich seinen Bewohnern schlechtere Lebensbedingungen bietet.

Die Turbine treibt einen Generator an, der die gewünschte elektrische Energie erzeugt. Über Transformatoren hochgespannt wird diese in ein Hochspannungsnetz eingespeist.

Die von einem Kraftwerk abgegebene elektrische Energie stammt aus der **chemischen Energie** des verwendeten Brennstoffes. Zur Natur chemischer Energie folgende Blitz-Visite in Biologie und Chemie:

Die Kohlenstoffverbindungen in organischem Material werden vor allem aus Wasser (H_2O) und Kohlendioxid (CO_2) aufgebaut. Dazu müssen H_2O - und CO_2 -Moleküle aufgespalten werden, wozu Energie erforderlich ist. Pflanzen beziehen diese benötigte Energie aus dem Sonnenlicht. Beim Verbinden der aufgespaltenen Bestandteile zu den brennbaren Verbindun-

gen wird zwar etwas Energie frei, aber netto muss in den gesamten Prozess eine gewisse Energie investiert werden. Ebendiese Energie wird andererseits wieder in Form von Wärme frei, wenn die Kohlenstoffverbindung zu H_2O und CO_2 verbrannt wird.

Wie viel Energie das Verbrennen eines Stoffes liefert, hängt damit von seinem chemischen Aufbau ab, und natürlich von der verbrannten Menge. Die folgende Tabelle listet für gängige Brennstoffe deren **Heizwert** auf, worunter man selbstverständlich die bei der Verbrennung freigesetzte Energie pro verbrannter Masse versteht. Die Dichte des Stoffes ist mit angegeben, weil man ja z.B. beim Kauf von Heizöl nach Volumen (z.B. in Litern) abrechnet und nicht nach Masse.

Leider setzt ein Kraftwerk nicht die gesamte zugeführte chemische Energie in nutzbare elektrische um, sondern nur etwa 40%. Diesen Anteil der tatsächlich genutzten Energie an der gesamten aufgewandten Energie nennt man – wie Du sicher noch aus Klasse 8 weißt – den **Wirkungsgrad** der Energieumwandlung, oder konkreter: des Kraftwerkes.

Die nicht genutzte Energie geht vor allem als Wärme des Kondenswassers verloren, daneben auch als Wärme der Abgase, als Rest-Energie in unverbranntem Brennstoff und als Reibungswärme an allen mechanischen Bauteilen. Nicht zuletzt benötigt ein Kraftwerk auch ein paar Prozent seiner Leistung zum eigenen Bedarf.

feste Stoffe		Schütt-	Flüssigkeiten			Gase		
Stoff	HzW. kJ/g	Dichte g/cm ³	Stoff	HzW. kJ/g	Dichte g/cm ³	Stoff	HzW. kJ/g	Dichte g/dm ³
trock. Holz	15	0,5	Spiritus	26	0,80	Erdgas	44	0,80
Braunkohle	20	0,75	Heizöl	42	0,85	Propangas	46	2,00
Koks	29	0,50	Benzin	45	0,80	Wasserstoff	120	0,09
Steinkohle	30	0,80						

27.1.2 Grobeinteilung von Energiequellen

Fossile Energiequellen

sind Kohle, Erdöl und Erdgas. Man verbrennt sie einfach und nutzt die dabei freiwerdende Wärme. Kohle entstand hauptsächlich durch Versteinierung urzeitlicher Wälder, deren Holz von Schlamm bedeckt in den Erdboden versunken ist, wo sich aus dem Holz infolge des hohen Druckes, der Temperatur und des Sauerstoffmangels Braun- und Steinkohle bilden konnte. Darin konserviert finden wir gelegentlich die Formen der damaligen Lebewesen als „Fossilien“. Das Wort „fossil“ geht auf das lateinische Verb *fossilis*=ausgegraben zurück.

Erdöl und Erdgas haben sich v.a. aus urzeitlichem Plankton gebildet. Plankton nennt man die Kleinstlebewesen an der Meeresoberfläche. Abgestorben setzen sich seine Reste am Meeresboden ab und werden durch

das Gewicht der nachfolgenden Sedimente (lateinisch *sedere*=sitzen, hier: absetzen) und des Meerwassers unter Sauerstoffmangel verdichtet, woraufhin durch verschiedene chemische Prozesse Erdöl und -gas entstehen können.

Die Bildung von fossilen Brennstoffen erfordert erdgeschichtlich lange Zeiten, so dass wir praktisch nur einen endlichen Vorrat an fossilen Brennstoffen zur Verfügung haben. Sind die Vorräte erschöpft, – was je nach Schätzung bereits in einigen Generationen der Fall sein kann, – ist Schluss damit. Insbesondere Erdöl fiele dann nicht nur als Energieträger, sondern auch als Grundstoff der petrochemischen Industrie weg, die uns mit Kunststoffen, Düngemitteln, Farbstoffen, Medikamenten, Waschmitteln und anderem versorgt.

Beim Verbrennen fossiler Brennstoffe wird einerseits Sauerstoff verbraucht, den wir zum Atmen benötigen, andererseits entsteht vor allem Kohlendioxid, das über den Treibhauseffekt unser Klima drastisch verändern kann. Es behindert die Abstrahlung von Wärme von der Erde in den Weltraum, so dass die Erde sich mit dem eingestrahnten Sonnenlicht im Durchschnitt erwärmt. (Ähnlich wie das Glas eines Treibhauses das Sonnenlicht von außen herein-, die Wärmestrahlung aus dem Inneren jedoch nicht herauslässt.)

Kernenergie

kann aus manchen chemischen Elementen unter erheblichen Risiken für Mensch und Natur gewonnen werden. Wir kommen darauf im nächsten Schuljahr ausführlicher zurück. Im Prinzip ist auch hier nur ein begrenzter Vorrat verfügbar, denn die heute nutzbaren Elemente (v.a. Uran) entstehen in großer Menge nur bei der Explosion großer Sterne. Allerdings ist ein Ende der irdischen Vorräte erst in weiter Ferne.

Grob beschrieben sind Kernkraftwerke Dampfkraftwerke, die die Energie zur Erhitzung des Wassers aus

der Spaltung von Atomkernen gewinnen. Dabei wird kein Sauerstoff verbraucht, und es entsteht kein klimafährdendes Kohlendioxid. Allerdings entstehen schon beim regulären Betrieb reichlich radioaktive Stoffe, die zum Teil Jahrtausende lang schwer abschirmbare und lebensgefährliche (weil u.a. krebsauslösende) Strahlung aussenden.

Regenerative Energiequellen

werden ständig erneuert. Es sind z.B. die mechanische Energie des Windes und des Wassers, die chemische Energie in nachwachsenden Organismen, die Wärmestrahlung der Sonne und die Wärme des Erdkörpers. Als Physiker wissen wir natürlich, dass auch diese Energien nicht aus dem Nichts entstehen: Letzten Endes stammt alle Energie, die wir nutzen können, aus unserer (oder einer anderen) Sonne. Insofern sind die regenerativen Energien dann doch erschöpft, wenn die Sonne in geschätzten 6 Milliarden Jahren aus geht. Vermutlich wird jedoch schon vorher niemand mehr auf der Erde sein, um dies zu bemerken.

27.1.3 Wasserkraftwerke

...wandeln mechanische Energie in elektrische Energie um. Das Grundprinzip ist dabei immer das gleiche: Wasser in erhöhter Lage strömt nach unten und treibt über **Turbinen** (=Schaufelräder) einen **Generator** an, der elektrische Spannung und Strom erzeugt.

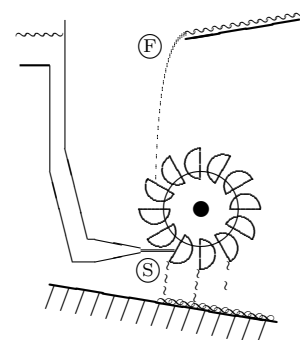
Modellhaft könnte am Fuße eines hohen Wasserfalls ⑤ ein Mühlrad vom schnell fallenden Wasser angetrieben werden. Typischerweise lässt man aber das erhöht gelegene Wasser nicht einfach fallen, sondern nutzt den Schweredruck, den hoch stehendes Wasser unter sich erzeugt. In vielen Kraftwerken mit großem Höhenunterschied lässt man das unter hohem Schweredruck stehende Wasser zunächst aus engen Düsen als schnellen Strahl ⑥ austreten und dann gegen die Schaufeln einer (Pelton-) Turbine spritzen. In anderen Kraftwerken lässt man das Wasser mit seinem Schweredruck direkt auf die (z.B. Kaplan- oder Francis-) Turbinen einwirken.

So wird aber letzten Endes immer Lageenergie des Wassers – evtl. mit Umweg über Bewegungsenergie – in elektrische Energie umgewandelt. Je nach Bauart des Kraftwerkes und Betriebsbedingungen erreicht man bei dieser Umwandlung Wirkungsgrade um 85%. Man unterscheidet nach der Art des genutzten Gewässers vor allem:

Flusskraftwerke an gestauten Flüssen: Das Wasser verliert an einer Staustufe zwar nur wenige Meter an Höhe, aber es fließt dauernd und viel Wasser pro Zeit durch die Staustufe.

Speicherkraftwerke hinter hochgelegenen (Stau-)Seen, v.a. im Gebirge: Im Gebirge führen die Flüsse zwar noch nicht so viel Wasser, dazu noch jahreszeitlich recht stark schwankend, jedoch sind die nutzbaren Gefälle größer. In den Seen lässt sich zudem in Form des hochgelegenen Wassers Energie speichern, bis sie benötigt wird.

Gezeitenkraftwerke nutzen die Höhenunterschiede des Meeresspiegels infolge der Gezeiten aus. Dazu kann man bei Flut das steigende Meerwasser durch turbinenbestückte Rohre in einen landwärts gelegenen Wasserspeicher strömen lassen und bei Ebbe zurück. Als Wasserspeicher kommen breite Flussmündungen in Frage.



27.2 Der Energieverbrauch eines Landes

...lässt sich nach vielen Kriterien aufgeschlüsselt darstellen, zum Beispiel wie im folgenden Diagramm, grob angelehnt an den Verbrauch Deutschlands im Jahr 2021. Breiten sollen den Energiebeträgen entsprechen. Von oben nach unten sind so die Energien dargestellt, wie sie in die Wirtschaft hineinkommen (**Primärenergieträger**), wie sie den Verbrauchern zur Verfügung gestellt werden (**Sekundär- oder Endenergieträger**, Vorsilbe „End“ aus Sicht der Energiewirtschaft), und wofür sie tatsächlich genutzt werden.

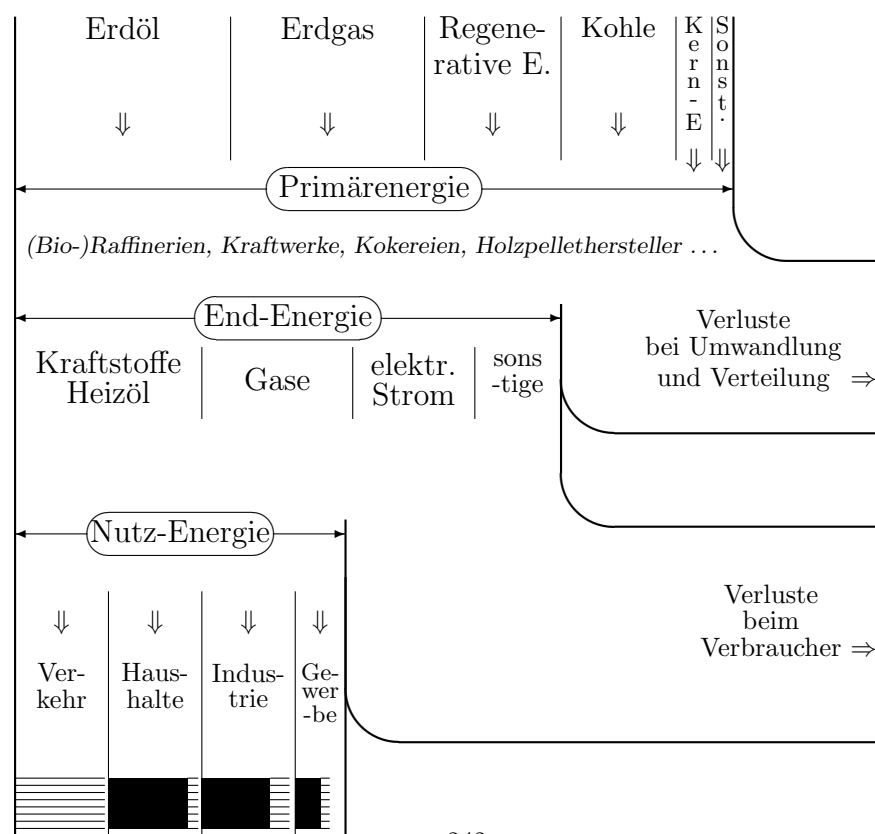
Die Zusammensetzung der Primärenergie variiert deutlich von Land zu Land und im Laufe der Zeit. 2021: Frankreich nutzt Kernenergie stark und Kohle kaum, Polen umgekehrt. Island kommt sehr weitgehend mit regenerativer Energie aus.

Bei den fossilen Energieträgern wie Erdöl entspricht die Breite im Diagramm dem Brennwert des gesamten verbrauchten Materials. Große Teile dieses Brennwertes gehen jedoch z.B. bei der Umwandlung in elektrische Energie und Fernwärme je nach Wirkungsgrad des Kraftwerkes ungenutzt verloren. Bei regenerativen Energien rechnet man hier ohne die physikalisch eintretenden Verluste bei der Gewinnung. Öl&Co kosten nämlich Geld, aber sie behalten ihren Brennwert, wenn man sie nicht verbrennt. Sonnenschein&Co hingegen kosten an sich kein Geld, und sie verlieren ihren Wert für die Energiewirtschaft sowieso, auch wenn man sie gar nicht nutzt. Zudem wäre es schwierig zu beziffern, wie viel ungenutzte Energie z.B. in einem Wind in der Umgebung eines Windrades steckt.

Die Verbraucher werden oft zusammengefasst zu

den vier Bereichen Verkehr, private Haushalte, Industrie und Gewerbe/Dienstleister. Die schwarzen Rechtecke am unteren Rand stellen die *allein* als Wärme genutzte Energie dar, die schraffierten Flächen daneben die für mechanische Arbeit genutzte Energie. Wärme umfasst Raumheizung, Warmwasserbereitung und – vor allem in der Industrie – Prozesswärme. Dies ist Wärme, die zur Durchführung eines technischen Prozesses benötigt wird, z.B. die Wärme, mit der in Hochöfen Metall geschmolzen wird. Im Haushalt fiel darunter z.B. die Wärme, die zum Kochen oder Bügeln benötigt wird.

Bei Endverbrauchern kann in der Regel nicht erfasst werden, wie viel Verlust entsteht, schon deshalb nicht, weil man „Verlust“ unterschiedlich interpretieren kann. Eine Heizungsanlage soll möglichst viel der chemischen Energie eines Brennstoffes zu Wärme des Wassers in den Heizkörpern umwandeln. Bei hohem Wirkungsgrad der Anlage gibt es nur geringe Umwandlungsverluste, vor allem als Wärme der Abgase. Ist aber das beheizte Haus schlecht isoliert, bleibt die Heizungswärme nicht lange im kalten Haus, so dass sie – gemessen am Möglichen – kaum den gewünschten Nutzen – nämlich warmes Haus – bringt, in diesem Sinne also als verloren („verschwendet“) angesehen werden kann. Umgekehrt ist die Erwärmung einer Glühlampe ein Verlust im Sinne der Umwandlung von elektrischer Energie zu Licht, kann aber zur Raumheizung beitragen. Das Diagramm ist deshalb nur unter der vereinfachenden Annahme erstellt, dass alle Verbraucher eine gleiche, geschätzte Verlustquote haben.



27.3 Wertigkeit von Energieformen

27.3.1 Nutzbarkeit

Der Energieerhaltungssatz besagt zwar, dass Energie weder erzeugt noch vernichtet werden kann, aber alle Welt spricht von *Energieverbrauch*, *-erzeugung* und *-knappheit*, von *endlichen* Ressourcen, die eines Tages *erschöpft* sein werden. Wie verträgt sich das?

Der Energieerhaltungssatz stimmt natürlich, aber er sagt nichts über die unterschiedliche *Nutzbarkeit* der Energie aus. Offensichtlich ist die chemische Energie in einem Liter Heizöl besser nutzbar und damit wirtschaftlich interessanter als dieselbe Energiemenge, die nach Verbrennen des Öls als innere Energie in der erwärmten Umgebung vorliegt.

Besser nutzbar und damit wirtschaftlich mehr wert

ist eine Energieform dann, wenn man mit ihr mehr Dinge tun kann. Mit der elektrischen Energie aus der Steckdose lässt sich zum Beispiel ein Motor betreiben, eine chemische Elektrolyse durchführen, Licht machen oder ein Zimmer heizen. Mit der inneren Energie warmer Zimmerluft lässt sich hingegen kaum mehr ein Motor betreiben, sondern bestenfalls ein anderes Zimmer heizen. Und selbst dabei gilt die Einschränkung, dass das zweite Zimmer ohne weiteren Aufwand niemals dieselbe Temperatur wie das erste Zimmer erreichen kann. Wärme wird ja bekanntlich nur vom wärmeren zum kälteren Körper übertragen.

27.3.2 Wärmepumpen

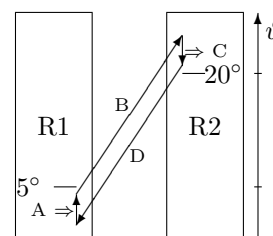
Mit folgendem Aufwand gelingt es jedoch, aus einem z.B. 5°C warmen Raum R1 Energie zu entnehmen und einem 20°C warmen Raum R2 zuzuführen: In einem Rohrsystem durchströmt ein 0°C kaltes Gas bei niedrigem Druck den 5°C kalten Raum, wobei es diesem Raum Wärme entzieht (A): maximal so viel, bis das Gas ebenso warm wie der Raum ist.

Mit einer Pumpe wird das Gas nun so stark komprimiert, dass seine Temperatur auf 25°C ansteigt (B). In diesem Zustand hohen Druckes durchströmt es den nur 20°C warmen Raum, dem es folglich Wärme abgibt (C). Das Gas selbst kühlt dabei wieder ab. Lässt man das Gas nun wieder bis auf den anfänglichen Druck expandieren (D), so kühlt es wieder auf 0°C ab und kann wieder dem kälteren Raum Wärme entziehen.

In der Energiebilanz nimmt das Gas Energie aus dem kälteren Raum auf, und die Pumpe verrichtet beim Komprimieren Arbeit an ihm. Das Gas gibt andererseits an den wärmeren Raum Energie ab und verrichtet beim Expandieren Arbeit an der Pumpe. Genaue Rechnungen ergeben, dass die Pumpe beim Komprimieren mehr Energie abgibt als sie beim Expandieren zurückerhält. In den wärmeren Raum gelangt also die Energie aus dem kälteren Raum zuzüglich der Differenz von der Pumpe.

Für praktische Zwecke benutzt man anstelle des Gases eine Substanz, die während eines Zyklus siedet und wieder kondensiert. Bei diesen Umwandlungen kann die Substanz nämlich bei selben Temperaturverhältnissen relativ große Energiebeträge aufnehmen bzw. abgeben.

Anwendungen sind vor allem Wärmepumpen mit dem Erdreich als Raum 1 und dem geheizten Wohnraum als Raum 2, und außerdem Kühlschränke mit ihrem Innenraum als Raum 1 und der restlichen Küche als Raum 2.



27.3.3 Unvermeidliche Entwertung

Mit einer Wärmepumpe ließe sich also z.B. mit mäßig warmer Luft doch wieder Luft so stark erhitzen, dass man mit dem so erzeugbaren Druck Turbinen zur Erzeugung elektrischer Energie betreiben könnte. Man würde damit aus der sonst schlecht nutzbaren Wärmeenergie wieder die vielseitig nutzbare elektrische Energie gewinnen. Der Haken an der Sache ist allerdings der Energieumsatz an der Pumpe: Die Pumpe würde ja z.B. auch mit elektrischer Energie betrieben, die umgewandelt würde in Wärme von Raum 2 oder gar in Verlustwärme.

Die Umwandlung wertarmer Energie in wertvolle gelingt nur, wenn dabei andere wertvolle Energie entwertet wird. Der Gesamtwert aller Energie wird niemals erhöht.

Die wertvollsten Energieformen sind erfahrungsgemäß elektrische und mechanische Energien. Mittlere Wertigkeit hat chemische Energie. Den niedrigsten Rang nimmt Wärme ein, wobei sie umso wertloser ist, je geringer die Temperatur ist.

Leider produzieren sehr viele Geräte vor allem Wärme, auch wenn sie gar nicht dazu gedacht sind: Glühlampen sind vom Energieumsatz her betrachtet fast reine „Heizlampen“, Automotoren nutzen nur etwa ein Drittel der chemischen Energie des Treibstoffes für den Antrieb des Fahrzeuges, ähnliches gilt für Dampfturbinen.

27.3.4 Reversibilität

Aus dem letzten Merksatz folgt, dass ein Vorgang, bei dem die Energie entwertet wird, nicht rückwärts ablaufen kann. Dabei würde ja die Energie aufgewertet, was aber nicht vorkommt.

Da bei allen Vorgängen mit Reibung aus einer mechanischen Energie Wärme entsteht, sind alle Vorgänge mit Reibung irreversibel. Gleiches gilt für alle Vorgänge mit Entwicklung Joulescher Wärme an elektrischen Widerständen. In der Praxis gibt es daher streng genommen überhaupt keine reversiblen Vorgänge.

27.4 Referate

einige Themen:

Wasserkraft, Solarenergie, Windkraft, Erdwärme, Müll, Kernspaltung, Kernfusion, Spitzenlast, dezentrale Kraftwerke, Kraftwerke in unserer Gegend, Kraft-Wärme-Kopplung, Energiereserven, „Desertec“, Elektroautos, „peak oil“, Biosprit, Klima, Energiesparen, Energieströme, ...

In diesem Kapitel gibt es viele ergänzende Themen, über die sich gut ein kurzes Referat halten lässt. Folgendes ist dabei zu beachten:

Ein Referat soll Mitschülern die Arbeit ersparen, sich anhand verschiedener Quellen selbst einen Ein- und Überblick zu einem Thema zu erarbeiten. Der Referent muss insbesondere nicht beweisen, dass er alle Einzelheiten zu einem Thema parat hat, er ist kein Prüfling. Seine Leistung besteht darin, aus seinen Unterlagen das Wesentliche herauszuarbeiten und dieses möglichst leicht verständlich den Zuhörern darzustellen. Das Referat ist geglückt, wenn der Zuhörer eine Stunde nach dem Referat das Wesentliche wiedergeben könnte.

Bedenke dazu: Der Zuhörer hat wenig Vorkenntnisse zum Thema, er hat sich noch nicht damit beschäftigt. Er durchlebt das Referat nur einmal, kann keinen Satz noch einmal nachlesen, keine Daten nachschlagen. Der Zuhörer merkt sich nicht alles auf Anhieb und vergisst auch manches. Alles, was man selbst sich nach einmaligem Hören nicht merken würde (weil es zu viel, zu langweilig, zu schwierig ist), ist vergeblich gesagt.

Wenn trotzdem das Wesentliche hängen bleiben soll, muss ebendieses klar und deutlich hervortreten, dürfen nicht viele Einzelheiten von ihm ablenken, muss es eher um Zusammenhänge gehen als um pure Auflistung von Daten. Dazu ist hilfreich, dem Zuhörer zunächst das Thema in groben Zügen und die Gliederung des Referates mitzuteilen, Zusammenhänge bildlich darzustellen, Wichtiges zu wiederholen oder dauerhaft sichtbar zu halten. Nackte Zahlen werden interessanter, wenn sie im Vergleich oder Bezug zu Bekanntem vorgestellt werden.

Es ist weiterhin zu beachten, dass jemand, der ein fertig ausformuliertes Manuskript vorliest, weniger Aufmerksamkeit hält als ein freier Redner mit Blickkontakt zum Publikum. Beim Schreiben eines Textes steht außerdem keine Gestik als Ausdrucksmittel zur Verfügung, so dass der Text durch Ersatz-Formulierungen schwerfälliger wird als im mündlichen Vortrag nötig.

27.5 Wahrscheinlichkeit von Zuständen

Um die Entwertung von Energie zu verstehen, betrachten wir einige Systeme von Teilchen einmal auf mikroskopischer Ebene. Auf mikroskopischer Ebene betrachtet ist die innere Energie eines Körpers nichts anderes als mechanische und elektrische Energie seiner Atom(-verbänd)e. Was macht nun z.B. die Bewegungsenergie thermisch bewegter Atome wertloser als die Bewegungsenergie eines als Ganzes bewegten Körpers?

27.5.1 Verteilung von Tinte in Wasser

Zum Einstieg betrachten wir ein verwandtes, einfacheres Problem, das gar nichts mit Energie zu tun hat: Ein Wasserbecken ist durch eine Trennwand in zwei Kammern – eine linke und eine rechte – geteilt. Links mischt man etwas Tinte ins Wasser und lässt das Wasser zur Ruhe kommen.

Entfernt man nun die Trennwand – sogar so vorsichtig, dass keine nennenswerten Wirbel entstehen, – so verteilt sich im Laufe der Zeit die Tinte dennoch gleichmäßig im ganzen Becken. Den Grund kennst Du eigentlich auch schon: Durch die thermische Bewegung

aller Teilchen im Becken werden die Tinten-Moleküle regellos hin- und hergeschubst und können so überall hin geraten.

Wir vereinfachen das Problem noch weiter, indem wir das Wasser ganz außer Acht lassen: Eine gewisse Anzahl Tinten-Moleküle verteilt sich in einem Raum. Jedes Molekül hält sich dabei mit gleicher Wahrscheinlichkeit in der linken wie der rechten Raumhälfte auf, keine Hälfte wird bevorzugt. Um den Überblick zu behalten, beschränken wir uns zunächst auf nur 4 gedanklich durchnummerierte Moleküle.

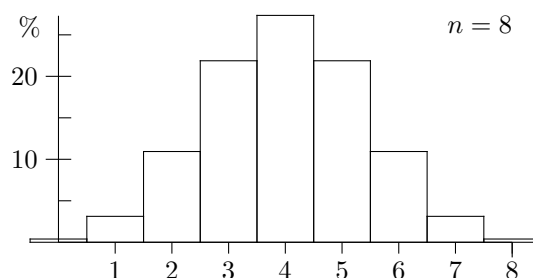
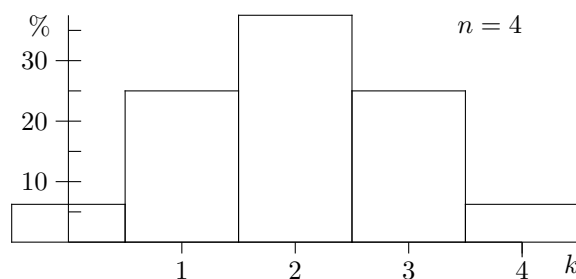
Die nebenstehende Tabelle listet alle Möglichkeiten auf, wie sich die 4 Moleküle auf die beiden Raumhälften aufteilen können. Das Kürzel RLRR z.B. soll bedeuten, dass sich nur Molekül Nr.2 (L)inks aufhält, die übrigen (R)echts.

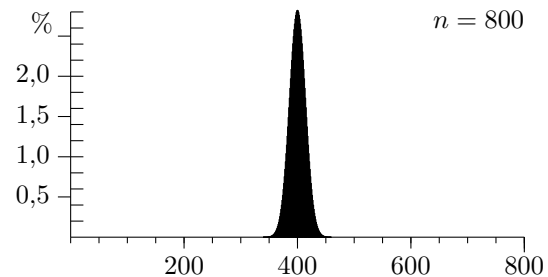
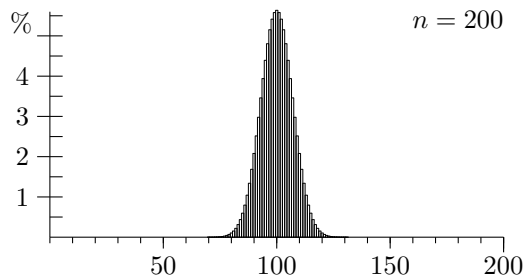
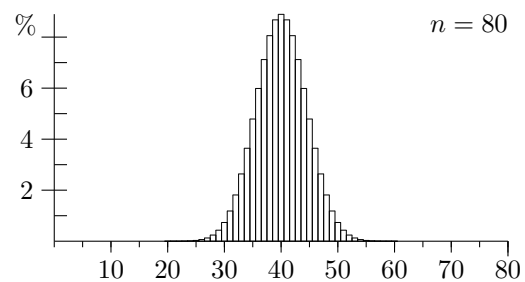
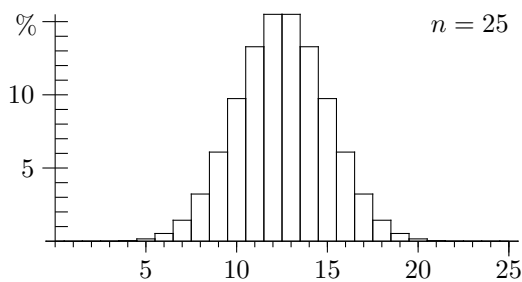
Aufteilungen	LLLL	LLLR LLRL LRLR RLLL	LLRR LRLR LRLR RLLR RLRL RLLR	LRRR RLRR RRLR RRRL	RRRR
wie viele R	0	1	2	3	4
wie viele L	4	3	2	1	0
Anzahl Auftlg	1	4	6	4	1

Da sich z.B. Molekül Nr.2 unabhängig von den übrigen Molekülen ebenso wahrscheinlich links wie rechts aufhält, ist die Aufteilung **rLrr** ebenso wahrscheinlich wie **rRrr**. Da dasselbe auch für Molekül Nr.3 gilt, tritt dieses **rrRr** auch ebenso wahrscheinlich wie **rrLr** auf. Insgesamt sind überhaupt alle 16 aufgelisteten Aufteilungen gleich wahrscheinlich, nämlich jeweils zu $\frac{100}{16}\% = 6,25\%$

Schaut man jetzt nicht darauf, *welches* Molekül wo ist, sondern nur darauf, *wie viele* Moleküle wo sind, stellt man fest: Es gibt sechs Aufteilungen, bei denen je zwei Moleküle links und rechts sind. Hingegen sind bei nur vier Aufteilungen drei Moleküle links und eines rechts. Die gleichmäßige Aufteilung 2 : 2 ist also mit $6 \cdot 6,25\% = 37,5\%$ etwas wahrscheinlicher als die ungleichmäßige Aufteilung 3 : 1 mit $4 \cdot 6,25\% = 25\%$.

Das obere Säulendiagramm zeigt diese Wahrscheinlichkeiten, mit denen die Anzahlen $k \in \{0, 1, 2, 3, 4\}$ von insgesamt $n = 4$ Molekülen links vorgefunden werden. Das untere Diagramm zeigt die entsprechenden Wahrscheinlichkeiten bei $n = 8$ Molekülen, und die Diagramme auf der nächsten Seite für weitere, immer größere n .





Der Trend ist unübersehbar: Je mehr Teilchen es werden, desto mehr konzentriert sich die Wahrscheinlichkeit um die gleichmäßige Verteilung mit je der Hälfte aller Teilchen links und rechts. Beispiel: Bei insgesamt 8 Teilchen tritt die Aufteilung 3 : 1 (entspricht bei 8 Teilchen dann 6 : 2) noch mit über 10%iger Wahrscheinlichkeit auf. Bei 800 Teilchen kommen Aufteilungen in der Gegend von 3 : 1 = 600 : 200 praktisch gar nicht mehr vor.

In einem wirklichen Experiment betragen die Teilchenzahlen leicht über 10^{20} . Nennenswerte Abweichungen von der gleichmäßigen Verteilung der Tinte im ganzen Becken werden dabei praktisch unmöglich.

Generell werden bei sehr großen Teilchenzahlen große relative Abweichungen von der wahrscheinlichsten Verteilung extrem unwahrscheinlich. Systeme aus vielen Teilchen nehmen daher immer praktisch sicher den wahrscheinlichsten Zustand ein.

27.5.2 Verteilung von Energieportionen

Kommen wir nun zur Energie: Angenommen, Energie stehe nur in gleich großen Portionen zur Verfügung. Es sollen nun 6 Energie-Portionen auf drei Teilchen verteilt werden. Auf wie viele Arten ist dies möglich?

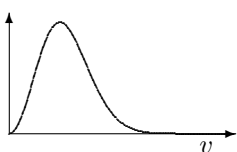
Energie-Portionen sind natürlich keine unterscheidbaren Individuen, es gibt keine 1., 2., ... Portion. Wir

betrachten daher von vornherein nur die Anzahlen der Portionen eines Teilchens ohne Unterscheidung der Portionen untereinander. Soll Teilchen Nummer 1 vier Portionen erhalten, die übrigen Teilchen je eine, so schreiben wir dafür kurz 4-1-1. Es gibt dann folgende 28 – gleich wahrscheinliche – Möglichkeiten:

6-0-0	5-1-0	4-2-0	4-1-1	3-3-0	3-2-1	2-2-2
0-6-0	5-0-1	4-0-2	1-4-1	3-0-3	3-1-2	
0-0-6	1-5-0	2-4-0	1-1-4	0-3-3	2-3-1	
	1-0-5	2-0-4			2-1-3	
	0-5-1	0-4-2			1-3-2	
	0-1-5	0-2-4			1-2-3	

Auch hier zeigen sich manche Portionierungen wahrscheinlicher als andere. Bei diesen wahrscheinlichsten handelt es sich weder um die gleichmäßigste Aufteilung 2-2-2 noch um besonders „ungerechte“ wie 6-0-0.

Verfolgt man solche Betrachtungen (weit über Schulniveau hinaus) weiter, so ergibt sich unter anderem, wie viel Prozent der Teilchen eines Gases mit ihrer Geschwindigkeit in welchem Wertebereich liegen. Das nebenstehende Diagramm zeigt sinngemäß diese „Maxwellsche Geschwindigkeitsverteilung“. Je höher die Kurve bei einer Geschwindigkeit ist, desto wahrscheinlicher tritt diese Geschwindigkeit auf.



Für das Folgende benötigen wir allerdings nur noch die Gesamtzahl aller Aufteilungsmöglichkeiten bei gegebener Anzahl von Portionen und Empfängern. Diese sind für einige Fälle in folgender Tabelle aufgelistet. Es ist übrigens kein Zufall, dass in dieser Tabelle lauter Zahlen stehen, die Du evtl. aus dem „Pascalschen Dreieck“ kennst.

Zahl der Empfänger	Portionen									
	2	3	4	5	6	7	8	9	10	
1	2	3	4	5	6	7	8	9	10	
2	3	6	10	15	21	28	36	45	55	
3	4	10	20	35	56	84	120	165	220	
4	5	15	35	70	126	210	330	495	715	
5	6	21	56	126	252	462	792	1287	2002	
6	7	28	84	210	462	924	1716	3003	5005	
7	8	36	120	330	792	1716	3432	6435	11440	
8	9	45	165	495	1287	3003	6435	12870	24310	
9	10	55	220	715	2002	5005	11440	24310	48620	
10	11	66	286	1001	3003	8008	19448	43758	92378	
11	12	78	364	1365	4368	12376	31824	75582	167960	
12	13	91	455	1820	6188	18564	50388	125970	293930	
13	14	105	560	2380	8568	27132	77520	203490	497420	
14	15	120	680	3060	11628	38760	116280	319770	817190	
15	16	136	816	3876	15504	54264	170544	490314	1307504	
16	17	153	969	4845	20349	74613	245157	735471	2042975	
17	18	171	1140	5985	26334	100947	346104	1081575	3124550	
18	19	190	1330	7315	33649	134596	480700	1562275	4686825	
19	20	210	1540	8855	42504	177100	657800	2220075	6906900	
20	21	231	1771	10626	53130	230230	888030	3108105	10015005	

27.5.3 Systeme aus Teilsystemen

Wir betrachten nun ein Teilchensystem, das aus zwei Teilsystemen A und B besteht. Es könnten zum Beispiel zwei Eisenkugeln als Teilsysteme aus Eisenatomen betrachtet werden. Die beiden Teilsysteme sollen zunächst keine Energie untereinander austauschen können. Die Energie eines Teilsystems soll auf N_A bzw. N_B mikroskopisch verschiedene Weisen auf die Einzelteilchen des Teilsystems aufgeteilt sein können.

Besteht System A z.B. aus 3 Teilchen, die insgesamt 14 Energieportionen enthalten, dann gibt es laut voriger Tabelle dafür $N_A = 120$ verschiedene Aufteilungen. System B bestehe aus 5 Teilchen mit insgesamt 10 Energieportionen und folglich $N_B = 1001$ Aufteilungsmöglichkeiten.

Wie viele mikroskopische Energie-Aufteilungen gibt es im Gesamtsystem?

?

Ganz einfach: Jede der 120 Varianten in System A lässt sich mit jeder der 1001 Varianten in System B kombinieren, so dass es im Gesamtsystem $120 \cdot 1001 = N_A \cdot N_B = 120120$ mikroskopisch unterschiedliche Energieaufteilungen gibt.

Wie wäre es, wenn ein Energiefluss zwischen den Systemen möglich wäre?

?

Es könnte z.B. eine Portion von B nach A fließen. In B gäbe es dann nur noch 715 Aufteilungen der verbleibenden 9 Portionen, in A stiege die Zahl der Aufteilungen dank der zusätzlichen 15. Portion auf 136. Im Gesamtsystem gäbe es dann $136 \cdot 715 = 97240$ Aufteilungen, und damit weniger als in der Ausgangslage. Eine Aufteilung der Gesamtenergie mit 15 : 9 Portionen in den Teilsystemen tritt folglich unwahrscheinlicher ein als die ursprüngliche Aufteilung 14 : 10.

Flösse umgekehrt eine Energieportion von A nach B, ergäben sich mit insgesamt $105 \cdot 1365 = 143325$ Aufteilungsmöglichkeiten mehr Möglichkeiten als in der Ausgangslage. Folglich findet dieser Energiefluss eher statt als nicht.

Am wahrscheinlichsten wird das Gesamtsystem sogar in den Zustand übergehen, in dem System A 8 und System B 16 Portionen enthält. Dabei erreicht nämlich $N_A \cdot N_B$ seinen maximal möglichen Wert von $45 \cdot 4845 = 218025$.

27.5.4 Entropie

Bei jedem Gesamtsystem aus Teilsystemen ist die Gesamtmasse die *Summe* der Teilmassen, die Gesamtenergie ist die *Summe* der Teilenergien, die Gesamtladung ist die *Summe* der Teilladungen, ..., aber die Gesamtzahl an Aufteilungen ist das *Produkt* der entsprechen-

den Zahlen der Teilsysteme. Die Multiplikation ist an einer solchen Stelle recht ungewohnt, wie sie auch überhaupt gegenüber der Addition recht unanschaulich ist. Abhilfe schafft ein Logarithmus, sogar mit beliebiger Basis. Vorher fragst Du Dich aber wohl:

?

Was ist ein Logarithmus?

Betrachte eine **Potenz** $P = B^E$ mit Basis $B > 0$ und Exponent $E > 0$! In Mathematik hast Du vielleicht schon gelernt, wie man diese Gleichung durch Ziehen einer **Wurzel** nach der Basis B auflösen kann: $B = \sqrt[E]{P}$. Das Wurzelsymbol ist dabei eigentlich nur als Kurzschreibweise definiert für „die nichtnegative Zahl, deren E . Potenz P ergibt“. Zum Auflösen nach dem Exponenten E definiert man „die Zahl, mit der B potenziert P ergibt“ als **Logarithmus** von P zur Basis B , mit der Schreibweise $E = \log_B(P)$. Solange man es nur mit positiven Zahlen zu tun hat, gilt also einfach:

$$B = \sqrt[E]{P} \Leftrightarrow P = B^E \Leftrightarrow E = \log_B(P)$$

Für unsere Zwecke brauchen wir nun Folgendes: Zwei positive Zahlen x und y lassen sich stets als Potenzen einer beliebigen positiven Basis b mit geeigneten Exponenten n bzw. m schreiben. Es ist also $x = b^n$ und $y = b^m$, anders ausgedrückt $n = \log_b(x)$ und $m = \log_b(y)$. Nun gilt:

$$\begin{aligned} \log_b(x \cdot y) \\ = \\ \log_b(x) + \log_b(y) \end{aligned}$$

$$\log_b(x \cdot y) = \log_b(b^n \cdot b^m) = \log_b(b^{n+m}) = n + m = \log_b(x) + \log_b(y)$$

Angewandt auf unsere Berechnung der Anzahl an mikroskopischen Realisierungsmöglichkeiten liefert dies:

$$N_{ges} = N_A \cdot N_B \quad | \log$$

$$\log(N_{ges}) = \log(N_A) + \log(N_B)$$

$$S_{ges} = S_A + S_B$$

Mit diesen Logarithmen haben wir im Wesentlichen eine neue physikalische Größe kennengelernt, nämlich die **Entropie S**. Die Einschränkung „im Wesentlichen“ rührt von zweierlei her: Erstens kümmern wir uns hier nicht um die mathematische Schwierigkeit, das Ausmaß an mikroskopischen Variationsmöglichkeiten des

Zustandes eines Systems auch dann zu erfassen, wenn man es nicht portions- und teilchenweise abzählen kann. Zweitens ignorieren wir zusammen mit der Basis einen Faktor vor dem Logarithmus, durch den die neue Größe sich eleganter in das Gefüge aller physikalischer Größen einpassen lässt.

Die Entropie S eines Systemzustandes ist im Wesentlichen ein Logarithmus aus dem Ausmaß an mikroskopisch verschiedenen Realisierungsmöglichkeiten dieses Zustandes.

Weiteres über Logarithmen

Es gelten folgende Beziehungen:

$$\log_b(x \cdot y) = \log_b(x) + \log_b(y)$$

$$\log_b\left(\frac{x}{y}\right) = \log_b(x) - \log_b(y)$$

$$\log_b(x^y) = y \cdot \log_b(x)$$

Der Logarithmus zur Basis 10 kann mit $\lg$ statt $\log_{10}$ geschrieben werden und spielt vor allem in der Chemie z.B. in Form von pH-Werten eine Rolle.

Physiker arbeiten lieber mit dem „natürlichen“ Logarithmus $\ln$, dessen Basis zwar eine irrationale Zahl ist (– die „Eulersche Zahl“ $e \approx 2,71828$ –), der aber in der Differenzial- und Integralrechnung leichter anzuwenden ist.

Zur Berechnung von Logarithmen mit dem Taschenrechner ist nützlich zu wissen:

$$\log_b(x) = \ln(x) / \ln(b) = \lg(x) / \lg(b)$$

ohne Entropie formuliert:	mit Entropie formuliert:
Ein System aus vielen Teilchen nimmt praktisch sicher den Zustand ein, für den es die meisten Realisierungsmöglichkeiten gibt.	... die Entropie maximal ist.
Von allen Energieumwandlungen finden nur diejenigen statt, bei denen der Gesamt-Wert der Energie nicht steigt.	... die Entropie nicht abnimmt.
Bei reversiblen Vorgängen wird Energie nicht entwertet.	... bleibt die Entropie gleich.
Bei irreversiblen Vorgängen wird Energie entwertet.	... nimmt die Entropie zu.
Ein System A gibt einem System B netto Energie ab, wenn dadurch insgesamt Energie entwertet wird.	... System A weniger Entropie verliert als B gewinnt.

27.5.5 Temperatur

Wärme geht bei Kontakt vom wärmeren zum kälteren Körper über. Da dieser Vorgang freiwillig und irreversibel abläuft, muss dabei am wärmeren Körper die Entropie weniger abnehmen, als sie beim kälteren zunimmt. Bei höherer Temperatur führt also dieselbe Änderung in der Energie des Körpers zu einer kleineren Änderung der Entropie. Anders ausgedrückt: Bei höherer Temperatur führt erst eine größere Energie-Änderung zur selben Entropie-Änderung.

Ist der Temperatúrausgleich erfolgt, wird (netto) keine Energie mehr zwischen den Körpern übertragen. In diesem Zustand muss die Abgabe einer Energieportion bei dem einen Körper die Entropie um denselben Betrag ändern wie die Aufnahme dieser Portion beim anderen Körper. Gleiche Temperatur bedeutet also gleiche Änderung der Entropie bei selber Änderung der Energie.

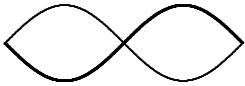
Im Einklang mit diesen Überlegungen lässt sich Temperatur verstehen als Energie-Änderung pro Entropie-Änderung: $T := \frac{\Delta E}{\Delta S}$

Dabei dürfen allerdings nur kleine Änderungen betrachtet werden, weil ja die Temperatur zu einer bestimmten Energie gehören soll und nicht zu einem ausgedehnten Energie-Bereich. Bei graphischer Darstellung des Zusammenhanges zwischen Entropie und Energie wäre die Temperatur die Steigung an einer Stelle des Graphen.

Die so definierte Temperatur deckt sich mit der absoluten Temperatur. Aus der genannten Beziehung ist nun auch die Einheit der Entropie ersichtlich: J/K. (Joule/Kelvin)

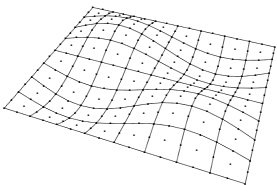
28.1 Grundlegendes über Atomhüllen

In der Quantenmechanik wird jedes Teilchen durch Wellen beschrieben. Stabile Zustände entsprechen dabei stehenden Wellen. (Das Wort „stabil“ kommt ja auch von lateinisch *stare*=stehen.) Eine primitive stehende Welle erhält Du, indem Du ein langes Seil mit einem Ende festbindest und das andere Ende mit der richtigen Frequenz schüttelst. Das Seil sollte dann zwischen den beiden hier dargestellten Extremformen hin- und herschwingen. Auch Saiten schwingen derartig.

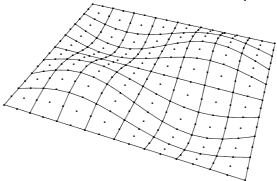


Man erkennt *Punkte*, an denen das Seil stillsteht (=Knoten, die Du noch farblich hervorheben solltest) und schwingende *Strecken*-Abschnitte dazwischen (=Bäuche). Die Stellen in der Mitte von Bäuchen schwingen mit der größten Amplitude, zu den Knoten hin nimmt die Amplitude bis auf Null ab. Die Wellen der Quantenmechanik beschreiben für jede Stelle mit ihrer dortigen Amplitude, wie wahrscheinlich das Teilchen an dieser Stelle angetroffen wird.

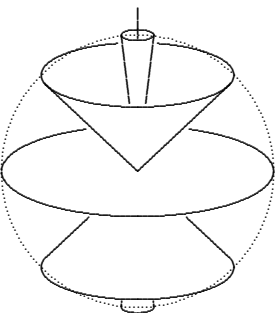
Eine Welle in der Quantenmechanik unterscheidet sich von der Seilwelle grob durch zweierlei: Erstens findet in ihr kein räumliches Hin- und Her-Schwingen von irgend-etwas statt, sondern es schwankt eine reine Rechengröße in ihrem Zahlenwert. (So wie in Schallwellen der Druck in seinem Wert schwankt.) Zweitens sind diese Wellen nicht entlang eines *eindimensionalen* Trägers wie dem Seil ausgebreitet, sondern im *dreidimensionalen* Raum.



Betrachten wir übergangsweise noch *zweidimensionale* stehende Wellen: Ein am Rand eingespanntes elastisches Netz schwinde zwischen den hier nebenstehend gezeichneten Extremlagen hin und her. (Vergleichbar mit der Membran einer Pauke) Man erkennt, dass es nun Knotenlinien gibt, zwischen denen es schwingende Teilflächen gibt. Gegenüber den eindimensionalen Wellen haben Knoten und Bäuche hier also eine Dimension mehr.



Dieser Trend setzt sich auch beim Übergang zu dreidimensionalen Wellen fort: Hier gibt es Knotenflächen und schwingende Raumbereiche dazwischen. Die Wellen für Elektronen einer Atomhülle können sich in der Anzahl, der Form und der Lage ihrer Knotenflächen unterscheiden, sowie in ihrer räumlichen Größe. An Formen haben wir Kugeloberflächen und Kegelmäntel zu betrachten, wobei die „Kegel“ als Öffnungswinkel auch mal 0° oder 90° haben können, so dass es eigentlich nur noch Achsen bzw. Ebenen sind.



Je nach Anzahl der kegeligen Knotenflächen¹ nennt man eine Schwingungsform ein s-, p-, d- oder f-**Orbital**. s-Orbitale haben gar keine kegeligen Knotenflächen und sind daher einfach kugelförmige Aufenthaltsbereiche. p-Orbitale haben eine einzige

¹Manche Flächen müssen dabei mehrfach gezählt werden: ähnlich überlegt, wie man Null eine doppelte Lösung der Gleichung $x^2 = 0$ nennen kann, weil ja das Produkt $x \cdot x$ Null ist, wenn einer der *beiden* Faktoren Null ist.

Knotenebene und bieten beiderseits dieser Ebene je einen keulenförmigen Aufenthaltsbereich für Elektronen. Da diese Ebene in drei zueinander senkrechten Raumrichtungen liegen kann, treten p-Orbitale immer im Dreierpack auf. Für die komplizierter geformten d-Orbitale ergeben sich fünf Variationen, für die f-Orbitale sogar sieben.

Orbitale können in verschiedenen Größen vorkommen. Dies führt zu der Sprechweise, eine Atomhülle sei aus **Schalen** aufgebaut. (Von innen nach außen K-, L-, M- ... Q-Schale genannt) Auf jeder Schale gibt es ein s-Orbital. Ab der L-Schale kommen drei p-Orbitale dazu, ab der M-Schale fünf d-Orbitale und ab der N-Schale sieben f-Orbitale. Ein p-Orbital der vierten Schale zum Beispiel wird als 4p-Orbital bezeichnet.

In unterschiedlichen Orbitalen haben Elektronen unterschiedlich viel Energie, die auch davon abhängt, welche anderen Orbitale des Atoms wie stark besetzt sind. In jedem Orbital können bis zu zwei Elektronen enthalten sein. Im Grundzustand eines Atoms besetzen seine Elektronen die Orbitale mit niedrigst möglicher Energie, was ja allermeist der stabilste Zustand ist.

Das chemische Verhalten eines Atoms ergibt sich weitgehend daraus, dass eine gleichmäßigere Besetzung der Orbitale den Zustand der Hülle stabiler macht. Vollbesetzte Schalen reagieren praktisch gar nicht mehr, so dass die chemischen Eigenschaften eines Elementes v.a. davon abhängen, wie die Orbitale seiner äußersten Atom-Schalen teilweise besetzt sind.

Es können sich z.B. zwei Atome mit jeweils einem erst einfach besetzten Orbital zu einem **Molekül** verbinden, indem die beiden halbbesetzten Atom-Orbitale sich zu einem vollbesetzten Molekül-Orbital um beide Kerne herum vereinigen. (= Kovalente oder Elektronenpaar-Bindung) Die Anzahl an ungepaarten Elektronen, die ein Atom enthält, oder die Anzahl an freien Plätzen in „angebrochenen“ Orbitalen legt die chemische Wertigkeit des Atoms fest.

28.2 Halbleiter

28.2.1 Leitungsmechanismen in einfachem Modell

Wie gut ein Stoff elektrisch leitfähig ist, hängt stets davon ab, in welchem Ausmaß er leicht bewegliche Ladungsträger enthält. Die Betonung muss dabei auf der leichten Beweglichkeit liegen, denn Ladungsträger überhaupt enthält ja jeder Stoff mit den Protonen seiner Atomkerne und den Elektronen der Hüllen reichlich.

Weiterhin darf unter Beweglichkeit nur eine *großräumige* Beweglichkeit verstanden werden. Eigentlich ist nämlich jedes Teilchen ständig in Bewegung. Je höher die Temperatur eines Systems ist, desto heftiger sind dabei bekanntlich die Bewegungen seiner Bestandteile im Durchschnitt. Gebundene Teilchen können sich natürlich nur innerhalb eines kleinen Bereiches bewegen und nur unwahrscheinlich diesen Bereich verlassen. (Andernfalls würde man sicher nicht von einer Bindung sprechen.)

Leiter

Die Protonen sind in aller Materie um uns herum stets fest in den Atomkernen eingebunden. Auch halten sich in allen festen Körpern die Atome gegenseitig in bestimmten räumlichen Anordnungen fest. So stehen die Protonen nur in Ausnahmen als bewegliche Ladungsträger zur Verfügung. Lediglich, wenn Ionen sich in einem flüssigen Lösungsmittel frei bewegen können, spielen Protonen für den elektrischen Strom eine größere

Rolle. Bekanntlich sind wässrige Lösungen von Salzen, Säuren und Basen elektrisch leitfähig.

Elektronen sind zwar in den Atomhüllen bei weitem nicht so stark eingebunden wie die Protonen im Kern, aber in der Regel jedenfalls doch in ihr jeweiliges Atom eingebunden. Ausnahme: Aus jedem Metall-Atom sind die ein bis drei äußersten Elektronen nicht mehr an dieses eine Atom gebunden, sondern nur noch an den gesamten Atomverband, aus dem das Metallstück besteht. Solange man ohne Quantenmechanik auskommen muss, kann man sich das so erklären, dass die äußersten Elektronen eines Atoms von dessen Kern kaum stärker angezogen werden als von den Kernen benachbarter Atome.

Man nennt ein Atom ohne diese lockeren äußersten Elektronen einen **Atomrumpf**. Atomrümpfe sind demnach stets positiv geladen. In Metallen bilden die Rümpfe ein räumliches Gitter, innerhalb dessen sich die losen Elektronen nahezu frei bewegen können. Diese Elektronen bilden ein sogenanntes **Elektronengas**, über dessen negative Ladung die positiven Atomrümpfe auch zusammenhalten. Diese Art von Zusammenhalt nennt man **metallische Bindung**. Es ist eigentlich auch diese Elektronen-Struktur, die einen Stoff erst zu einem Metall macht. Da diese sehr zahlreichen Elektronen für die gute Leitfähigkeit von Metallen verantwortlich sind, nennt man sie auch **Leitungselektronen**.

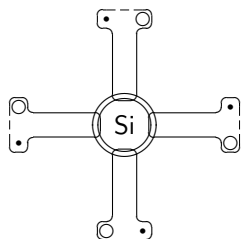
Isolatoren

Isolatoren (=Nicht-Leiter) haben keine frei beweglichen Ladungsträger: Auch die äußersten Elektronen werden jeweils von „ihrem“ Atomkern in der Atomhülle gehalten oder sind als Bindungselektronen in einem Molekülorbital eingebunden. Isolatoren sind zum Beispiel Glas, Porzellan, Plastik, Gummi und trockenes Holz.

Halbleiter

Wie der Name vermuten lässt, liegen Halbleiter mit ihrem elektronischen Innenleben zwischen Leitern und Isolatoren. Grob beschrieben sind es im Grundzustand Isolatoren, in denen es jedoch Elektronen gibt, die mit relativ wenig Energie aus ihrem gebundenen Zustand austreten und dann relativ frei im gesamten Atomverband umherschwirren können.

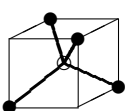
Als halbleitende Elemente sind vor allem Silizium und Germanium zu nennen. Das nebenstehende Bild zeigt schematisch die bindungsrelevanten Aspekte eines Silizium-Atoms. Es hat vier Elektronen (•) auf der dritten Schale seiner Atomhülle, wo sich das eine s - und die drei p -Orbitale zu vier gleichartigen sp^3 -Hybridorbitalen vermischt haben. So kann das Siliziumatom über vier halbbesetzte Orbitale Bindungen eingehen. Man erinnere sich: In jedem Orbital können sich zwei Elektronen aufhalten. In den Hybridorbitalen wären also noch vier freie Plätze (○) für Elektronen vorhanden.



Da Elektronen in teilweise besetzten Orbitalen die Wertigkeit eines Atoms bestimmen, nennt man solche Elektronen **Valenzelektronen**. (lateinisch: *valere*=wert sein)

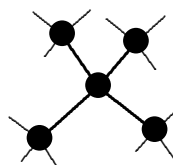
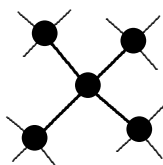
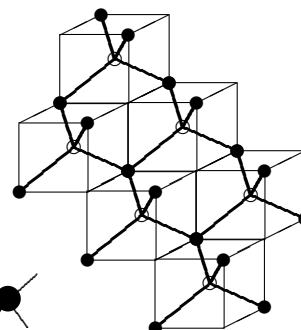
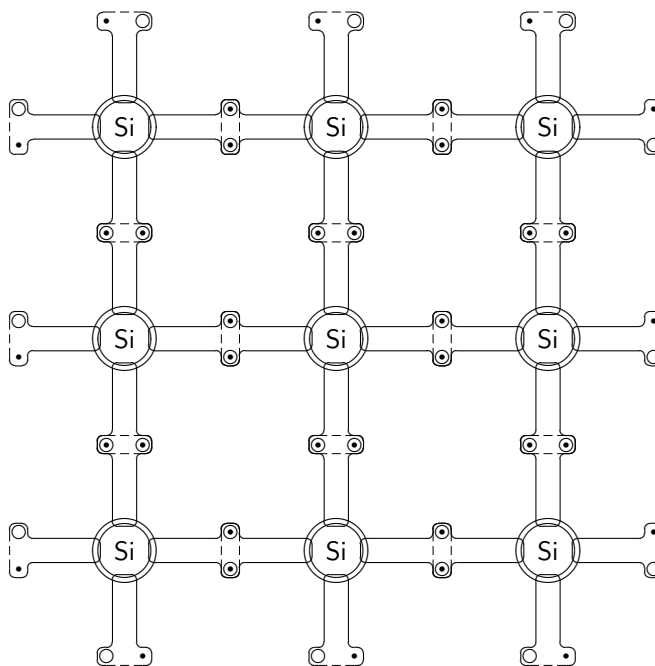
Bei der Entstehung eines Silizium-Kristalles verbinden sich die Silizium-Atome mit jeweils vier anderen Atomen. Jede dieser Bindungen kommt dadurch zustande, dass sich von jedem Bindungspartner jeweils eines der halbbesetzten Orbitale mit einem Orbital des Partners vereinigt, so dass ein vollbesetztes gemeinsames Orbital entsteht.

Was hier rechts als ebenes Gitter gezeichnet ist, ist dabei natürlich in Wirklichkeit ein räumliches Gitter. Auch diese Zeichnung ist nur schematisch zu sehen.



Die räumliche Anordnung eines Siliziumatoms und seiner vier Bindungspartner ist in Wirklichkeit so wie hier links angedeutet: Um ein Siliziumatom als Mittelpunkt eines Würfels herum sitzen die Bindungspartner so als Eckpunkte des Würfels, dass sie auch die Eckpunkte eines Tetraeders (=dreiseitige Pyramide) bilden.

Rechts ist angedeutet, wie sich ein ganzer Kristall aus solchen würfelförmigen Zellen zusammensetzt. Die folgende Zeichnung zeigt die Anordnung der Bindungspartner eines Atoms bei „magischem“ Blick noch einmal in 3D.



Elektronen- und Löcherleitung

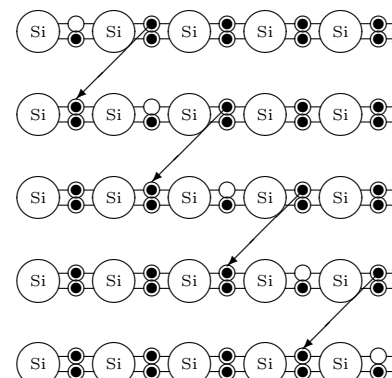
Durch relativ geringe Energiezufuhr können die bindenden Elektronen aus den Bindungen herausgelöst werden, so dass sie sich durch den ganzen Kristall bewegen können. Durch diese Erhöhung der beweglichen negativen Ladungsträgerzahl erhöht sich die Leitfähigkeit des Kristalles. Elektrische Leitung aufgrund solcher Elektronen heißt **n-Leitung**.

Mit jedem herausgelösten Elektron erhält der Kris-

tall noch einen zweiten beweglichen Ladungsträger: Die ehemaligen Bindungselektronen hinterlassen in der Orbitalstruktur eine Lücke, ein Loch. In dieses Loch kann ein Elektron aus einer benachbarten Bindung übertreten, wodurch das Loch in diese benachbarte Bindung wandert. Natürlich kann nun aus deren Nachbarschaft wieder ein anderes Elektron in das Loch übertreten, so dass das Loch weiterwandert.

Die nebenstehende Bildergeschichte zeigt diesen Vorgang anhand einer Kette von Atomen, Schritt für Schritt von oben nach unten erzählt. Siehst Du, wie das Loch nach rechts wandert, indem Elektronen nacheinander nach links springen?

Obwohl ein Loch eigentlich ein Nichts ist, erweist es sich oft als nützlich, dieses Nachrücken von Elektronen in ein entstandenes Loch nicht als Bewegung von Elektronen, sondern als entgegengesetzt gerichtete Bewegung eines Lochs zu beschreiben. Dazu muss man das Loch als einen positiven Ladungsträger ansehen, denn seine Bewegung muss ja der entgegengesetzten Bewegung von Elektronen entsprechen. Leitung durch Bewegung von Löchern heißt **p-Leitung**.



Besetzt nun wieder ein Leitungselektron ein solches Loch, so nennt man dies **Rekombination**. Elektron und Loch verschwinden dabei beide als bewegli-

che Ladungsträger. Bei der Rekombination wird genauso viel Energie frei, wie zum Herauslösen eines Elektrons aus dem Loch erforderlich ist. Diesen Energiebetrag nennt man Rekombinationsenergie.

28.2.2 Das ausgereiftere Bändermodell

In der Hülle eines einzelnen Atoms ist jeder mögliche stabile Zustand eines Elektrons eine bestimmte Schwingungsform, die als bestimmtes „Orbital“ bezeichnet wird, eine gewisse räumliche Ausdehnung hat und mit einem bestimmten Energiegehalt des Elektrons einhergeht. Dies ist im ersten der folgenden Bilder dargestellt.



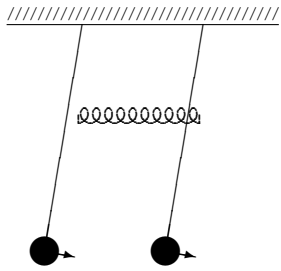
Der dicke Punkt stellt einen Atomkern dar. Die Striche darüber stellen jeweils ein Orbital dar. Die Länge eines Striches entspricht der räumlichen Ausdehnung des Orbitals. Die Höhe eines Striches stellt die zugehörige Energie dar.

Sind zwei Atome miteinander verbunden (mittleres Bild), so ergeben sich aus jeder Schwingungsform im Einzelatom zwei neue mögliche Schwingungsformen mit verschiedener Energie.

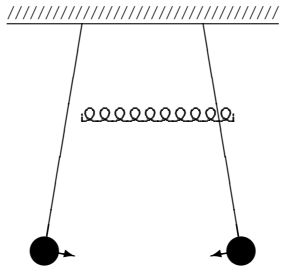
Bei N verbundenen Atomen spaltet sich jedes Energieniveau des Einzelatoms in N Niveaus auf. Bei den riesigen Atom-Anzahlen in einem Körper liegen diese Niveaus dann so dicht aneinander, dass sie praktisch ein Energieintervall vollständig abdecken. (rechtes Bild) Die energiereicheren Orbitale sind dabei auch nicht mehr an einzelnen Atomen lokalisiert, sondern

verschmolzen über den gesamten Kristall ausgedehnt: Sie bilden ein **Band**.

Bei einem Kristall aus N Atomen wird aus N gleichartigen Atom-Orbitalen für je 2 Elektronen ein einziges Band für $2N$ Elektronen. Das energetisch höchstgelegene mit Elektronen besetzte Band heißt **Valenzband**, das niedrigste unbesetzte Band darüber heißt **Leitungsband**. Im obigen Bild ist die Besetzung als Schwärzung dargestellt: Das oberste gezeichnete (Leitungs-)Band ist vollständig unbesetzt, die Bänder darunter voll besetzt.



Schon an einem gewöhnlichen Fadenpendel lässt sich ein ähnlicher Effekt erkennen: Ein einzelnes Pendel hat nur eine mögliche Schwingungsfrequenz. Koppelt man zwei ruhende Pendel über eine entspannte Feder, kann man unterschiedliche Schwingungsfrequenzen beobachten, je nach dem, ob die Pendel gleich oder entgegengesetzt schwingen. Bei gleichphasigem Schwingen bleibt die Feder dauernd entspannt und spielt folglich keine Rolle: Die Pendel schwingen ebenso, wie sie es einzeln täten. Bei gegenphasigem Schwingen wird die Feder abwechselnd gedehnt und gestaucht, übt folglich Kräfte auf die Pendel aus und verändert (genauer: erhöht) daher deren Frequenz.

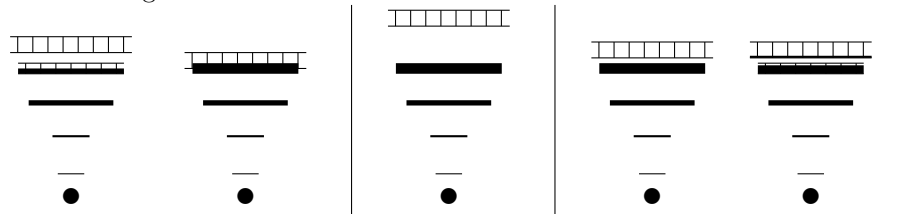


Ein vollbesetztes Band ermöglicht keinen Stromfluss. Bei quantenmechanischer Betrachtung entspricht ein Band nämlich einer Menge aller möglicher Wellenformen, wobei schon aus Symmetrie zu jeder Welle auch die genau entgegengesetzt laufende Welle enthalten ist. Bei voller Besetzung kann es daher keine überwiegende Bewegungsrichtung der Elektronen geben.

Leiter haben daher entweder ein nur halbbesetztes Valenzband (Alkalimetalle, Kupfer, Silber, Gold) oder eine Überlappung von Valenz- und Leitungsband (Erdalkalimetalle), so dass Elektronen des Valenzbandes ohne Energieaufwand in das Leitungsband übertreten können. Bei Isolatoren sind das vollbesetzte Valenzband und das Leitungsband durch eine so große Energielücke getrennt, dass es nur selten einen solchen Übertritt gibt. Bei Halbleitern ist diese Lücke relativ klein, und für die Valenzelektronen folglich leichter überwindbar.

zu ergänzen:

- Beschriftung



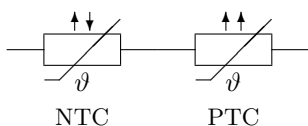
28.2.3 Anwendungen undotierter Halbleiter

Heißeleiter / NTC

Die wohl einfachste Art, durch Energiezufuhr Elektron-Loch-Paare zu erzeugen, ist eine Erwärmung des Halbleiters. Bei Zimmertemperatur beträgt das Verhältnis von freien Ladungsträgern zu Atomen typischerweise $1 : 10^9$. Bei einer Erwärmung um einige hundert Grad steigt dieses Verhältnis leicht um einige Zehnerpotenzen, erreicht allerdings bei Weitem keine Werte wie bei Metallen, wo ja eher $1 : 1$ typisch ist. Mit wachsender Temperatur wächst jedenfalls die Leitfähigkeit des Halbleiters deutlich, sein Widerstand sinkt.

Metalle verhalten sich gerade umgekehrt: Ihr Widerstand steigt mit der Temperatur. Bei Erwärmung kommt es nämlich allgemein durch die heftigeren Bewegungen der Teilchen zu verstärkter Behinderung der Bewegung von Ladungsträgern. Die gleichzeitige Vermehrung von freien Ladungsträgern macht dies bei Halbleitern aber mehr als wett, während in Metallen praktisch keine neuen freien Ladungsträger entstehen.

Ein Halbleiter, der erst in heißem Zustand gut leitet, wird naheliegend **Heißeleiter** genannt. Daneben ist auch die Bezeichnung **NTC**-Widerstand gebräuchlich: Der Zusammenhang zwischen Temperatur ϑ und Widerstand R eines Leiters kann zumindest als Näherung in kleinen Temperaturbereichen durch $R(\vartheta) = R_0 + \alpha \cdot \vartheta$ beschrieben werden. Bei einem Halbleiter muss der Temperatur-Koeffizient α dann natürlich negativ sein, weil ja zum größeren ($\uparrow$) ϑ ein kleineres ($\downarrow$) R gehört. Halbleiter haben im Englischen daher einen „negative temperature-coefficient“.

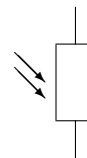


Fotowiderstand / LDR

Beleuchtung ist eine weitere Möglichkeit, Elektron-Loch-Paare zu erzeugen. In einem **Light Dependent Resistor** löst die Energie des auftreffenden Lichtes Elektronen aus ihren Bindungen. Eine Anwendung hierfür sind natürlich Messgeräte und Sensoren

für Helligkeit, ganz analog zur Verwendung von Heißeilitern zur Temperaturmessung.

Neben der Helligkeit des Lichtes spielt auch dessen Farbe eine Rolle. Licht kann nämlich als ein Strom von besonderen Teilchen, den Photonen, beschrieben werden. Die Energie eines Photons hängt dabei nur von der Farbe des Lichtes ab: Sie wächst entlang eines Spektrums in Richtung von rot nach violett. Ein Photon kann beim Auftreffen auf einen Halbleiter natürlich nur dann ein Elektron-Loch-Paar erzeugen, wenn es die dazu nötige Energie mitbringt.

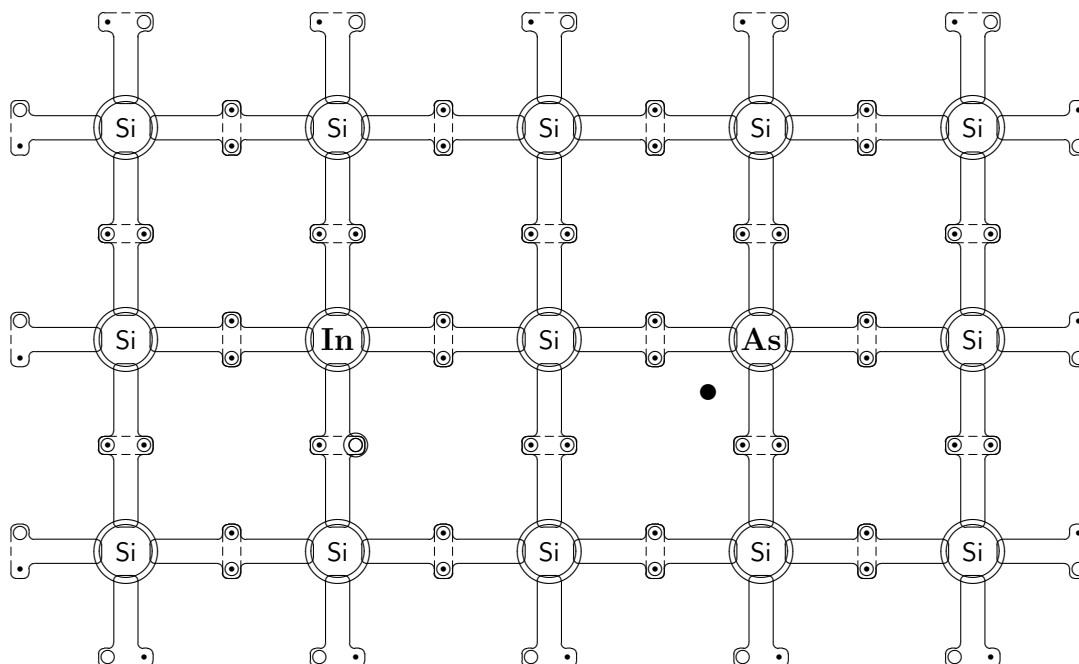


28.2.4 Leitfähigkeit durch Dotierung

Einen Halbleiter dotieren heißt, Fremdatome in sein Kristallgitter einzubauen. (1 Fremdatom auf etwa 10^6 Atome) Man unterscheidet dabei zwei Arten von Dotierungen, je nach dem, ob die Fremdatome mehr oder weniger Valenzelektronen haben als die Atome des Halbleiters.

Indium-Atome zum Beispiel sind 3-wertig, haben 3 Außenelektronen. In einem Gitter aus Silizium-Atomen bleibt daher am Indium-Atom ein Bindungsloch. In diesem Fall spricht man von einer **p-Dotierung**, weil ja mit dem Loch im Bindungsgitter ein positiver Ladungsträger für Stromfluss zur Verfügung steht. Die Fremdatome sind Elektronen-Akzeptoren.

Arsen-Atome im Silizium-Gitter haben hingegen mit 5 Außenelektronen ein Elektron übrig. Dies ist eine **n-Dotierung**, weil ja mit dem nicht eingebundenen Elektron ein negativer Ladungsträger für Stromfluss zur Verfügung steht. Die Fremdatome sind Elektronen-Donatoren.



28.3 Grenzsichten

Man dotiert Halbleiter i.d.R. nicht deshalb, um bessere Leiter aus ihnen zu machen, sondern wegen der interessanten Phänomene, die sich an Grenzflächen zwischen unterschiedlich dotierten Bereichen ergeben. (Wer einen besseren Leiter braucht, sollte Kupfer nehmen.)

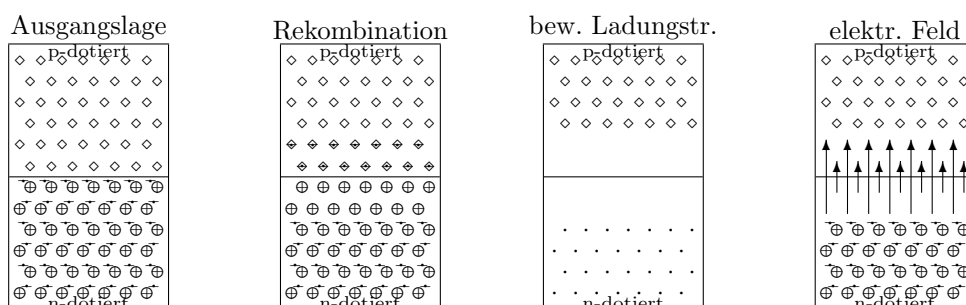
28.3.1 Diode

Aufbau (ohne Bändermodell)

Dioden bestehen aus zwei aneinandergrenzenden Halbleiterschichten, von denen die eine p-dotiert, die andere n-dotiert ist, wie es das erste der folgenden schematischen Bilder zeigt. Darin steht im n-dotierten Teil ein $\ominus$ für das nicht eingebundene Außenelektron eines Elektronen-Donators mit seiner negativen Ladung, ein $\oplus$ für den restlichen, folglich positiv geladenen Atomrumpf des Donators. Im p-dotierten Teil steht ein $\circ$ für ein Bindungsloch.

Natürlich kann diese Situation nicht so bleiben: Die

beweglichen Ladungsträger – wir betrachten vor allem die Elektronen – bewegen sich in der Umgebung ihres Ursprungsatoms zufällig im Kristall umher, auch über die pn-Grenzfläche hinweg, die ja nur eine von uns gedachte Grenze im durchgehenden Kristallgitter ist. Trifft dabei ein Elektron aus dem n-dotierten Bereich auf ein Loch des p-dotierten Bereiches, kommt es zu Rekombination: Das Elektron „bleibt im Loch hängen“. Das zweite Bild zeigt die Lage nach solchen Übertritten mit Rekombination.



Die Rekombination in der Grenzschicht hat zwei Folgen:

1. Wie das dritte Bild aller beweglichen Ladungsträger zeigt, gibt es nun in der Grenzschicht keine beweglichen Ladungsträger mehr, so dass diese Schicht nicht mehr elektrisch leitfähig ist.
2. An der Grenzfläche ist der p-dotierte Bereich durch die dorthin übergetretenen Elektronen negativ aufgeladen worden, im n-dotierten Bereich überwiegt nun die positive Ladung der Atomrümpfe. Zwischen beiden Zonen herrscht ein elektrisches Feld, das weitere Elektronen-Übertritte erschwert. Das vierte Bild zeigt dieses Feld.

Unterscheide: In einem Halbleiter entsteht durch n-Dotierung kein negativer Ladungsüberschuss: Er enthält ebenso viele Elektronen wie Protonen. Es sind lediglich einige Elektronen mehr da als für die Bindungen im Kristallgitter erforderlich. Entsprechend enthält ein p-dotierter Halbleiter keinen Mangel an Elektronen im Vergleich zur Protonenzahl, sondern einen Mangel an Valenzelektronen im Vergleich zur Anzahl benötigter Bindungselektronen. Dotierung ist keine elektrische Aufladung.

Wie üblich stellt sich auch in einer Diode ein Zustand kleinstmöglicher Energie ein: Einerseits können Elektronen aus dem n-Leiter, die durch Stöße mit anderen Teilchen per Zufall ausreichend schnell geworden sind, das elektrische Feld der Grenzschicht durchqueren und im p-Leiter rekombinieren. Das elektrische Feld wird durch diese zusätzliche Ladungstrennung stärker, der Übertritt weiterer Elektronen ist da-

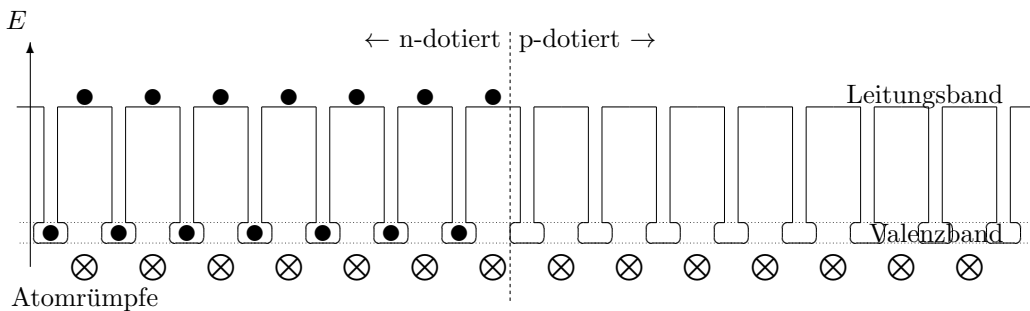
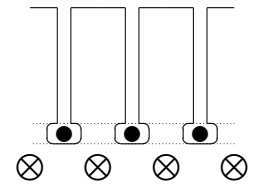
nach also erschwert.

Andererseits können die in der Grenzschicht gebundenen Elektronen ebenso zufällig genügend Energie erhalten, um ihren Bindungsplatz zu verlassen. Der hierfür erforderliche Energiebetrag hängt dabei vom Material ab und ist Dir bereits namentlich bekannt: Rekombinationsenergie. Das elektrische Feld treibt die frei gewordenen Elektronen dann in den n-Leiter und das gleichzeitig entstandene Loch in den p-Leiter. Dieser Ladungsausgleich schwächt das Feld und erleichtert damit Übertritte von Elektronen aus dem n-Leiter.

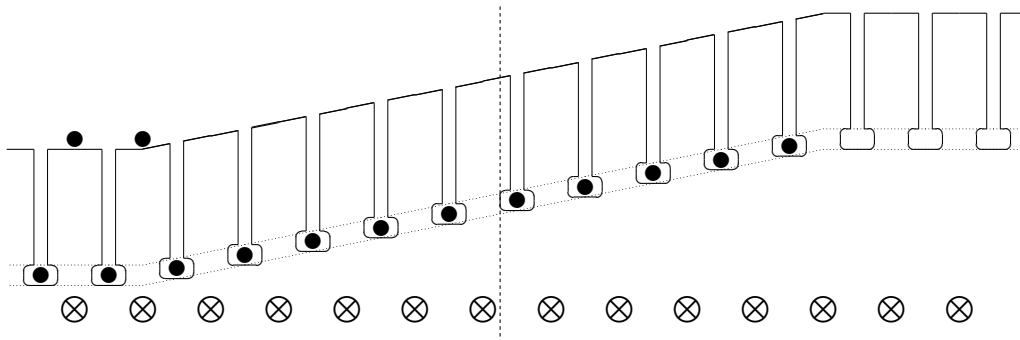
So stellt sich ein Zustand ein, bei dem die Übertritte in beiden Richtungen sich die Waage halten. Da die Wahrscheinlichkeiten beider Vorgänge von ihrem Energieumsatz abhängen, stellt sich ein elektrisches Feld ein, dessen Überwindung gerade genauso viel Energie kostet, wie zum Erzeugen eines Elektron-Loch-Paares erforderlich ist.

Aufbau (mit Bändermodell)

Das folgende Bild zeigt schematisch links einen n- und rechts einen p-dotierten Bereich. Eingezeichnet sind allerdings nur Fremdatome, auf die es ja bei dieser vereinfachten Betrachtung vor allem ankommt. Für die undotierten Abschnitte zwischen zwei Fremdatomen sähe das Bild so aus wie hier am Rand zu sehen: Die Bindungselektronen zwischen zwei Atomrümpfen füllen das Valenzband gerade voll auf, das Leitungsband ist leer. Zwischen den Bändern können sich Elektronen nicht dauerhaft befinden. („verbotene Zone“)



Bei Kontakt dieser beiden Bereiche „fallen“ die umherschweifenden Leitungselektronen des n-Bereiches in die Löcher des p-Bereiches. Dieser Ladungsübertritt lädt den p-Bereich an der Kontaktfläche negativ, den n-Bereich positiv auf. Aus Sicht eines Elektrons im n-Bereich wird durch das so entstehende elektrische Feld ein Übertritt in das p-Gebiet erschwert: das gesamte p-Gebiet ist auf ein höheres Energieniveau angehoben.



Es stellt sich die hier gezeichnete Situation ein, bei der das p-Gebiet durch das elektrische Feld ebenso hoch energetisch angehoben ist, wie die beiden Bänder auseinanderliegen. Dass ein Elektron aus dem n-Gebiet es schafft, im Leitungsband die „Steigung“ zu überwinden um dann in ein freies Loch des p-Gebietes zu fallen, ist in dieser Situation nämlich ebenso wahrscheinlich wie der Fall, dass ein Bindungselektron in der Grenzschicht es schafft, vom Valenzband ins Leitungsband aufzusteigen um darin dann in das p-Gebiet „hinabzurutschen“. (Du könntest diese beiden Elektronenbewegungen in verschiedenen Farben einzeichnen.)

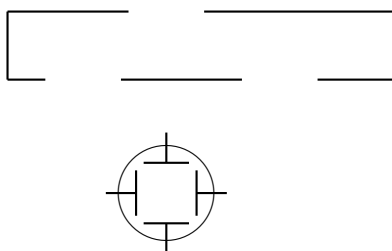
Verhalten einer Diode im Stromkreis

V 28.1 U-I-Kennlinie einer Halbleiterdiode

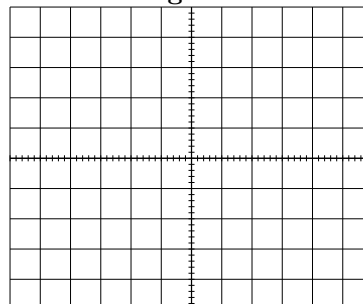
zu ergänzen:

- Wechselspannungsquelle
- Diode
- Widerstand
- Messkabel
- Oszilloskop-Bild
- vereinfachte Kennlinie

Aufbau:



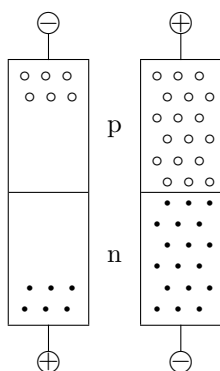
Beobachtung:



Ist der Pluspol einer Spannungsquelle am n-dotierten und der Minuspol am p-dotierten Bereich einer Diode angeschlossen, so werden im n-Bereich Elektronen und im p-Bereich Löcher von der Grenzschicht weg gezogen. Die isolierende Grenzschicht wird folglich breiter: Es fließt praktisch kein Strom. Die Diode ist in **Sperr-Richtung** geschaltet.

Ein unbedeutender **Sperrstrom** kommt dadurch zustande, dass hin und wieder Elektronen aus der Sperrschicht genügend Energie erhalten, um sich aus ihren Bindungen zu lösen. Es stehen dann ein Elektron und ein Loch als bewegliche Ladungsträger für einen Stromfluss zur Verfügung. Die Stärke des Sperrstromes hängt also davon ab, wie viele Elektron-Loch-Paare pro Zeit in der Sperrschicht gebildet werden.

Bei umgekehrter Polung mit Pluspol am p-Bereich und Minuspol am n-Bereich werden Elektronen des n-Bereiches und Löcher des p-Bereiches aufeinander zugetrieben. Die isolierende Grenzschicht wird dünner. Bei ausreichend hoher Spannung verschwindet sie ganz und die Diode leitet den Strom. Die von außen angelegte Spannung reicht dann aus, um Ladungsträger gegen das auf Seite 256 im vierten Bild eingezeichnete Feld im Inneren der Diode über den pn-Übergang zu treiben. Die Diode ist in **Durchlass-Richtung** geschaltet.



Im Bändermodell erklärt sich dieses Verhalten wie folgt: Legt man an die Diode eine Spannung in Durchlasspolung an, so erleichtert das damit einhergehende elektrische Feld die Elektronenübertritte vom n- ins p-Gebiet. Das p-Gebiet wird also energetisch erniedrigt. Wird es soweit erniedrigt, dass die Bänder wieder waagerecht verlaufen, kostet ein Elektronenübergang gar keine Energie mehr und die Diode leitet.

Bei entgegengesetzter Polung werden die Elektro-

den zwar vom (– dann noch mehr als sowieso schon –) erhöhten p-Gebiet in Richtung zum energetisch tiefergelegten n-Gebiet getrieben, haben aber auf ihrem Weg eine sehr steile Energie-Stufe zu überwinden: Auf ihrem Weg von rechts nach links im Bild auf Seite 257 gelangen sie durch das Valenzband des p-Gebietes zur Grenzschicht, wo aber das Valenzband schon gefüllt ist. Hier müssten die Elektronen ins Leitungsband hoch wechseln, wozu ihnen jedoch die Energie fehlt.



Eine Halbleiter-Diode lässt also nur in einer Richtung Strom fließen. Ein ähnliches Verhalten hast Du schon auf Seite 197 an Röhrendioden kennengelernt. Das Schaltzeichen einer Halbleiterdiode siehst Du hier am Rand. Die zusätzlichen Polbezeichnungen geben an, wie die Diode in Durchlass-Richtung geschaltet wird.

Bei Erreichen der Schwellenspannung wird eine Diode fast schlagartig sehr gut leitfähig: Die von außen angelegte Spannung ist dann gerade groß genug, um das oben beschriebene elektrische Feld in der Grenzschicht der Diode zu kompensieren. Da die Stärke dieses Feldes von der materialabhängigen Rekombinationsenergie abhängt, ist auch die Schwellenspannung materialabhängig.

Da eine Diode als einzelnes Bauteil bei höheren Spannungen dem Strom kaum mehr Widerstand entgegensetzt, würde die Stromstärke schon bei geringem Überschreiten der Schwellenspannung so groß werden, dass die Diode zerstört werden könnte. Dioden werden daher allermeist in Serie mit einem Schutzwiderstand geschaltet, der die Stromstärke begrenzt.

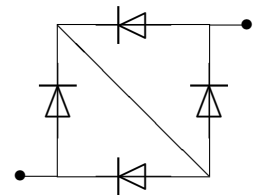
Die Schwellenspannung liegt für Siliziumdioden bei etwa 0,6 V, bei Germaniumdioden bei etwa 0,3 V. Die maximal zulässige Stromstärke liegt je nach Dientyp meist im Bereich 10...500 mA. Will man also z.B. eine Siliziumdiode mit $I_{max} = 50 \text{ mA}$ an einer Spannungsquelle von 2 V betreiben, kann der Diode nichts passieren, wenn man ihr einen Widerstand von

$R = 2 \text{ V} / 50 \text{ mA} = 40 \Omega$ in Serie dazuschaltet. Selbst wenn die Diode überhaupt keinen Widerstand hätte, würde dieser ohmsche Widerstand alleine den Strom auf 50 mA begrenzen. Berücksichtigt man, dass bei Stromfluss von den angelegten 2 V an der Diode bereits 0,6 V abfallen, kommt man mit einem Schutzwiderstand von $1,4 \text{ V} / 50 \text{ mA} = 28 \Omega$ aus.

Gleichrichter-Diode

Da eine Diode nur Stromfluss in einer Richtung erlaubt, kann man mit einer Diode aus einem Wechselstrom eine der beiden Richtungen herausfiltern. Schließt man eine Serienschaltung aus Diode und ohmschem Widerstand an eine Wechselspannung an, fließt nur in den Zeiten Strom, zu denen die angelegte Spannung mit ihrer Polung in Durchlassrichtung der Diode liegt und betragsmäßig deren Schwellenwert übersteigt.

So erhält man trotz Verwendung einer Wechselspannungsquelle einen Gleichstrom. Natürlich ist es bei dieser Art der Gleichrichtung ärgerlich, dass man (mehr als) die halbe Zeit über gar keinen Strom hat, quasi eine Hälfte der Wechselspannung verschenkt. Dies ist nur eine Ein-Phasen-Gleichrichtung, eine Einwegschaltung. Eine Zwei-Phasen-Gleichrichtung gelingt mit nebenstehender Schaltung. Egal, welche Polung die zwischen die dick gezeichneten Anschlusspunkte angelegte Spannung hat, die schräg gezeichnete Leitung wird immer in derselben Richtung von Strom durchflossen, nämlich?



Leuchtdiode / LED

Bei Stromfluss durch eine Diode rekombinieren in der Grenzschicht Elektronen mit Löchern. Die dabei frei werdende Energie kann in Form von Licht ausgesandt werden, dessen Farbe von der Rekombinationsenergie abhängt und daher auch mit der Schwellenspannung zusammenhängt.

Farbe	rot	gelb	grün	blau
Schw.-Spannung [V]	1,6...1,8	2,0...2,2	2,2...2,4	2,6...2,8

Da Dioden aus Germanium oder Silizium mit ihren relativ kleinen Schwellenspannungen nur infrarotes Licht erzeugen könnten, verwendet man als Grundmaterial für **Light Emitting Diodes** andere Stoffe, vor allem Gallium-Arsenid oder Gallium-Phosphid mit geeigneten Dotierungen.

Falls Du auf Seite 257 den Weg eines Elektrons von links nach rechts eingezeichnet hast, könntest Du nun dort, wo das Elektron in ein Loch „hinabfällt“, die Aussendung von Licht einzeichnen.



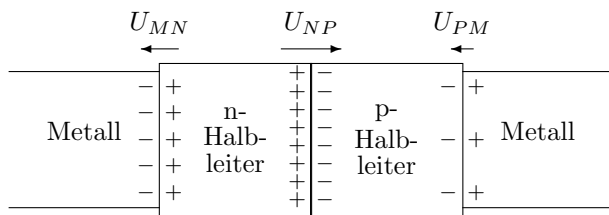
Solarzelle / Fotoelement

Solarzellen sind gewissermaßen die Umkehrung von Leuchtdioden: Bestrahlt man diese Art von Dioden mit Licht, treiben sie einen Strom an. Um die Entstehung dieses Fotostromes zu verstehen, müssen wir uns zunächst mit **Kontaktspannungen** beschäftigen.

Immer, wenn zwei verschiedene Materialien in Kontakt kommen, treten Elektronen vom einem zum anderen über. Dies liegt einfach daran, dass unterschiedliche Materialien Elektronen unterschiedlich stark binden. Für einen Kontakt zwischen p- und n-dotierten Halbleitern hast Du das vor kurzem kennengelernt. Durch den Übertritt von Elektronen entsteht über dem pn-Übergang ein elektrisches Feld und folglich auch eine Spannung zwischen den geladenen Gebieten:

die Kontaktspannung.

Die Kontaktspannung U_{NP} über der Grenzschicht einer Diode kann man jedoch nicht direkt messen oder nutzen. Zum Messen muss man ja erst einmal Messkabel an die Enden der Diode anschließen, die z.B. aus Kupfer bestehen. Beim Anschluss entstehen dann aber an den Anschluss-Stellen die Kontaktspannungen U_{MN} und U_{PM} . Das an freien Elektronen reiche Metall kann einerseits an der Kontaktstelle zum p-Halbleiter dessen Löcher stopfen, andererseits die im n-Halbleiter ungebundenen Elektronen mitaufnehmen. Bei der hohen Dichte an freien Elektronen im Metall fällt dort beides kaum ins Gewicht, verringert aber in den halbleitenden Gebieten die Energie erheblich.



Alle drei Spannungen von links nach rechts in Serie zusammen ergeben dieselbe Spannung, die auch bei direktem Kontakt zwischen dem Kupfer links und dem Kupfer rechts entstünde: Null. Selbst bei Kurzschluss der beiden Kupferkabel würde folglich kein Strom fließen. Die Spannung im Inneren der Diode kann nicht zum Antrieb eines Stromes genutzt werden.

Falls Du auf Seite 257 den Weg eines Elektrons von rechts nach links eingezeichnet hast, könntest Du dort, wo das Elektron aus dem Valenzins Leitungsband „hochspringt“, den Einfall von Licht einzeichnen.

Bei einer Solarzelle ist zumindest eine ihrer beiden Halbleiterschichten so dünn, dass Licht durch diese Schicht in die Grenzschicht fallen kann. Dort kann das Licht Elektron-Loch-Paare erzeugen. Elektron und Loch werden vom dort herrschenden Feld getrennt: das Elektron wird in den positiv geladenen n-Leiter bewegt, das Loch in den negativ geladenen p-Leiter. Diese Ladungsverschiebung (– von ihrer Richtung her ein Sperrstrom –) verringert natürlich die Aufladung der Grenzgebiete. Es sinkt dort folglich auch die elektrische Feldstärke und mit ihr U_{NP} . In der Serienschaltung der drei Kontaktspannungen überwiegt nach dieser Abschwächung von U_{NP} die Summe $U_{MN} + U_{PM}$. Die angeschlossene Solarzelle ist bei Beleuchtung eine Spannungsquelle! Ihr Pluspol ist der p-Halbleiter, Minuspol ist der n-Halbleiter.

Ihre maximale Spannung erreicht die Zelle, wenn das Licht so viele Elektronen pro Zeit freisetzt, dass U_{NP} praktisch ganz abgebaut wird. Die abgegebene Spannung ist dann vom Betrage her gleich U_{NP} , anders ausgedrückt: gleich der Schwellenspannung der Solarzelle.

Ein vollständiger Abbau von U_{NP} ist dabei eigentlich nicht möglich, da bei geschwächtem Feld Elektronen aus dem n-Gebiet leicht ins p-Gebiet übertreten und rekombinieren können, was ja das Feld wieder aufbaut. Beim Freisetzen von Elektronen aus den Löchern konkurriert das Licht also gegen die Tendenz der freien Elektronen, sich in den Löchern festzusetzen.

Fotodiode

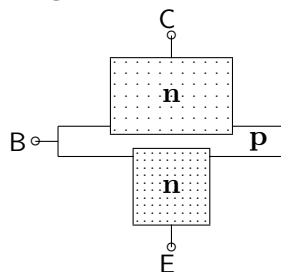
Eine Fotodiode ist im Prinzip eine Solarzelle, die jedoch nicht als Spannungsquelle benutzt wird, sondern als Schalter. Man betreibt sie in Sperr-Richtung, so dass bei Dunkelheit kein Strom fließt. Bei Helligkeit erzeugt das Licht in der Grenzschicht bewegliche Elektronen und Löcher, die – wie bei der Solarzelle beschrieben – vom elektrischen Feld in der Grenzschicht zu einem Strom in Sperr-Richtung angetrieben werden, wie ihn die äußere Spannungsquelle natürlich auch unterstützt.

Kurz beschrieben: Das Licht ermöglicht einen nennenswerten Sperrstrom, indem es in der sonst nicht leitenden Sperrschicht Elektron-Loch-Paare erzeugt.



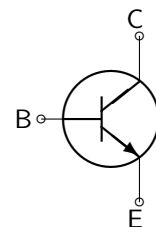
28.3.2 Transistoren

Eigenschaften eines Bipolartransistors



Ein npn-Bipolartransistor besteht aus drei Halbleiter-Schichten: einer nur wenige Mikrometer dünnen, schwach p-dotierten **Basis B**, dem stark n-dotierten **Emitter E** und dem n-dotierten **Kollektor C**.

(← Querschnitt Schaltzeichen →)



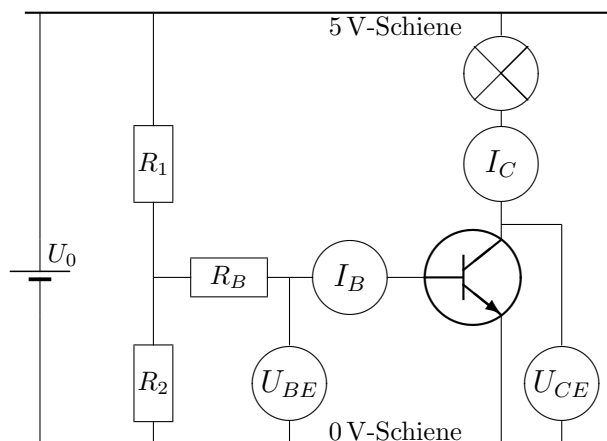
Zwischen Kollektor und Emitter liegen also zwei pn-Übergänge in einander entgegengesetzten Richtungen. Beim Anlegen einer Spannung U_{CE} zwischen Kollektor und Emitter fließt folglich kein Strom, weil ja bei jeder möglichen Polung einer der beiden Übergänge in Sperr-Richtung geschaltet ist.

Der npn-Transistor ist dazu gedacht, den Kollektor an Plus und den Emittor an Minus anzuschließen. Die Diode aus Basis und Emittor (im Bild also der untere pn-Übergang) ist dann in Durchlass-Richtung, die Diode aus Basis und Kollektor (oberer pn-Übergang) in Sperr-Richtung betrieben. Um trotzdem einen Strom zwischen Basis und Kollektor zu ermöglichen, müssten bewegliche Ladungsträger in diese Sperrschicht eingebracht werden.

Dazu legt man eine zweite Spannung U_{BE} zwischen Basis und Emittor an, und zwar mit Plus an der Basis. Der Basis-Emittor-Übergang ist also in Durchlass-Richtung geschaltet. Ist die Spannung U_{BE} größer

als die Schwellenspannung der Basis-Emittor-Diode, so treibt sie einen Strom I_B durch den Basis-Anschluss.

Für diesen Basisstrom treten also Elektronen aus dem Emittor in die Basis über. Da der Emittor wesentlich stärker dotiert ist als die Basis, können bei weitem nicht alle übergetretenen Elektronen dort rekombinieren. Da die so mit Elektronen überschwemmte Basis nur sehr dünn ist, gelangen von diesen Elektronen auch viele in die Sperrschicht zum Kollektor hin, zumal ihr Weg zum Basis-Anschluss relativ weit entlang dieser Sperrschicht verläuft. Mit diesen Elektronen wird nun auch der Basis-Kollektor-Übergang leitend.



Diese Schaltung ist Grundlage unserer weiteren Betrachtungen. Über den Spannungsteiler R_1 - R_2 wird ein Teil der Versorgungsspannung U_0 (z.B. $= 5\text{ V}$) abgegriffen und an die Basis-Emittor-Diode angelegt. (eventuell mit einem Schutzwiderstand R_B)

V 28.2 I_B - I_C -Kennlinie eines Transistors

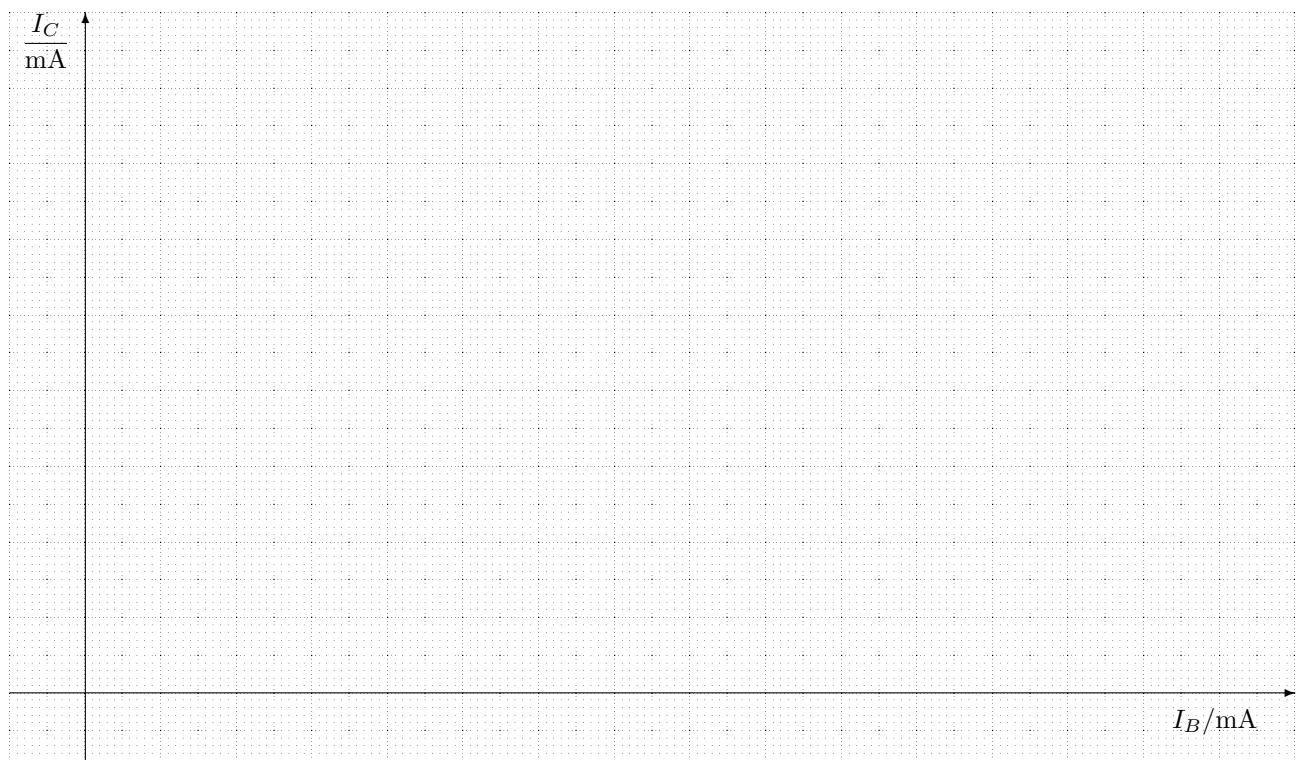
Aufbau: Als Spannungsteiler R_1 - R_2 wird ein Schiebewiderstand verwendet, dessen Mittelabgriff den Basisstrom liefert.

Durchführung: Durch Verstellen des Schiebewiderstandes werden verschiedene Basisströme I_B eingestellt und jeweils I_C gemessen.

Messwerte:

zu ergänzen:

- Messkurve
- Verstärkungsfaktor



Der Clou am Transistor ist also, dass ein schwacher Basisstrom I_B ausreicht, um einen starken Kollektorstrom I_C zu ermöglichen, der darüberhinaus auch noch weitgehend proportional zum Basisstrom I_B ist. Insbesondere werden auch die schwachen zeitlichen Schwankungen des schwachen Basisstromes in kräftige Schwankungen des stärkeren Kollektorstroms übersetzt. In anderer Formulierung ist der Transistor zwischen Kollektor- und Emitteranschluss ein Widerstand, dessen Größe sich mit dem Basisstrom variieren lässt.

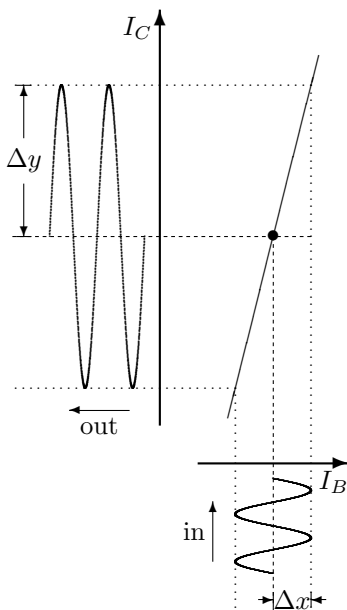
Anwendung zum Verstärken

Dies lässt sich überall dort anwenden, wo man mit einem schwachen elektrischen Signal (in Form einer geringen Spannung oder eines schwachen Stroms) größere Ströme steuern möchte. Zeichne doch im Schaltplan auf Seite 261 die Wege des steuernden Stromes (I_B) und des gesteuerten Stromes (I_C) in verschiedenen Farben ein!

Bekanntestes Beispiel einer Anwendung: In der Antenne eines Radios entstehen beim Empfang von Radiowellen schwache Ströme und Spannungen, deren Schwankungen auf einen Lautsprecher übertragen werden sollen. Natürlich ist die Antenne als Spannungsquelle viel zu schwach um den Lautsprecher direkt zu speisen. Sie ist aber stark genug, um den schwachen

Basisstrom eines Transistors schwanken zu lassen. Der Kollektorstrom weist dann dieselben Schwankungen in verstärkter Form auf, so dass mit ihm der Lautsprecher betrieben werden kann.

Nach allem bisher Besprochenen genügt es, die Antenne als Spannungsquelle direkt oder indirekt irgendwo bei R_2 , R_B oder sogar R_1 seriell oder parallel dazu zu schalten: einen Einfluss auf I_B hat man in jedem Fall. Aus allerlei anderen Erwägungen, die wir hier nicht im Detail diskutieren können, sind jedoch in verschiedenen Anwendungsfällen nur bestimmte Schaltungsvarianten geeignet. Allgemein ist für eine Verstärkerschaltung jedoch die folgende Überlegung wichtig:



Damit das verstärkte Ausgangssignal in seinem zeitlichen Verlauf dieselbe (z.B. Sinus-)Form wie das Eingangssignal hat, muss die Verstärkung in einem linearen Bereich der Kennlinie stattfinden. Also stellt man die Schaltung so ein, dass beim Mittelwert des Eingangssignals bereits („Ruhe-“)Ströme fließen, die einem Punkt (•) mitten in einem linearen Abschnitt der Kennlinie entsprechen. („Arbeitspunkt“) So ist gewährleistet, dass sich Ab- und Zunahmen des Eingangssignals gleichermaßen verstärkt auf das Ausgangssignal auswirken. (Zumindest, solange alles im linearen Bereich der Kennlinie bleibt.)

Nebenbei: Es ist oft nützlich, die Steigung $m = \frac{\Delta y}{\Delta x}$ eines Graphen als eine Art Verstärkungsfaktor zu sehen, mit dem (– zumindest winzig kleine –) Änderungen Δx einer Größe x in eine Änderung $\Delta y = m \cdot \Delta x$ einer Größe y übersetzt werden.

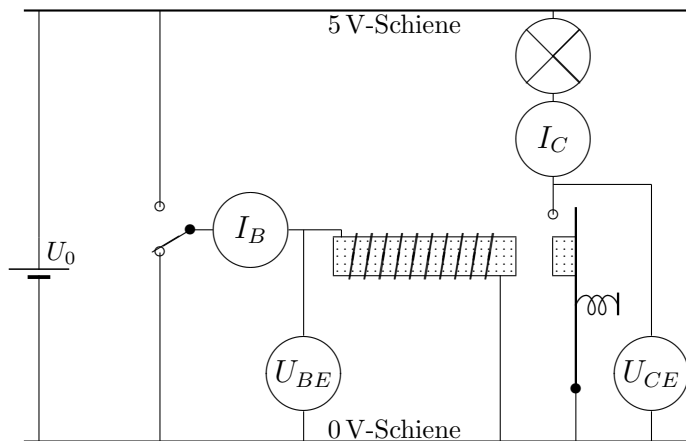
Schwache Signale, die man verstärken will, erzeugen als weitere Beispiele auch der Tonabnehmer eines Schallplattenapparates, der Lesekopf eines Bandlaufwerks, Mikrofone und andere empfindliche Messapparaturen aller Art.

Anwendung zum Steuern und Schalten

Ähnlich wie ein elektromagnetisches Relais kann ein Transistor auch als fernbedienter Schalter benutzt werden. Mit relativ kleinen Steuerströmen lassen sich dann relativ große Arbeitsströme ein- oder ausschalten. Gegenüber einem Relais bietet der Transistor jedoch einen grundsätzlichen Vorteil:

Während bei dem Relais im folgenden Bild mit dem Wechselschalter der linke Elektromagnet-Anschluss nur

auf 0 V- oder 5 V-Potential gelegt werden kann, kann bei der entsprechenden Transistor-Schaltung mit dem Spannungsteiler (Schiebewiderstand) R_1 - R_2 der Basis-Anschluss auch auf jedes Potenzial dazwischen gelegt werden. Damit kann der Widerstand des Transistors stufenlos in einem gewissen Wertebereich eingestellt werden, während der Schalter im Relais nur die Stellungen *an* und *aus* kennt.



Aufgrund der mechanischen Funktionsweise ist ein Relais natürlich auch langsamer und verschleißfreudiger als ein Transistor. Kannst Du Dir auch Vorteile von Relais vorstellen?

Feldeffekttransistor / FET

... sind eine jüngere und sehr beliebte Variante von Transistoren.

Ein „n-Kanal-Sperrschicht-Feldeffekt-Transistor“ zum Beispiel besteht aus einem dünnen n-dotierten Kanal, der von p-dotiertem Material umgeben ist. Seine drei Anschlüsse heißen *Gate* (G), *Source* (S) und *Drain* (D). Ein physischer Unterschied zwischen Drain und Source besteht nicht. Vereinbarungsgemäß wird derjenige Anschluss Drain genannt, der an den Pluspol der Spannungsquelle angeschlossen ist.

Der n-dotierte Kanal zwischen Source und Drain ist auf den ersten Blick einfach ein Stück Leiter, dessen Widerstand von seiner Länge, seiner Querschnittsfläche und dem Material abhängt. Zwischen dem n-dotierten Kanal und dem Gate bildet sich jedoch wie an jedem pn-Übergang durch Rekombination eine ladungsträgerarme Schicht aus, die den leitfähigen Querschnitt des Kanals verringert. Bei negativer Spannung von Gate gegenüber dem Kanal ist der pn-Übergang in Sperr-Richtung betrieben, die Sperrschicht daher dicker, der Kanal effektiv dünner und der Widerstand zwischen Source und Drain höher. Im Extremfall wird der Kanal ganz abgeschnürt.

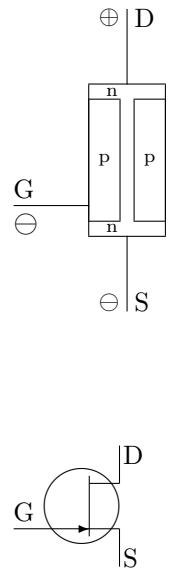
Mit dem Potenzial am Gate lässt sich daher der Strom zwischen Source und Drain steuern. Im Unterschied zu (und als Vorteil gegenüber) dem Bipolartransistor erfordert diese Steuerung nahezu gar keinen Gate-Strom, da der pn-Übergang am Gate ja in Sperr-Richtung betrieben wird, während durch eine Basis-Emitterdiode in Durchlass-Richtung ein merklicher Basisstrom fließen muss.

Für eine Polung des Gate-Kanal-Überganges in Durchlass-Richtung ist ein FET nicht konstruiert und könnte dabei zerstört werden. Das Gate muss daher mindestens ebenso stark negativ gepolt sein als der am stärksten negative Abschnitt des gesamten Kanals, also als der Source-Anschluss. Das Potenzial variiert ja innerhalb des Kanals von der negativ gepolten Source-Elektrode zur positiv gepolten Drain-Elektrode.

Deshalb wird auch die Sperrschicht zwischen Kanal und Gate mit zunehmender Stromstärke im Kanal dicker: Bei höherer Stromstärke im Kanal ist der Spannungsabfall entlang des Kanals ebenfalls größer, so dass zwischen Drain und Gate eine größere Spannung entsteht. Mit dieser Spannung wächst die Sperrschicht und schnürt bei einer bestimmten Stromstärke den Kanal ganz ab. Die Stromstärke in einem FET kann daher einen bestimmten maximalen Wert nicht überschreiten.

Neben npn-Bipolartransistoren gibt es natürlich ebenso pnp-Bipolartransistoren und neben n-Kanal-FETs auch p-Kanal-FETs.

Ein FET ist gut mit einer Triode vergleichbar: Eine Triode ist eine Röhrendiode mit einem zusätzlichen dritten Anschluss für ein Drahtgitter zwischen Katode und Anode, das bei negativer Polung gegenüber der Katode den Elektronenstrom zur Anode hemmt. Hier steuert also die Spannung zwischen Gitter und Katode die Stromstärke zwischen Kathode und Anode.



28.4 Digitalelektronik

28.4.1 Zahlen, Logik und Elektrik

Elektronische Bauteile ermöglichen bekanntlich eine sehr effektive Datenverarbeitung. Daten sind dabei alle Arten von Informationen und letzten Endes immer als Zahlen auffassbar. Auch Texte, Bilder, Klänge und Vorgänge lassen sich durch Zahlen darstellen. Man kann ja bei jeder Information, die sich aus einer endlichen Menge in Frage kommender Elemente zusammensetzen lässt, diese Elemente durchnummerieren. Mit den Nummern der Buchstaben im Alphabet als Code stellt z.B. die Zahlenfolge 16-8-25-19-9-11-0-9-19-20-0-20-15-12-12 einen kleinen Text dar. Andere Informationen lassen sich oft in Form der Zahlenwerte physikalischer Größen darstellen. Ein Ton lässt sich z.B. durch drei Zahlenangaben für Frequenz, Lautstärke und Zeitdauer beschreiben.

Um Zahlen elektronisch zu verarbeiten, muss man sie natürlich erst einmal elektronisch darstellen. Recht einfach und zuverlässig geht dies, indem man alle Zahlen im Zweiersystem ausdrückt und für die Ziffern 0 und 1 zwei gut unterscheidbare physikalische Zustände verwendet.

Zunächst zum Zweiersystem: Im Alltag verwenden wir ja das Zehnersystem (auch Dezimalsystem ge-

nannt), bei dem jede Stelle einer Zahl den 10-fachen Wert der Stelle rechts daneben hat. Im Zweiersystem (Dualsystem) unterscheiden sich benachbarte Stellenwerte um den Faktor 2.

So wie die Dezimalzahl 702,9 für

...	10^2	10^1	10^0	10^{-1}	...
0	7	0	2	9	0

also für $7 \cdot 100 + 0 \cdot 10 + 2 \cdot 1 + 9 \cdot \frac{1}{10}$ steht,

so steht die Dualzahl 1011,01 für

...	2^3	2^2	2^1	2^0	2^{-1}	2^{-2}	...
0	1	0	1	1	0	1	0

also für $1 \cdot 8 + 0 \cdot 4 + 1 \cdot 2 + 1 \cdot 1 + 0 \cdot \frac{1}{2} + 1 \cdot \frac{1}{4} = 11,25$ im Zehnersystem.

So wie im Zehnersystem 10 Ziffern (0...9) benötigt werden, braucht man im Zweiersystem nur 2 Ziffern (0 und 1), was die Sache schon recht einfach macht. (Das „kleine Einmaleins“ ist im Zweiersystem wirklich kaum mehr als $1 \cdot 1$.) Elektrisch kann man diese beiden Ziffern als Schalterstellungen „auf“ oder „zu“ interpretieren, als Potenziale +5 V oder 0 V, als Stromfluss oder nicht, als magnetischen Nord- oder Südpol, als elektrisch aufgeladen oder nicht, als Licht an oder aus, ...

Zur Verarbeitung von Dualziffern hat sich die Sichtweise bewährt, die „0“ als den logischen Wahrheitswert *falsch* und „1“ als *wahr* anzusehen. Man kennt aus der Aussagenlogik als grundlegende Verarbeitungen von Wahrheitswerten die Negation (Verneinung), die Konjunktion (und-Verknüpfung) und die Disjunktion (oder-Verknüpfung).

Für die Negation $\bar{A}$ (gelesen als „nicht-A“, „non-A“) eines Wahrheitswertes / einer Dualziffer $A \in \{0, 1\}$ gelten dabei die beiden Gleichungen

$$\bar{0} = 1 \quad \text{und} \quad \bar{1} = 0$$

$\wedge$	0	1
0	0	0
1	0	1

Die Verknüpfung $A \wedge B$ („A und B“) und die Oder-Verknüpfung $A \vee B$ („A oder B“) liefern Ergebnisse nach diesen Tabellen hier.

$\vee$	0	1
0	0	1
1	1	1

In der Elektronik werden oft die englischen Bezeichnungen **not**, **and** und **or** benutzt. Weitere Verknüpfungen heißen **nand**, **nor** und **xor**, die für „Not-AND“ (= $\overline{A \wedge B}$), „Not-OR“ (= $\overline{A \vee B}$) und „eXclusive-OR“ (= $(A \vee B) \wedge \overline{(A \wedge B)}$), „A oder B, aber nicht beide“, *entweder* A oder B) stehen. Negieren heißt in der Elektronik auch invertieren.

Bemerkenswert sind ferner die beiden De Morganschen Regeln

$$\overline{A \wedge B} = \bar{A} \vee \bar{B}$$

und

$$\overline{A \vee B} = \bar{A} \wedge \bar{B}$$

28.4.2 Grundschaltungen

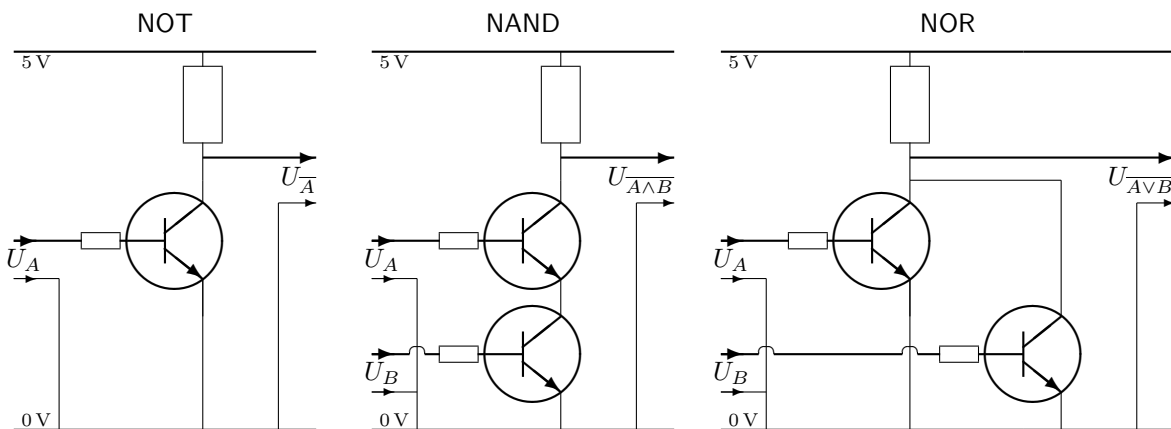
Um die Verneinung und die genannten Verknüpfungen elektronisch auszuführen, sind die folgenden Schaltungen geeignet. Logische Werte sind darin als Potenziale bezüglich der 0 V-Schiene repräsentiert: 0 V für „0“, 5 V für „1“. Man würde diese Potenziale auch als „low“ bzw. „high“ bezeichnen.

Betrachte die NOT-Schaltung: Will man A negieren, so legt man das zu A gehörige Potenzial an die dick gezeichnete und mit U_A beschriftete Leitung links an. Ist $A = 0$, so liegen dann Basis und Emitter des Transistors auf gleichem (nämlich low-) Potenzial. Es ist also $U_{BE} = 0$, folglich auch $I_B = 0$, so dass der Transistor sperrt. Es fließt demnach kein Kollektorstrom durch den Widerstand. Es fällt folglich auch keine Spannung über dem Widerstand, sondern die ganze Versorgungsspannung über dem Transistor ab. Die Leitung zwischen Widerstand und Transistor ist also auf hohem Potenzial. Logisch betrachtet kann daher an ihr „1“ abgegriffen werden, was ja mit $\bar{A}$ übereinstimmt. Die dick gezeichnete Leitung rechts ist zurecht mit $U_{\bar{A}}$ beschriftet.

Ist $A = 1$, so schaltet der Transistor durch. Die

Kollektorleitung ist dann sehr gut leitend mit der 0 V-Schiene verbunden und damit auf deren low-Potenzial. Die Versorgungsspannung fällt über dem stromdurchflossenen Widerstand in der Kollektorleitung ab. Auch in diesem Fall entspricht das Potenzial der dicken Leitung rechts dem Wert $\bar{A}$.

Grob beschrieben liegt diese Ausgangsleitung immer auf dem Potenzial derjenigen Schiene, mit der sie durch den geringeren Widerstand verbunden ist. Dies ist die „0“-Schiene, wenn der Transistor wie ein geschlossener Schalter zu dieser Schiene durchschaltet. Es ist die „1“-Schiene, wenn der Transistor wie ein offener Schalter die Verbindung zur „0“-Schiene sperrt und nur die leitende Verbindung zur „1“-Schiene über den stromlosen Widerstand bleibt.

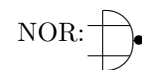
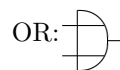
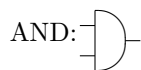


Die NAND- und die NOR-Schaltung dürftest Du Dir jetzt selbst erklären können.

Diese sogenannten RTL-Schaltungen (Resistor-Transistor-Logic) sind zwar gut verständlich, werden aber u.a. wegen ihrer langen Schaltzeiten heute kaum mehr eingesetzt. Man benutzt lieber die – leider etwas schwieriger zu durchschauende – Transistor-Transistor-Logic, kurz TTL.

Natürlich erhält man insgesamt z.B. eine AND-Schaltung, indem man den Ausgang einer NAND-Schaltung mit dem Eingang einer NOT-Schaltung verbindet.

Um die Logik komplizierterer Schaltungen besser überschauen zu können, führt man für die logischen Grundsaltungen („Gatter“) eigene Symbole ein und zeichnet nur noch die Leitungen ein, die ein logisches Signal übertragen. So spart man in der Zeichnung sämtliche Versorgungsleitungen, Transistoren und Widerstände. Weit verbreitet sind folgende halbkreisförmige Symbole für verknüpfende Gatter, an deren Ein- und Ausgängen dicke Punkte für Negationen stehen. Die Eingänge sind am geraden Rand, der Ausgang am gebogenen.



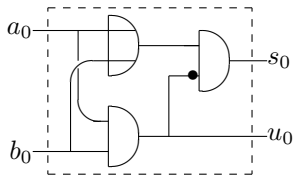
Färbe doch noch bei den NAND- und NOR-Gattern die Anschlüsse in jeweils derselben Farbe, die sich im ausführlichen Schaltplan und im Symbol einander entsprechen!

28.4.3 Komplexe Schaltungen

Eine Rechenschaltung: Addierer

Natürlich muss man natürliche Zahlen addieren können. Im Dualsystem würde man schriftlich folgendermaßen addieren:

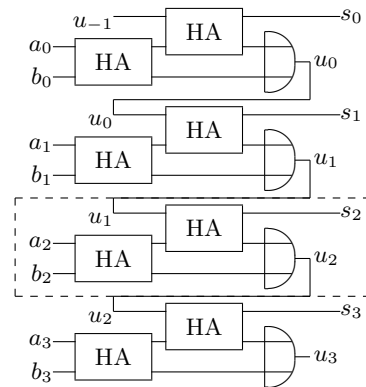
Stellen-Nr.	7	6	5	4	3	2	1	0
Summand A	0	1	1	1	1	0	1	0
Summand B	0	0	1	1	0	0	1	1
Überträge	1	1	1			1		
Ergebnis	1	0	1	0	1	1	0	1



Betrachte die Addition nur der beiden letzten Ziffern a_0 und b_0 , die die Summanden A und B an der 0. Stelle, also der Einerstelle, enthalten. Man erhält dort eine 1 als Übertrag (u_0), wenn a_0 und b_0 beide gleich 1 sind. In das Ergebnis kommt nur dann eine 1 (als s_0) hin, wenn a_0 oder b_0 gleich 1 ist, aber nicht beide. Ein Gatter, das diese Rechnung ausführt, sieht daher so aus wie am Rand zu sehen.

Ein solches Gatter könnte auch a_1 und b_1 addieren, würde jedoch den Übertrag aus der 0. Stelle nicht berücksichtigen. Es erledigt nur knapp die halbe Arbeit, die bei jeder Stelle anfällt, und heißt deshalb Halbaddierer. Wir zeichnen es kurz als **HA**.

Einen Volladdierer erhält man aus zwei Halbaddierern und einem ODER-Gatter. In nebenstehender Zeichnung ist eine Schaltung gezeichnet, die zwei vierstellige Dualzahlen addieren kann. Gestrichelt ist der Volladdierer für die 2. Stelle eingrahmt. Auch für die 0. Stelle kann ein Übertrags-Eingang vorgesehen werden, der beim Subtrahieren genutzt werden kann. Der Übertrags-Ausgang in der letzten Stufe wird als „carry-flag“ genutzt.



Dieser Addierer („carry-chain-adder“) ist nicht der schnellste, weil jede Stelle erst dann berechnet werden kann, wenn der Übertrag („carry“) aus der vorigen Stelle bekannt ist. Die Rechenzeit steigt daher linear mit der Stellenzahl. Schnellere Addierschaltungen wie der „carry-look-ahead-adder“ erfordern deutlich mehr Gatter. Ihre Rechenzeit steigt dafür aber nur mit dem Logarithmus der Stellenzahl. Die Zeiterparnis kommt daher, dass an vorderen Stellen bereits „Vorüberlegungen“ über die Auswirkungen von Überträgen angestellt werden, bevor die Überträge aus hinteren Stellen dann tatsächlich eintreffen.

Zum **Subtrahieren** benötigt man keinen eigenen Schaltkreis, denn Subtraktion ist ja bloß die Addition der Gegenzahl. Bei geschickter Darstellung von Gegenzahlen wird Subtraktion dann ganz leicht.

Ein Rechner rechnet immer nur mit einer begrenzten Anzahl an Stellen, (der „Wortbreite“ oder Bit-Zahl des Rechners,) z.B. mit 4 Stellen. Addiert man nun zu 1111 eine 0001, so ist das Ergebnis, 10000, 5-stellig. Der Übertrag aus der vierten Stelle geht in einem 4-bit-Rechner einfach verloren. (Nur am „carry-flag“ macht sich der Fehler bemerkbar.) Folglich ist 1111 diejenige Zahl, die zu 1 addiert 0 ergibt, mit anderen Worten also die -1 . Du kennst dieses Phänomen im Zehnersystem vielleicht vom Bandzähler an einem Kassettenrekorder, der beim Rückwärtsdrehen auch von 0000 auf 9999 springt.

Natürlich ist in dieser „Zweierkomplement“-Darstellung $-2 \rightarrow 1110$, $-3 \rightarrow 1101$, $\dots -8 \rightarrow 1000$. Ohne Interpretation als Zweierkomplement wäre 1000 natürlich die $+8$. Dualzahlen, die mit einer 1 beginnen, werden hier jedoch stets als negativ interpretiert, Zahlen mit führender 0 als positiv. Dies ermöglicht es nebenbei auch, dass die Stellenanzahl der Zahldarstellung ohne Änderung der darge-

stellten Zahlenwerte leicht erhöht werden kann, indem die bisher vorderste angegebene Ziffer der Zahl beliebig oft nach links wiederholt wird.

Wie man sich leicht überzeugen kann, erhält man zu einer Zahl B ihre Gegenzahl $-B$ als $\overline{B} + 1$, wobei sich $\overline{B}$ durch Negieren jeder einzelnen Ziffer ergibt. Zur Berechnung von $A - B$ mit obiger Addierschaltung müssen also nur die Eingänge b_i invertiert werden und eine zusätzliche 1 addiert werden, die man ohne großen Aufwand als u_{-1} einspeisen kann.

Eine Bestätigung dieses Verfahrens am Beispiel einer kleinen 4-bit-Rechnung:

$A - B = A + (1111 - B) - 1111 - 1 + 1 = A + \overline{B} - 10000 + 1$
 Von 1111 kann man B sicher ohne Überträge abziehen, wobei an jeder Stelle $1 - 0 = 1$ oder $1 - 1 = 0$ ergibt und somit praktisch nur B invertiert wird. Die Subtraktion von 10000 findet bei begrenzter Stellenzahl automatisch statt.

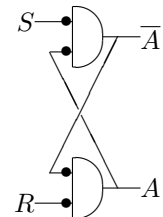
Multiplizieren lässt sich wie beim schriftlichen Multiplizieren auf das Addieren von Teilprodukten zurückführen. Für $A \cdot B$ werden ja Teilprodukte $A \cdot b_0$, $A \cdot b_1$, $\dots$ gebildet und stellenverschoben addiert. Die Teilprodukte sind im Zweiersystem offensichtlich sehr leicht zu bilden.

Dividieren und weitere Operationen wie Wurzelziehen sind schwieriger. Man kann aber z.B. für $0 < B < 2$ den Kehrwert $1/B$ erhalten, indem man in die Gleichung $y = (1 - B) \cdot x + 1$ zunächst einen beliebigen Wert aus dem angegebenen Intervall für x einsetzt und y berechnet. Den erhaltenen Wert setzt man wieder für x ein und wiederholt das ganze so oft, bis (in der Genauigkeit des Rechners) immer dieselbe Zahl herauskommt. Diese Zahl ist $1/B$, berechnet über lauter Additionen, Subtraktionen und Multiplikationen.

Die dabei erforderlichen Berechnungen mit Kommazahlen lassen sich auf die schon besprochenen Verfahren mit ganzen Zahlen zurückführen, wenn man Kommazahlen als Gleitkommazahlen darstellt. Im Zehnersystem erklärt: Man schreibt für 471,1 „wissenschaftlich“ $4,711 \cdot 10^2$ (oder unüblich, für prinzipielles Verständnis aber besser geeignet: $(+1) \cdot 4711 \cdot 10^{-1}$) und rechnet dann intern mit dem Vorzeichen +1, der natürlichen „Mantisse“ 4711 und dem ganzzahligen Exponenten 2 (oder -1).

Eine Speicherzelle: Bistabile Kippstufe / Flip-Flop

Wenn die Eingänge S und R der Schaltung auf dem Seitenrand standardmäßig auf logisch 0 liegen, so hat die Schaltung die Eigenschaft, nach einem kurzen Anlegen von logisch 1 auf Eingang S (wie: set) bzw. R (wie reset) bis auf Weiteres selbständig ihren Ausgang A dauerhaft auf 1 bzw. 0 zu halten. Der ursprüngliche Zustand spielt dabei keine Rolle. Die Schaltung „merkt sich“ einfach, auf welchem Eingang sie zuletzt eine 1 bekam. Anders formuliert: Eingang S dient zum Setzen des Ausgangs A auf 1, Eingang R dient zum Zurücksetzen von A auf 0.



Eine Steuerschaltung: Dekodierer

Um einen Computer zu bauen, braucht man noch Schaltungen, die festlegen, wann Daten aus welchen Schaltungen zu welchen anderen Schaltungen gelangen sollen. Als Beispiel soll ein Datenbit auf Leitung D in eine bestimmte von vier Speicherzellen übernommen werden. Die Speicherzellen seien RS-Flip-Flops, die beim Anlegen eines 1-Pulses auf einen ihrer Eingänge R oder S den Zustand 0 bzw. 1 annehmen. (Die Flip-Flops sind als Rechtecke gezeichnet.)

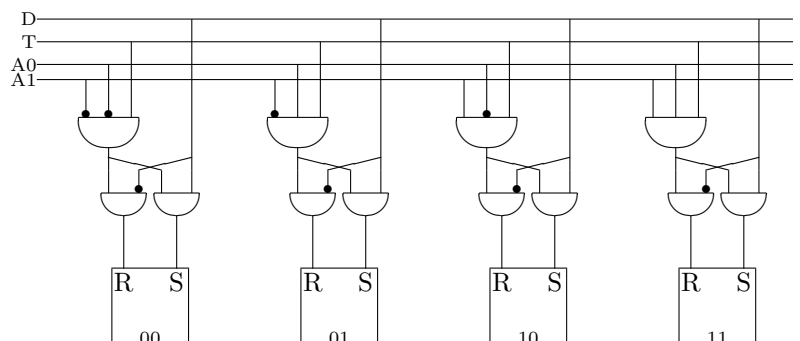
Welche der vier Speicherzellen den Wert übernimmt, wird durch die 2 Adressleitungen A0 und A1 festgelegt: Jedes Ziel habe eine zweistellige Adresse (=Nummer), die am unteren Rand jeder Zelle eingetragen ist. Die etwas größer gezeichneten Und-Gatter produzieren jeweils nur dann eine 1, wenn an ihren Eingängen eine ganz bestimmte 0/1-Kombination über-

A0 und A1 ankommt und zusätzlich die Taktleitung T eine 1 führt.

Beispiel: Bei der Kombination A0=1, A1=0 und T=1 ergibt nur das zweite große Gatter von links eine 1, weil an ihm mit A0, negiertem A1 und T lauter 1en anliegen. Die produzierte 1 gelangt zu zwei kleiner gezeichneten Und-Gattern, von denen nur dasjenige eine 1 produziert, dessen zweiter Eingang von D auch passend bedient wird. Je nach dem, ob D=1 oder D=0 ist dies nur das rechte bzw. linke, so dass nur der entsprechende Set- bzw. Reset-Eingang des Flip-Flops einen 1-Puls bekommt.

(Im folgenden Schaltplan haben nur Leitungen elektrischen Kontakt, die voneinander abzweigen. Sich kreuzende Leitungen haben keinen Kontakt.)

Zum Schreiben in eine Zelle würde man T erst einmal auf 0 lassen, damit keine unbeabsichtigten Übernahmen stattfinden: Ohne 1 auf T schaltet keines der Gatter durch. Dann würde man die Adresse der gewünschten Zelle an die Adressleitungen („auf den Adressbus“) legen und den zu schreibenden Daten-Wert auf die Datenleitung. Schließlich gibt man durch einen (ausreichend langen) 1-Puls auf der Taktleitung das Signal zur Übernahme der Daten.



271 E-Book

