



universität
wien

DIPLOMARBEIT

Titel der Diplomarbeit

„Anwendungsmöglichkeiten der (Schul-)Mathematik in der (Schul-)Physik“

Verfasser

Manfred Rusch

angestrebter akademischer Grad

Magister der Naturwissenschaften (Mag. rer. nat.)

Wien, 2015

Studienkennzahl lt. Studienblatt: A 190 412 406

Studienrichtung lt. Studienblatt: Lehramtsstudium UF Physik UF Mathematik

Betreuer: Mag. Dr. Andreas Ulovec

Es ist unmöglich, die Schönheiten der Naturgesetze angemessen zu vermitteln, wenn jemand die Mathematik nicht versteht. Ich bedaure das, aber es ist wohl so.

Richard Feynman

Inhaltsverzeichnis

1	Vorwort und Einleitung	1
2	Zur Koordination der beiden Unterrichtsfächer Mathematik und Physik	3
2.1	Was sagen die Lehrpläne zur Rolle der Mathematik im Schulunterricht?	3
2.2	Koordiniertes Unterrichten von Mathematik und Physik	8
2.3	Ziele und Grenzen dieser Diplomarbeit	13
2.4	Die Reifeprüfung vom 9. Mai 2014	17
3	Mathematische Anwendungen in der Physik nach Klassen geordnet	21
3.1	1. Klasse	23
3.1.1	Gleichungen und direktes Verhältnis	23
	Beispiel 1: Aufstellen einer Formel für die Geschwindigkeit	25
	Beispiel 2: Gewichtskraft	26
	Beispiel 3: Newton'sche Bewegungsgleichung	27
3.1.2	Maßeinheiten	28
	Beispiel 1: Dichte und elektrische Ladung in SI-Basiseinheiten	31
	Beispiel 2: Einheit der Gravitationskonstanten	32
	Beispiel 3: Falsifizieren von Gleichungen	33
	Beispiel 4: Vervollständigung von Formeln	34
3.1.3	Statistik, Mittelwerte	34
	Beispiel: Aufstellen einer Messreihe	35
3.1.4	Quader, Würfel, Raummaße	37
	Beispiel 1: Der Begriff der Dichte	37
	Beispiel 2: Auftrieb	39
	Beispiel 3: Eisberge	40
3.2	2. Klasse	43
3.2.1	Dreieck, Schwerpunkt, Flächeninhalt	43
	Beispiel 1: Flächeninhalte in Geschwindigkeits-Zeit-Diagrammen	44
	Beispiel 2: Spannungsarbeit und potenzielle Energie	45
	Beispiel 3: Zweites Kepler'sches Gesetz	47
3.3	3. Klasse	51
3.3.1	Potenzschreibweise	51
	Beispiel 1: Lichtjahr in Meter	51
	Beispiel 2: Anzahl der Atome im Universum	52

Beispiel 3: Energieerzeugung der Sonne	52
3.4 4. Klasse	55
3.4.1 Lehrsatz des Pythagoras	55
Beispiel 1: Flussüberquerung	56
Beispiel 2: Billardspiel.....	57
3.4.2 Graphen.....	59
Beispiel 1: s-t- und v-t-Diagramme in der Bewegungslehre.....	59
Beispiel 2: Bestimmung einer Regressionsgeraden	69
3.5 5. Klasse	77
3.5.1 Koordinatensysteme	77
Beispiel 1: Ebene Polarkoordinaten	78
Beispiel 2: Räumliche Polarkoordinaten (Kugelkoordinaten)	79
Beispiel 3: Zylinderkoordinaten.....	81
3.5.2 Sinus, Cosinus, Tangens.....	83
Beispiel 1: Entfernungsbestimmung von Sternen durch die Parallaxenmethode	83
Beispiel 2: Das Parsec als astronomische Entfernungsangabe	85
3.5.3 Gleichungen in zwei Variablen	87
Beispiel: Stoßprozesse.....	87
3.5.4 Vektoren in der Ebene.....	94
Beispiel 1: Vektoraddition am Beispiel des senkrechten Wurfes.....	95
Beispiel 2: Vektoraddition und Vektorzerlegung am Beispiel des schiefen Wurfes	97
Beispiel 3: Vektorzerlegung von Kräften am Beispiel des Pendels	100
3.6 6. Klasse	105
3.6.1 Exponentialfunktionen und Wachstumsprozesse.....	105
Beispiel: Radioaktiver Zerfall	107
3.6.2 Trigonometrische Funktionen	111
Beispiel 1: Schwingungen und Schwingungsgleichung	111
Beispiel 2: Interferenzen	115
Beispiel 3: Schwebungen.....	120
3.6.3 Zahlenfolgen und Reihen.....	123
Beispiel 1: Bücherüberhänge.....	123
3.6.4 Vektoren im Raum.....	129
Beispiel 1: Vektorprodukte in der Mechanik.....	131
Beispiel 2: Der Euler'sche Drehimpulssatz und die Drehimpulserhaltung.....	135
3.6.5 Statistik.....	139

Beispiel: Statistische Mechanik – Die Bernoulligleichung und die Gasgesetze.....	139
3.7 7. Klasse	149
3.7.1 Differentialrechnung	149
Beispiel 1: Fermat'sches Prinzip und Reflexionsgesetz	150
Beispiel 2: Fermat'sches Prinzip und Brechung.....	155
Beispiel 3: Definition der Geschwindigkeit und Beschleunigung	158
Beispiel 4: Einige weitere Größen in der Physik, die durch eine Ableitung definiert sind	161
Beispiel 5: Geschwindigkeit und Beschleunigung bei einer Drehbewegung	163
3.7.2 Kegelschnitte	167
Beispiel 1: Halley'scher Komet	167
Beispiel 2: Elliptische, hyperbolische und parabolische Spiegel	169
3.7.3 Komplexe Zahlen.....	178
Beispiel 1: Zeigerdiagramme	179
Beispiel 2: Zeigerdiagramme bei Widerständen, Spulen und Kondensatoren	182
Beispiel 3: Schaltkreise aus mehreren Bauteilen	189
Beispiel 4: Komplexe Widerstände beim Serienschwingkreis.....	191
3.8 8. Klasse	193
3.8.1 Integralrechnung	193
Beispiel 1: Arbeit in einem Kraftfeld	194
Beispiel 2: Dehnungsarbeit einer Feder	196
Beispiel 3: Der Massenmittelpunkt	197
Beispiel 4: Das Trägheitsmoment.....	201
Beispiel 5: Mittlere Leistung von Wechselstrom.....	204
3.8.2 Differentialgleichungen	207
Beispiel 1: Drei wichtige Differentialgleichungen der Physik.....	207
Beispiel 2: Lösen der Bewegungsgleichung für den schiefen Wurf	212
Beispiel 3: Aufstellen und Lösen der Exponentialgleichung für den radioaktiven Zerfall.....	217
3.8.3 Modelle und Modellbildung	221
Beispiel: Oberflächenform einer Flüssigkeit im rotierenden Eimer	222
Quellenverzeichnis	237
Abbildungsverzeichnis.....	241
Tabellenverzeichnis	245
Anhang A: Die Reifeprüfung vom 9. Mai 2014	
Anhang B: Abstract	
Anhang C: Lebenslauf	

1 Vorwort und Einleitung

Die beiden Wissenschaften, die Naturwissenschaft Physik und die Formalwissenschaft Mathematik sind auf dem Stand, den sie bis heute erreicht haben, nicht mehr voneinander zu trennen. Die Mathematik mag wohl auch als reine Wissenschaft denkbar sein, aber faktischerweise dient sie in vielfältigster Weise in den unterschiedlichsten Disziplinen als die entscheidende Hilfswissenschaft, die unsere heutige Technik und unser Wissen über die Natur erst möglich macht. Das heißt, dass die Mathematik entgegen der weitverbreiteten Meinung, keine Anwendungen zu haben und nur reiner Selbstzweck zu sein, ganz im Gegenteil äußerst praktisch ist und unsere heutige Gesellschaft, die zu einem sehr großen Teil auf Technik und Wissenschaft aufgebaut ist, erst ermöglicht hat. Die Mathematik ist sicherlich einer der Grundpfeiler unserer technisierten Zivilisation und es wäre wünschenswert, dass die Querverbindungen der Mathematik zu den diversen anderen naturwissenschaftlichen Disziplinen in der Schule den SchülerInnen besser vermittelt werden könnten, um der Wichtigkeit der Mathematik ihren verdienten Stellenwert zu geben.

In der Schule wird dieser praktische und anwendungsbezogene Teil der Schulmathematik oft nur in sehr eingeschränkter Weise an die SchülerInnen weitergegeben. Dabei wäre es sicher überaus motivierend für die SchülerInnen, wenn der Nutzen der Mathematik für sie im Unterricht erkennbar und erfahrbar sein könnte. Einerseits stellt die Mathematik eines der nach Meinung des Autors hervorragendsten Kulturgüter dar, die die Menschen in den vergangenen knapp 3000 Jahren hervorgebracht haben, andererseits ist die Kenntnis von mathematischen Fertigkeiten aus keiner ernstzunehmenden wissenschaftlichen Tätigkeit mehr wegzudenken.

Die Physik bietet sicherlich in ganz besonderem Maße die Möglichkeit, den SchülerInnen zu zeigen, dass Mathematik ungemein gewinnbringend angewendet werden kann. Inwieweit dies möglich ist, möchte ich in meiner Diplomarbeit darlegen, wobei natürlich von der jeweiligen unterrichtenden LehrerIn abgewogen werden muss, in welchem Ausmaß dies in ihrer Klasse möglich ist. Es ist in keinsten Weise das Ziel meiner Arbeit oder soll als Ziel angestrebt werden, dass alle Beispiele in dem dargelegten Ausmaß im Unterricht behandelt werden sollen. Die Schulphysik soll und darf natürlich nicht zu einem Vorstadium oder Ersatz eines Universitätsstudiums der Physik ausarten, das natürlich mittlerweile durch und durch mathematisiert ist. Es sollte aber doch durch eine an das Lernniveau angepasste

Mathematisierung den SchülerInnen das Wesen der Physik vorgestellt werden, damit sie sich ein Bild machen können, auch in Bezug darauf, wie sie sich ein Studium dieser Wissenschaft an einer Universität vorzustellen haben.

Ich werde in meiner Arbeit die wichtigsten Themen der Schulmathematik, die in der Physik angewandt werden können, behandeln und ausgewählte, nach meinem Dafürhalten interessante, Beispiele dazu auswählen und versuchen, diese möglichst genau, aber doch in einem vernünftigen Rahmen, vorzuzeigen. Dies beginnt bei sehr einfachen Beispielen, geht dann aber bis hin zu Beispielen der achten Klasse, die sicherlich in den meisten Klassen so nicht im Normalunterricht durchgenommen werden können. Da es aber auch immer ein Ziel einer LehrerIn sein muss, begabte und interessierte SchülerInnen einerseits zu fördern wie auch andererseits zu fordern, sollen diese Beispiele doch aufzeigen, was etwa in einer vorwissenschaftlichen Arbeit möglich sein kann und wie weit diese vom mathematischen Blickwinkel aus gehen kann. Auch sind diese Beispiele durchaus für SchülerInnen geeignet, die in Physik mündlich maturieren wollen. Es sollte auch bedacht werden, dass natürlich auch, speziell natürlich für die begabteren SchülerInnen, Herleitungen und weiterführende Methoden vorgestellt werden, die nicht unbedingt zum Teststoff für alle SchülerInnen der Klasse gehören müssen.

Diese Beispiele bilden dann auch die Grenze für die Schulmathematik und gleichzeitig den Übergang zur Hochschulmathematik beziehungsweise natürlich entsprechend zur Hochschulphysik.

2 Zur Koordination der beiden Unterrichtsfächer Mathematik und Physik

2.1 Was sagen die Lehrpläne zur Rolle der Mathematik im Schulunterricht?

Es soll zu Beginn in aller Kürze ein Blick in die Lehrpläne für Mathematik und Physik geworfen werden, um zu sehen, was denn der Gesetzgeber zur Rolle und Aufgabe der Mathematik im Schulunterricht zu sagen hat. Im Unterstufenlehrplan für Mathematik findet man unter dem Punkt „Beitrag zu den Aufgabenbereichen der Schule“ folgende zwei Ziele des Mathematikunterrichts:

Der Mathematikunterricht soll folgende miteinander vielfältig verknüpfte Grunderfahrungen ermöglichen:

- *Erscheinungen der Welt um uns in fachbezogener Art wahrzunehmen und zu verstehen;*
- *Problemlösefähigkeiten zu erwerben, die über die Mathematik hinausgehen.*¹

Der Gesetzgeber sagt hier sehr klar aus, dass Mathematik nicht auf das eigene Fachgebiet eingeschränkt werden darf, sondern vielmehr dazu dienen muss, „*Erscheinungen der Welt um uns in fachbezogener Art wahrzunehmen und zu verstehen*“ und „*Problemlösefähigkeiten zu erwerben, die über die Mathematik hinausgehen*“.

Im Oberstufenlehrplan für Mathematik findet man unter dem Punkt „Bildungs- und Lehraufgabe“ wiederum sehr klar formuliert, worauf der Mathematikunterricht hinzielen muss:

Beim Erwerben dieser Kompetenzen sollen die Schülerinnen und Schüler die vielfältigen Aspekte der Mathematik und die Beiträge des Gegenstandes zu verschiedenen Bildungsbereichen erkennen.

Die mathematische Beschreibung von Strukturen und Prozessen der uns umgebenden Welt, die daraus resultierende vertiefte Einsicht in Zusammenhänge

¹ Vgl. [3]

und das Lösen von Problemen durch mathematische Verfahren und Techniken sind zentrale Anliegen des Mathematikunterrichts.²

Unter dem Punkt „Mathematische Kompetenzen“ findet sich, welche Kompetenzen die SchülerInnen durch den Mathematikunterricht erwerben sollen:

Sie äußern sich in der Fähigkeit, mathematische Begriffe mit adäquaten Grundvorstellungen zu verknüpfen. Die Schülerinnen und Schüler sollen Mathematik als spezifische Sprache zur Beschreibung von Strukturen und Mustern, zur Erfassung von Quantifizierbarem und logischen Beziehungen sowie zur Untersuchung von Naturphänomenen erkennen.³

Schließlich wird sodann unter dem Punkt „Aspekte der Mathematik“ sehr klar formuliert, dass Mathematik immer einen erkenntnistheoretischen Aspekt miteinbeziehen soll:

Mathematik ist eine spezielle Form der Erfassung unserer Erfahrungswelt; sie ist eine spezifische Art, die Erscheinungen der Welt wahrzunehmen und durch Abstraktion zu verstehen; Mathematisierung eines realen Phänomens kann die Alltagserfahrung wesentlich vertiefen.⁴

Des Weiteren findet sich unter dem Punkt „Beiträge zu den Bildungsbereichen“ unter dem Unterpunkt „Natur und Technik“ eine Aussage über die Notwendigkeit mathematischer Begriffsbildungen und Methoden, um zu einer adäquaten Naturbeschreibung zu gelangen:

Viele Naturphänomene lassen sich mit Hilfe der Mathematik adäquat beschreiben und damit auch verstehen; die Mathematik stellt eine Fülle von Lösungsmethoden zur Verfügung, mit denen Probleme bearbeitbar werden.⁵

Gerade im Unterrichtsfach Physik ist es möglich, diesem Auftrag des Gesetzgebers zu entsprechen, Naturphänomene adäquat in der ihnen entsprechenden Sprache der Mathematik zu beschreiben und Naturphänomene dadurch besser und grundlegender verstehen zu können. Es ist bedauerlich, dass in der Schule zwar durch den Mathematikunterricht tatsächlich eine Fülle von Lösungsmethoden zur Verfügung gestellt wird, diese aber dann doch nicht genutzt werden, weil nicht wenige LehrerInnen der Meinung sind, den SchülerInnen in der Physik durch die Mathematik eine noch größere Belastung aufbürden zu wollen. Damit nehmen sie

² Vgl. [4]

³ Vgl. [4]

⁴ Vgl. [4]

⁵ Vgl. [4]

aber auch gleichzeitig den SchülerInnen die Möglichkeit, gelernte Methoden, die ja bereits zu ihrem Lösungsrepertoire gehören, auch tatsächlich nutzbringend anwenden zu können.

Wie sich der Gesetzgeber nun vorstellt, dass Mathematik unterrichtet werden soll, findet sich im Lehrplan unter dem Punkt „Lernen in anwendungsorientierten Kontexten“. Hier heißt es:

Anwendungsorientierte Kontexte verdeutlichen die Nützlichkeit der Mathematik in verschiedenen Lebensbereichen und motivieren so dazu, neues Wissen und neue Fähigkeiten zu erwerben. Vernetzungen der Inhalte innerhalb der Mathematik und durch geeignete fächerübergreifende Unterrichtssequenzen sind anzustreben. Die minimale Realisierung besteht in der Thematisierung mathematischer Anwendungen bei ausgewählten Inhalten, die maximale Realisierung in der ständigen Einbeziehung anwendungsorientierter Aufgaben- und Problemstellungen zusammen mit einer Reflexion des jeweiligen Modellbildungsprozesses hinsichtlich seiner Vorteile und seiner Grenzen.⁶

Der Gesetzgeber verlangt hier in aller Deutlichkeit, dass Mathematik in anwendungsbezogenen Kontexten gelehrt und von den SchülerInnen erfahren werden muss. Weiters geht der Gesetzgeber davon aus, dass eine Anwendung von Mathematik in verschiedenen Lebensbereichen der SchülerInnen auf die SchülerInnen motivierend wirkt, um neues Wissen und neue Fähigkeiten zu erwerben.

Inwieweit dies natürlich im Mathematikunterricht selbst verwirklicht werden kann, ist ein Problem an sich. Zumeist ist es das vorrangige Ziel der LehrerIn, zumindest den minimalen Unterrichtsstoff, der durch den Lehrplan für die jeweilige Schulstufe vorgegeben wird, im Schuljahr auch tatsächlich zu bewältigen. Mir persönlich ist keine LehrerIn bekannt, die am Ende des Schuljahres noch Zeit zur Verfügung hätte, um noch zusätzlichen Lehrstoff durchzunehmen. Es finden sich natürlich in den Mathematikbüchern sehr wohl Aufgaben, die außermathematische Anwendungen zum Thema haben. Etwa Aufgaben aus der Mechanik (insbesondere Aufgaben zu Geschwindigkeiten und Beschleunigungen), aus der Kernphysik (radioaktiver Zerfall, Altersbestimmung durch die Radiokarbonmethode) und Aufgaben, die Wachstumsprozesse etwa aus der Populationsdynamik zum Thema haben. Aus Gesprächen mit LehrerkollegInnen ist mir bekannt, dass sich oft MathematiklehrerInnen, die als Zweitfach kein naturwissenschaftliches Unterrichtsfach haben, mit gewissen weiterführenden außermathematischen Aufgaben überfordert fühlen, wie etwa beim physikalischen Begriff der

⁶ Vgl. [4]

Arbeit, der in den Mathematiklehrbüchern oft als physikalische Anwendung der Integralrechnung dient. Insofern wäre es dann natürlich zweckmäßig, wenn LehrerInnen, die etwa Physik oder Chemie unterrichten, verstärkt mathematische Methoden anwenden, um die Vorgaben des Lehrplanes aus dem Blickwinkel ihres Fachs heraus zu erfüllen. Hier bietet die Physik wohl sicherlich die meisten Möglichkeiten der direkten Anwendung der im Mathematikunterricht gelehrteten Methoden.

Das findet sich auch explizit im Lehrplan für Physik. Der Unterstufenlehrplan für Physik verlangt den Einsatz mathematischer Methoden natürlich noch sehr eingeschränkt, wie sich unter dem Punkt „Didaktische Grundsätze“ finden lässt:

Bei der Formulierung von Gesetzen ist auf qualitative Je-desto-Fassungen besonderer Wert zu legen. Nur an geeigneten Beispielen ist die Leistungsfähigkeit mathematischer Methoden für die Physik zu zeigen.⁷

Obwohl natürlich das Repertoire an mathematischen Methoden in der Unterstufe noch sehr eingeschränkt ist, lassen sich aber dennoch einige mathematische Grundfertigkeiten sehr gewinnbringend auch in den ersten Jahren in Physik anwenden, wie noch dargelegt werden wird (etwa Gleichungsumformungen oder einfache beschreibend-statistische Methoden). Damit soll den SchülerInnen bereits in der Unterstufe, im angepassten Ausmaß natürlich, auch vor Augen geführt werden, dass die Mathematik als Hilfswissenschaft und „Sprache der Physik“ durchaus ihre Berechtigung hat und Anwendungsmöglichkeiten bietet. Als MathematiklehrerIn ist man regelmäßig mit der Schülerfrage konfrontiert, wozu man denn überhaupt all die viele Mathematik, die man in der Schule lernen muss, brauchen kann und was für einen Nutzen die Mathematik denn eigentlich hat.

Es soll aber auch auf die Formulierung im Lehrplan Bedacht genommen werden, dass mathematische Methoden nur an geeigneten Beispielen angewandt werden sollen. Es geht nicht darum, die Mathematik zwanghaft auf jede physikalische Problemstellung anzuwenden, hier sind das Augenmaß und die Kompetenz der LehrerIn gefragt, geeignete Beispiele zu finden.

Im Oberstufenlehrplan für Physik ist zu lesen,

Mathematisierung als spezifische physikalische Arbeitsweise bedeutet das Durchlaufen verschiedener Stufen zunehmender Abstraktion von der

⁷ Vgl. [5]

Gegenstandsebene über bildliche, sprachliche und symbolische Ebenen zur formal-mathematischen Ebene. Für das Verständnis sind die nichtformalen Ebenen wichtig, während der mathematischen Ebene für die Anwendung (Vorhersage) besondere Bedeutung zukommt.⁸

Zusammenfassend lässt sich also schlussfolgern, dass vom Lehrplan, und das bedeutet natürlich in letzter Konsequenz vom Gesetzgeber, explizit gefordert wird, dass Mathematik und Physik sich gegenseitig ergänzend unterrichtet werden sollen. Das ist natürlich nicht in jedem Teilgebiet gleichermaßen möglich und auch gar nicht gewünscht und zielführend. Wie in der obigen Passage aus dem Lehrplan für Physik für die Oberstufe zu lesen ist, sind die nichtformalen Ebenen genauso wichtig wie die formal-mathematischen und es muss keineswegs alles in der Physik auf Biegen und Brechen mathematisiert werden. Aber in einzelnen, exemplarisch ausgewählten Themengebieten führt es sicherlich zu einer tieferen Verständnisebene, wenn auch Teile des durchgenommenen Physiklehrstoffes in formal-mathematischer Weise behandelt werden. Was dann letztlich und in welchem Ausmaß formal-mathematisch den SchülerInnen dargeboten und abverlangt wird, liegt aber im Ermessen der LehrerIn. Auf jeden Fall würde ein entscheidender Wesenszug der Physik vernachlässigt werden, wenn gänzlich auf mathematische Formulierungen und Herleitungen verzichtet würde.

⁸ Vgl. [6]

2.2 Koordiniertes Unterrichten von Mathematik und Physik

Idealerweise sollten die Lehrpläne für Mathematik und Physik aufeinander abgestimmt sein, was aber natürlich realerwise nicht immer der Fall ist und auch nicht sein kann. Oft hinkt die Mathematik den Anforderungen des Physikstoffes hinterher, sodass natürlicherweise nur unvollständig auf das an und für sich nötige Repertoire an mathematischen Methoden und Begriffen zurückgegriffen werden kann. Umgekehrt ist dieses Problem nicht so schwerwiegend, da die MathematiklehrerInnen nicht unbedingt ein spezielles physikalisches Themengebiet benötigen, um aus diesem Anwendungsbeispiele für den Mathematikunterricht schöpfen zu können. Im Normalfall werden sich die MathematiklehrerInnen auf die Beispiele in den Mathematikschulbüchern beschränken, viele verzichten sogar gänzlich darauf, physikalische Beispiele zu behandeln. Dies wird in Zukunft im Hinblick auf die neue Zentralmatura sicher so nicht mehr möglich sein, da für die zukünftigen Matura auch außermathematische Anwendungen der Mathematik eine größere Rolle spielen werden. Dies soll im übernächsten Kapitel kurz dargelegt werden.

Da an verschiedenen Schulen die Physikstundenanzahl nicht in allen Schulstufen gleich verteilt ist, werde ich, ohne dass sich dadurch etwas Grundlegendes an meinen Ausführungen ändert, von folgenden Stundenverteilungen ausgehen:

Unterstufe: 3 Jahre Physikunterricht, 2., 3. und 4. Klasse

Oberstufe: 3 Jahre Physikunterricht, 6., 7. und 8. Klasse

Ein Beispiel für diese Kluft zwischen der Physik und der benötigten, noch nicht zur Verfügung stehenden Mathematik, sei vorab kurz erwähnt. In der sechsten Klasse Physik wird zu Beginn die Mechanik, also die Bewegungslehre, behandelt und am Anfang werden die beiden grundlegenden Begriffe „Geschwindigkeit“ und „Beschleunigung“ eingeführt. Ist man in der Unterstufe noch mit gleichförmigen Geschwindigkeiten ausgekommen und somit mit der Gleichung $v = \frac{s}{t}$, wobei v die Geschwindigkeit bewegten des Körpers in Meter pro Sekunde, s der zurückgelegte Weg in Metern und t die benötigte Zeit in Sekunden ist, so findet man damit in der Oberstufe nicht mehr das Auslangen. Es wird zwar bereits in der Unterstufe darüber gesprochen, dass die Geschwindigkeit eine gerichtete Größe ist, also immer neben dem Betrag der Geschwindigkeit (oder der „Schnelligkeit“) auch eine Richtung angegeben werden muss. Daneben wird in der Unterstufe rein qualitativ noch der Begriff der

ungleichförmigen Geschwindigkeit erörtert, dass sich also der Betrag der Geschwindigkeit beziehungsweise die Richtung der Bewegung mit der Zeit ändert.

In der Oberstufe müssen diese beiden Komponenten der Geschwindigkeit, also ihr Betrag und ihre Richtung, mathematisiert werden. Für die Angabe der Richtung sind Vektoren notwendig, welche bereits in der fünften Klasse, zumindest im $\mathbb{R}^2$, eingeführt werden. Die Mathematisierung der Richtung der Geschwindigkeit stellt somit kein Problem dar. Anders sieht es aus mit der Änderung des Betrages der Geschwindigkeit, wenn diese keine gleichförmige Geschwindigkeit mehr ist, was ja bereits beim freien Fall oder dem schiefen Wurf der Fall ist. Hierfür wäre mathematisch korrekterweise der Differentialquotient beziehungsweise die erste Ableitung für die Geschwindigkeit und die zweite Ableitung für die Beschleunigung Voraussetzung, welche aber erst in der siebten Klasse im Rahmen der Differenzialrechnung unterrichtet werden. Sozusagen aus der Patsche hilft man sich im Unterricht, indem man den Begriff der Durchschnittsgeschwindigkeit $\bar{v}$ einführt, welche mathematisch dem Differenzenquotienten $\bar{v} = \frac{s_2 - s_1}{t_2 - t_1} = \frac{\Delta s}{\Delta t}$ entspricht, welcher aber natürlich

in der sechsten Klasse noch nicht so bezeichnet wird. Sodann wird qualitativ weiter so vorgegangen, dass man das Zeitintervall $\Delta t = t_2 - t_1$ immer kleiner werden lässt. Die Momentangeschwindigkeit v als Grenzübergang $v = \lim_{\Delta t \rightarrow 0} \frac{\Delta s}{\Delta t}$ kann man eigentlich noch nicht mathematisch genau definieren, da der Grenzwert (meist) ebenfalls noch nicht zur Verfügung steht⁹. Also baut man hier auf das intuitive Verständnis, dass sich die SchülerInnen diesen Grenzübergang auch vorstellen können, rechnen kann man damit natürlich keineswegs. Wenn dann im Mathematikunterricht in der siebten Klasse die Ableitung im Rahmen der Differentialrechnung durchgenommen wird, wird zumeist, aber jetzt nicht im Physikunterricht, sondern im Mathematikunterricht, im Nachhinein eine exakte Definition der Geschwindigkeit und analog dazu der Beschleunigung nachgereicht.

An diesem Beispiel sollte klar geworden sein, dass ein streng koordinierter Unterricht von Mathematik und Physik oftmals einfach nicht möglich ist. Weitere Beispiele wären etwa die Integralrechnung, die in der achten Klasse unterrichtet wird, aber im Grunde bereits in der sechsten Klasse etwa für die Definition und Berechnung der Arbeit benötigt wird, speziell wenn es um die Arbeit geht, die bei der Bewegung in einem Gravitationsfeld oder einem elektrischen Feld verrichtet wird. Wie man hier aber beispielsweise die Integralrechnung im

⁹ Der Grenzwert einer Folge wird in der sechsten Klasse eingeführt, in der siebten Klasse wird dieser exaktifiziert.

Physikunterricht „umgehen“ kann, wird in einem späteren Kapitel meiner Arbeit dargelegt werden.

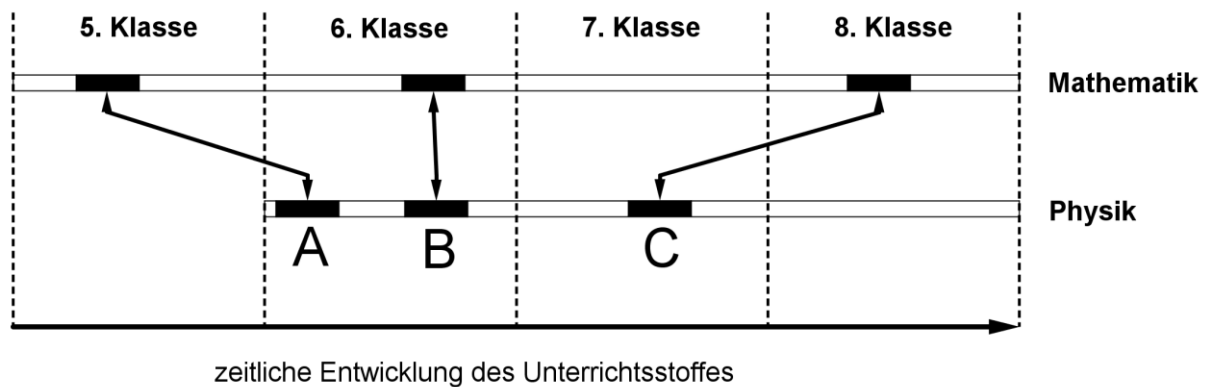


Abbildung 1: Zeitliche Koordination von Mathematik und Physik

Es können im Grunde drei Fälle auftreten, die ich in der Abbildung 1 zu veranschaulichen versucht habe, dabei entsprechen Rechtecke Kapiteln aus der Mathematik, die in einem entsprechenden Kapitel der Physik zur Anwendung kommen können.

- Fall A

Der Mathematiklehrstoff wurde bereits behandelt und steht für den Physikunterricht uneingeschränkt zur Verfügung. Da der Physikunterricht in der Oberstufe im Regelfall in der sechsten Klasse Oberstufe beginnt, ist dieser auch erst ab der sechsten Klasse angegeben. Dieser Fall stellt natürlich einen Idealfall dar und ist in vielen Fällen nicht erfüllt.

- Fall B

Der Mathematik- und Physiklehrstoff wird parallel behandelt. Dies ist in der sechsten Klasse in mehreren Bereichen möglich, etwa bei der Exponential- und Logarithmusfunktion. Hier wäre es natürlich im Idealfall von Vorteil, wenn die Behandlung der beiden einander entsprechenden Lehrstoffe in Physik und Mathematik nicht völlig parallel verlaufen, sondern ein gewisser zeitlicher Unterschied vorhanden ist. So kann der Stoff im Mathematikunterricht vorab eingeübt und gefestigt werden und dann später im Physikunterricht in der Anwendung wiederholt werden.

Falls Mathematik und Physik von verschiedenen Lehrpersonen unterrichtet werden, wäre hier ein Abgleich der Jahresplanung von Vorteil, sodass daraus beide Fächer profitieren würden. Falls dieselbe Lehrperson beide Fächer in derselben Klasse unterrichtet, ist es natürlich leicht möglich, den Unterricht entsprechend zu koordinieren, was auch den Vorteil hat, dass die LehrerIn genau weiß, was man in der Klasse voraussetzen darf in Hinblick auf das Wissen im jeweiligen Stoffgebiet. Für die SchülerInnen ist dann durch den kurzen zeitlichen Abstand direkt erkennbar, welchen Nutzen die Mathematik in der physikalischen Anwendung hat. Insbesondere, dass man eben nicht nur qualitativ über gewisse Dinge in der Physik sprechen kann, sondern durch die Mathematik ein Hilfsmittel in der Hand hat, um auch quantitative Aussagen treffen zu können. So erfahren die SchülerInnen, dass die Mathematik die geeignete Sprache darstellt, um im Rahmen der Physik Aussagen über die Naturerscheinungen treffen zu können.

Auch kann damit den SchülerInnen ein wichtiges Ziel der Physik vor Augen geführt werden, nämlich dass es eines der Grundziele der Physik ist, quantitative Voraussagen machen zu können. Andererseits erfahren die SchülerInnen auch, dass die Methoden der Mathematik sehr wohl nicht nur reiner Selbstzweck der Mathematik sind, sondern eine direkte Anwendung und einen gewinnbringenden Nutzen in vielen Anwendungsbereichen haben. Dies stellt dann in Summe eine auch vom Lehrplan geforderte fächerübergreifende Verschränkung der beiden Fächer dar.

- Fall C

Der Mathematiklehrstoff steht noch nicht zur Verfügung zu der Zeit, in der der entsprechende Physikstoff gelehrt wird. Für den Mathematikunterricht stellt dies natürlich kein Problem dar, da man ja umgekehrt auf den Physikstoff zurückgreifen kann und damit auch die Vorteile der behandelten mathematischen Methode herausstreichen kann.

Wieder sei etwa die Geschwindigkeit als Beispiel herausgegriffen. Wurde diese im Physikunterricht als Durchschnittsgeschwindigkeit eingeführt, kann jetzt im Mathematikunterricht der Schritt zur Momentangeschwindigkeit ganz selbstverständlich vollzogen werden. Voraussetzung hierfür ist natürlich, dass die LehrerIn sich das auch zutraut, was im Falle der Geschwindigkeit sicher kein Problem für jede kompetente MathematiklehrerIn darstellen sollte. Es ist mir als Physiklehrer schon einige Male passiert, dass MathematiklehrerInnen auf mich zugekommen sind mit der Frage, wie es

sich mit gewissen physikalischen Begriffen genauer verhält, explizit etwa der Arbeit als Integral über eine Kraft.

Umgekehrt stellt das für den Physikunterricht mitunter ein größeres Hindernis dar, da dann nicht mehr direkt auf die mathematischen Begriffe und Methoden zugegriffen werden kann, sondern indirekte und womöglich umständlichere Wege herangezogen werden müssen. Etwa wird in der sechsten Klasse oft ein kleines Kapitel über die Vektorrechnung in die Mechanik eingeschoben, falls der Begriff des Vektors noch nicht zur Verfügung steht. Wenn man sich Schulbücher ansieht, etwa „Big Bang“ von Apolin, so werden Vektoren eigens in einem eigenen Kapitel „Tool Time“¹⁰ behandelt.

Dies kann nun natürlich auch gewinnbringend verwendet werden, wenn hier durch die PhysiklehrerIn gezeigt werden kann, dass gewisse mathematische Begriffe ihren Ursprung in physikalischen Fragestellungen haben, das Beispiel der Vektoren ist hier sicher eines der „prominentesten“ Beispiele. Ob dann letztlich für die SchülerInnen auch ein Erkenntnisgewinn dahingehend resultiert, sprich „Wie funktioniert eigentlich Physik und wie spielen Mathematik und Physik zusammen und wie bedingen sie sich gegenseitig?“, das hängt aber in erster Linie von der unterrichtenden LehrerIn ab.

¹⁰ Siehe [18], Kapitel , Seite 3 ff.

2.3 Ziele und Grenzen dieser Diplomarbeit

Erstens

Wie ich im vorangegangenen Kapitel dargelegt habe, wäre eine zeitliche Koordination des Lehrplanes für Mathematik und des Lehrplanes für Physik zwar erwünscht, dies wird in der Praxis beim Unterrichten der beiden Unterrichtsgegenstände meistens aber nicht möglich sein. Mein Ziel in dieser Diplomarbeit ist es, für alle acht Schulstufen einer Allgemeinbildenden Höheren Schule auszuführen, welche mathematischen Methoden für welche physikalischen Inhalte geeignet sind. Ich werde dabei von der Mathematik ausgehen und die mathematischen Lerninhalte der jeweiligen Schulstufen ausgewählten physikalischen Themen zuordnen. Natürlich wird es in vielen Fällen nicht möglich sein, dass der Mathematiklehrstoff einer bestimmten Schulstufe unmittelbar dem Physiklehrstoff der entsprechenden Schulstufe entspricht, allerdings werde ich dies in dieser Arbeit nicht weiter berücksichtigen, da in dieser Diplomarbeit die fachliche und nicht die zeitliche Koordination das hauptsächliche Thema sein soll. Hier ist dann die unterrichtende LehrerIn gefordert, die zeitliche Koordination der Lehrinhalte selbst zu bewerkstelligen.

Zweitens

Es wird von mir auch nicht der Versuch unternommen werden, zu allen mathematischen Themen der Mathematik der Unter- und Oberstufe Beispiele zu konstruieren, sondern es sollen die mathematischen Themen mit Beispielen unterlegt werden, durch welche nach meinem Dafürhalten ein wirklicher Nutzen in der Physik erkennbar sein soll. Dadurch soll auch ersichtlich werden, dass die Mathematik, die die SchülerInnen in der Schule lernen, sehr wohl Sinn macht, indem man eben die für viele SchülerInnen trockene Schulmathematik tatsächlich anwendet und gebraucht, um quantitative Aussagen über die Natur zu treffen beziehungsweise physikalische Gleichungen und Gesetzmäßigkeiten mit Hilfe der Mathematik herzuleiten.

Leider ist dies gerade bei einigen der bemerkenswertesten und aus der Sicht vieler Physiker schönsten Gleichungen in der Schule nicht möglich. Doch sollte zumindest durch die PhysiklehrerIn darauf hingewiesen werden, wie diese für die Naturbeschreibung grundlegendsten Gleichungen der Physik aussehen und weshalb sie eine innere Schönheit besitzen, die oft erst durch tiefergehende mathematische Kenntnisse sichtbar wird. Ich denke

hier etwa an das Differentialgleichungssystem der vier Maxwellgleichungen, aus denen die gesamte Theorie des Elektromagnetismus hergeleitet werden kann oder die Schrödingergleichung, die etwa das Wasserstoffatom quantenmechanisch beschreibt¹¹.

Drittens

Meine Ambition ist es, dass diese Diplomarbeit eine Art kleines, wenngleich natürlich auch unvollständiges, Nachschlagewerk sein soll, in dem die Physik- beziehungsweise MathematiklehrerIn Material und Ideen findet, einen koordinierten und fächerübergreifenden Unterricht gestalten zu können, wobei ich natürlich keineswegs den Anspruch auf Vollständigkeit erhebe.

Infolge des gewählten Aufbaus der Diplomarbeit, dass ich mich nach dem Lehrplan für Mathematik beziehungsweise der Reihenfolge des Auftretens des Mathematiklehrstoffes anhand der Mathematikschulbücher richte, wird es natürlich des Öfteren vorkommen, dass die Mathematik für Anwendungen, die in der Oberstufe in Physik Unterrichtsstoff sind, bereits im Mathematiklehrbuch der Unterstufe, im Extremfall sogar im Schulbuch der ersten Klasse zu finden sind.

Viertens

Die Anwendungsbeispiele sind teilweise, wenn auch in der Minderzahl, direkt den Mathematikschulbüchern entnommen, in vielen Fällen sind es Beispiele und Aufgaben aus Physikbüchern, die ich als Grundgedanken übernommen und dann „mathematisiert“ habe. Manche Anwendungsbeispiele habe ich mir selbst ausgedacht oder sie gründen sich auf Ideen aus Aufgaben, die ich in der Vergangenheit in Büchern oder Fachmagazinen gelesen habe, deren Ursprung ich aber nicht mehr eruieren kann. Viele mathematische Herleitungen, die in den Physikbüchern nicht durchgeführt werden, habe ich ergänzt beziehungsweise die Herleitungen überhaupt erst durchgeführt. Es ist klar, dass die Ausführlichkeit vieler Herleitungen über das Ausmaß hinausgehen, die die PhysiklehrerIn im Unterricht den SchülerInnen vortragen wird, aber es ist meines Erachtens doch wichtig, dass die PhysiklehrerIn zumindest dieses mathematische Hintergrundwissen im Kopf hat, um auf Fragen der SchülerInnen fachlich korrekte Antworten bieten zu können.

¹¹ Zumindest nichtrelativistisch, eine relativistische Beschreibung ist durch die Diracgleichung gegeben.

Fünftens

Ich möchte ganz besonders darauf hinweisen, dass es nicht das Ziel sein kann, dass alle Beispiele im Physikunterricht beziehungsweise im Mathematikunterricht gebracht werden. Die LehrerIn muss eine Auswahl treffen, was natürlich von mehreren Faktoren abhängen kann. Die Tiefe der Mathematisierung muss etwa an das jeweilige Leistungsniveau der Klasse angepasst werden, wobei aber laut Lehrplan mindestens ein Thema pro Schulstufe eine tiefergehende Behandlung erfahren sollte. Im Lehrplan für Physik in der Oberstufe heißt es,

*Dabei ist exemplarisch an mindestens einer Thematik pro Schulstufe eine größere Erklärungstiefe anzustreben und vermehrte Möglichkeit zur eigenständigen Befassung zu geben. Dies ist nach Möglichkeit auch fächerübergreifend durchzuführen.*¹²

Dieses Thema könnte und sollte man natürlich, insbesondere in der Oberstufe, auch mathematisch tiefergehend beleuchten. Ich möchte aber doch die Warnung aussprechen, dass die Mathematisierung der Physik in der Schule in einem vernünftigen Rahmen bleiben muss, da dies nur einen von mehreren möglichen Zugängen zur Physik darstellt. Es ist dies zwar meines Erachtens ein durchaus wichtiger, er darf aber trotzdem nur einer unter verschiedenen sein. Im Lehrplan heißt es dazu:

*Mathematisierung als spezifische physikalische Arbeitsweise bedeutet das Durchlaufen verschiedener Stufen zunehmender Abstraktion von der Gegenstandsebene über bildliche, sprachliche und symbolische Ebenen zur formal-mathematischen Ebene. Für das Verständnis sind die nichtformalen Ebenen wichtig, während der mathematischen Ebene für die Anwendung (Vorhersage) besondere Bedeutung zukommt.*¹³

Es darf also der mathematische Zugang keineswegs vernachlässigt werden und muss seinen Platz neben anderen, nicht minder bedeutsamen Zugängen einnehmen. Dass aber gänzlich auf eine Mathematisierung verzichtet wird, wie es sicherlich auch geschieht, ist nicht zu rechtfertigen, insbesondere wenn sich die LeserIn die Passage aus dem Lehrplan vor Augen hält, dass „der mathematischen Ebene für die Anwendung (Vorhersage) besondere Bedeutung zukommt“.

¹² Vgl. [6]

¹³ Vgl. [6]

Ich werde aus den zahlreichen Schulbüchern, die es naturgemäß gibt, diejenigen heranziehen, die ich aus meiner eigenen Unterrichtspraxis kenne. Dies soll aber keineswegs eine wie auch immer geartete Wertung der Schulbücher darstellen, aber diese Diplomarbeit hat ja nicht das Ziel, die diversen Mathematik- und Physikschulbücher zu vergleichen, auch habe ich versucht, den Stoff nicht auf ein vorgegebenes Schulbuch beziehungsweise eine Schulbuchreihe einzuschränken. Ich werde auch durchaus Herleitungen, Beispiele und Anwendungen bringen, die nicht so in den Schulbüchern stehen, da es ja grundsätzlich nur um die Möglichkeiten der Anwendung der Mathematik im Physikunterricht geht, was schulbuchunabhängig zu sehen ist. Die Wahl eines konkreten Schulbuches dient einfach dazu, ein Gerüst zu haben, an dem man sich leichter orientieren kann.

2.4 Die Reifeprüfung vom 9. Mai 2014

Mit dem Schuljahr 2014/15 wird in ganz Österreich erstmals flächendeckend die neue Reifeprüfung an den Allgemeinbildenden Höheren Schulen und im Schuljahr 2015/16 an Berufsbildenden Höheren Schulen durchgeführt werden. Bereits im Schuljahr zuvor, im Mai 2014, nahmen 50 Klassen an einem Schulversuch zur standardisierten kompetenzorientierten schriftlichen Reifeprüfung in Mathematik an den Allgemeinbildenden Höheren Schulen teil.

Was nach Bekanntgabe der Probereifeprüfung mich als Mathematik- und Physiklehrer und insbesondere KollegInnen, die nur das Unterrichtsfach Mathematik und nicht Physik unterrichten, erstaunte, war der überraschend hohe Anteil an Aufgaben, die aus dem Bereich der Physik stammten. Wenn man sich dann aber dahingehend informiert, was denn die neuen Zielvorgaben sind, die mit der standardisierten kompetenzorientierten schriftlichen Reifeprüfung eingefordert werden, so verwundert es nicht, dass die Aufgabenstellungen aus dem Bereich der Mathematik einem großen Wandel unterzogen sind.

Die standardisierten Reifeprüfungsaufgaben beinhalten sogenannte Typ-1 und Typ-2-Aufgaben. Auf der Homepage des bifie findet man Typ-2-Aufgaben folgendermaßen beschrieben:

„Typ-2-Aufgaben sind Aufgaben zur Anwendung und Vernetzung der Grundkompetenzen in definierten Kontexten und Anwendungsbereichen. Es handelt sich dabei um umfangreichere kontextbezogene oder auch innermathematische Aufgabenstellungen, bei denen eine selbstständige Anwendung von Wissen und Können erforderlich ist.“¹⁴

Es geht also insbesondere um Aufgaben, die die Grundkompetenzen voraussetzen und diese dann anwenden und vernetzen sollen, wobei seine selbstständige Anwendung derselben verlangt wird. Dieses Anwenden kann aber sicher nicht nur auf rein mathematische Aufgabenstellungen beschränkt sein, sondern damit soll den SchülerInnen nahegebracht werden, dass eben die Mathematik, als Hilfswissenschaft verstanden, nicht nur theoretisches Wissen beinhaltet, sondern insbesondere anwendungsbezogenes. Letztlich sollen die SchülerInnen die Mathematik auch als eine Wissenschaft und ein Werkzeug kennenlernen, das dazu dient, anderen Wissenschaften ein Handwerkszeug zur Verfügung zu stellen, wodurch diese Wissenschaften erst zur wirklichen Entfaltung kommen können. David Hilbert,

¹⁴ Vgl. [2]

einer der ganz Großen der Mathematik, hat einen wunderbaren Satz geschrieben, der dies unübertrefflich auszudrücken vermag:

*Die Mathematik ist das Instrument, welches die Vermittlung bewirkt zwischen Theorie und Praxis, zwischen Denken und Beobachten: Sie baut die verbindende Brücke und gestaltet sie immer tragfähiger. Daher kommt es, daß unsere ganze gegenwärtige Kultur, soweit sie auf der geistigen Durchdringung und Dienstbarmachung der Natur beruht, ihre Grundlage in der Mathematik findet.*¹⁵

Wenn man sich an die eigenen Maturaaufgaben zurückerinnert oder sich Maturaaufgaben vor Einführung der standardisierten kompetenzorientierten schriftlichen Reifeprüfung ansieht, insbesondere wenn diese eine längere Zeit zurückliegen, fällt doch sehr eklatant die neue Zielvorgabe auf, nicht mehr wie bisher rein mathematische Aufgabenstellungen lösen zu müssen, sondern Aufgabenstellungen, die zu einem nicht unbeträchtlichen Teil aus dem Bereich der Physik oder auch anderer Naturwissenschaften herrühren.

Beim Schulversuch vom 9. Mai 2014 kamen von 24 gestellten Typ-1-Aufgaben vier und von fünf gestellten Typ-2-Aufgaben drei aus dem Bereich der Physik beziehungsweise der Chemie/Biologie. Das sind immerhin sieben von 29 Aufgaben, die einen konkreten Physik/Chemie/Biologie-Bezug hatten, also beinahe 25%, was doch sehr bemerkenswert ist! Jedem Mathematik- beziehungsweise Physiklehrer ist die Frage von SchülerInnen „Machen wir jetzt Mathematik oder Physik?“ hinlänglich bekannt, sobald im Mathematikunterricht physikalische Beispiele als Anwendungen der Mathematik gebracht oder andererseits im Physikunterricht mathematische Methoden angewandt werden. In Zukunft ist die Anwendung mathematischer Methoden im Physikunterricht genauso wie physikalische Fragestellungen im Mathematikunterricht eine Herausforderung, die durch die neue Reifeprüfung an alle SchülerInnen gestellt werden wird und die diese bewältigen werden müssen. Das bedeutet natürlich eine Umstellung im Verständnis, was sowohl das Unterrichtsfach Mathematik als auch das Unterrichtsfach Physik in Zukunft werden leisten müssen. Es sollte aber nicht als nur als eine zusätzliche und neue Forderung verstanden werden, sondern auch als eine Chance, den Unterricht interessanter und in größerem Ausmaß als bisher fächerübergreifend zu gestalten.

¹⁵ David Hilbert, aus der Rede *"Naturerkennen und Logik"*, gehalten auf der Versammlung Deutscher Naturforscher und Ärzte in Königsberg, 8. September 1930

Folgende Aufgabenstellungen wurden im Schulversuch vom 9. Mai 2014 gestellt¹⁶:

Typ-1-Aufgaben:

- Aufgabe 2: Punktladungen
- Aufgabe 7: Zerfallsprozess
- Aufgabe 15: Nikotin
- Aufgabe 18: Pflanzenwachstum

Typ-2-Aufgaben:

- Aufgabe 2: Zustandsgleichung idealer Gase
- Aufgabe 3: Chemische Reaktionsgeschwindigkeit
- Aufgabe 5: Sportwagen

Daran anknüpfend möchte ich mit meiner Diplomarbeit exemplarisch darlegen, inwieweit die Mathematik im Physikunterricht eine Rolle spielen kann beziehungsweise welche physikalischen Themen gewinnbringend auf die Mathematik als geeignetes Hilfsmittel zurückgreifen können.

Grundsätzlich bietet es sich an, entweder von der Mathematik auszugehen und die physikalischen Themen zu untersuchen oder umgekehrt die physikalischen Themen im Schulunterricht als Ausgangsbasis herzunehmen und davon ausgehend sich die Mathematik anzusehen, die zur Anwendung kommen kann. Ich habe mich dafür entschieden, von der Seite der Mathematik das Thema anzugehen und davon ausgehend zu untersuchen und darzulegen, welche physikalischen Themen diese mathematischen Methoden verwenden können oder zumindest die Möglichkeit dazu haben, was natürlich immer vom jeweiligen Niveau der Klasse abhängt. Einige der gezeigten Beispiele, insbesondere natürlich diejenigen der höheren Klassen, können auch im Rahmen von Wahlpflichtfächern verwendet werden beziehungsweise als Spezialthemen für die mündliche Matura.

¹⁶ Siehe [2]; die gesamte Reifeprüfung vom Mai 2014 ist im Anhang beigelegt.

3 Mathematische Anwendungen in der Physik nach Klassen geordnet

Da sich, wie bereits ausgeführt, grundsätzlich zwei Möglichkeiten der Behandlung des Themas anbieten, entweder eine Orientierung nach dem Mathematik- oder dem Physiklehrplan, musste ich eine Entscheidung treffen und habe die erste Möglichkeit gewählt, werde also bei der Auswahl der Mathematik nach den Klassen aufsteigend vorgehen, wobei ich aus jeder Schulstufe relevante Stoffkapitel auswählen werde. Sodann werde ich passende Beispiele aus der Physik präsentieren, wobei dies dann meistens aber nicht mehr mit den gleichen Schulstufen wie der durchgenommene Mathematikstoff korreliert, wie ich auch bereits dargelegt habe.

Werden etwa Maßeinheiten der Mathematik der ersten Klasse zugeordnet, so entsprechen die dazu gehörigen Beispiele etwa der Physik der sechsten Klasse der Oberstufe. Viele zur Physik der Oberstufe passende mathematische Themen werden bereits in der Unterstufe im Mathematikunterricht behandelt, wie etwa der pythagoreische Lehrsatz. Es ist einfach meist nicht möglich, die Physik und Mathematik einer Schulstufe parallel zu behandeln und die Unterrichtsstoffe in Einklang zu bringen, da nur in sehr wenigen Fällen der durchgenommene Stoff in Physik mit der benötigten Mathematik und umgekehrt zusammenfällt.

3.1 1. Klasse

3.1.1 Gleichungen und direktes Verhältnis

Gleichungen werden bereits sehr frühzeitig im ersten Unterrichtsjahr im Mathematikunterricht eingeführt. Wird in der ersten Klasse zunächst der Begriff der Variablen erklärt¹⁷ und Lösungen von einfachsten Gleichungen noch in einer unsystematischen Weise durch Probieren gefunden¹⁸, so wird das Konzept der Gleichung in der zweiten Klasse in bereits systematischerer Weise fortgesetzt¹⁹. In der ersten Klasse werden zunächst Anwendungen von einfachsten Gleichungen in Form von Textaufgaben und Formeln behandelt²⁰. Es geht dabei in erster Linie darum, dass die SchülerInnen einen ersten Einstieg in eine doch für die erste Klasse recht abstrakte Methode der Handhabung und der Lösung von Textaufgaben erhalten. Dabei wird besonderer Wert gelegt auf die „Übersetzung“ eines mathematischen Textes in eine adäquate mathematische Sprache.

In der dritten Klasse werden Terme eingeführt²¹ und das Rechnen mit Termen²² behandelt und auch Gleichungen durch Äquivalenzumformungen²³ gelöst, wobei ab jetzt das Repertoire an Gleichungen und Lösungsverfahren mit jedem Schuljahr erweitert wird, indem im Laufe der Zeit sukzessive neue Rechentechniken und Funktionen eingeführt werden. Vom Wurzelziehen in der dritten²⁴ und insbesondere in der vierten²⁵ Klasse über das Lösen von zwei oder mehr Gleichungen in mehreren Unbekannten²⁶ bis zum Logarithmus²⁷ und dem Lösen von Exponentialgleichungen²⁸.

¹⁷ Siehe [7]; Seite 83

¹⁸ Siehe [7]; Seite 84

¹⁹ Siehe [8]; das Kapitel „Gleichungen und Formeln“ ab Seite 80

²⁰ Siehe [7]; Seite 87 und besonders das Kapitel „Gleichungen aus Textaufgaben“ ab Seite 88

²¹ Siehe [9]; das Kapitel „Terme“ ab Seite 60

²² Siehe [9]; das Kapitel „Rechnen mit Termen“ und die Unterkapitel dazu ab Seite 69

²³ Siehe [9]; das Kapitel „Gleichungen und Formeln“ ab Seite 99

²⁴ Siehe [9]; das Kapitel „Berechnungen mit dem Satz von Pythagoras – Quadratwurzel“ ab Seite 226

²⁵ Siehe [10]; das Kapitel „Quadratwurzeln“ und die Unterkapitel dazu ab Seite 16

²⁶ Siehe [10]; das Kapitel „Lineare Gleichungen mit zwei Variablen“ ab Seite 98;

[11] das Kapitel „Lineare Gleichungen und Gleichungssysteme in zwei Variablen“ ab Seite 194;

[12] das Kapitel „Lineare Gleichungssysteme in drei Variablen; Lage von drei Ebenen“ ab Seite 193

²⁷ Siehe [12] das Kapitel „Logarithmen“ ab Seite 28 und das Kapitel „4 Exponential- und Logarithmusfunktionen“

²⁸ Siehe [12] das Kapitel „Anwendungen von Exponentialfunktionen“ ab Seite 60 und das Kapitel „Exponentialgleichungen“ ab Seite 75

In der ersten Klasse wird auch das direkte Verhältnis eingeführt²⁹, wobei eine genauere Behandlung erst in der zweiten Klasse gemeinsam mit dem indirekten Verhältnis erfolgt³⁰. Zumindest aber die Idee, dass zwei Größen bei einem direkten Verhältnis derart zusammenhängen, dass eine Verdoppelung, Verdreifachung der einen Größe eine entsprechende Verdoppelung, Verdreifachung der anderen Größe nach sich zieht, sollte in der ersten Klasse bereits verstanden werden.

Im Mathematikunterricht geschieht eine Lösung von direkten und indirekten Proportionalitätsaufgaben meist durch das Aufstellen einer Tabelle³¹, mittels der dann der Wert der gesuchten Größe ermittelt wird. Das Aufstellen einer expliziten Formel, der ein direktes Verhältnis zugrunde liegt, spielt eigentlich im Mathematikunterricht keine Rolle und wird auch so nicht behandelt. Das ist aber wiederum genau der Schritt, der in der Physik die entscheidende Rolle spielt. Hier steht als Ausgangsgröße und zugeordnete Größe meist eine Tabelle zur Verfügung, die etwa aus einer Messserie gewonnen wurde und aus der dann abgeleitet werden soll, dass es sich um ein direktes oder indirektes Verhältnis handelt. Der darauf und daraus folgende Schluss, dass ein direktes Verhältnis so gedeutet werden muss, dass die eine Größe aus der anderen durch Multiplikation mit einem konstanten Faktor hervorgeht, stellt oft ein großes Problem für SchülerInnen dar, ist aber natürlich in der Physik von grundlegender Bedeutung.

Zum ersten Mal tritt dieses Problem bereits ganz am Anfang des ersten Physikjahres auf, wenn es um die Einführung der Geschwindigkeit und den Zusammenhang der drei Größen Weg, Geschwindigkeit und Zeit bei einer geradlinig gleichförmigen Bewegung geht³². Das folgende Beispiel hierzu mit den entsprechenden Zahlenwerten basiert auf einem Versuch, den der Autor dieser Arbeit in einer zweiten Klasse in Physik standardmäßig durchführt³³.

²⁹ Siehe [7] das Kapitel „Direktes Verhältnis“ ab Seite 69

³⁰ Siehe [8] das Kapitel „Direkte und indirekte Proportionalität“ ab Seite 98

³¹ Siehe [8] Seite 100 (direkte Proportionalität) und 105-106 (indirekte Proportionalität)

³² Siehe [15] das Kapitel „Die Geschwindigkeit“ ab Seite 12

³³ Der Versuch, dem die hier aufgeführten Zahlenwerte zugrunde liegen, wurde in einer zweiten Klasse, die der Autor in Physik unterrichtet hat, im Evangelischen Gymnasium und Werkschulheim im Schuljahr 2014/15 durchgeführt.

Beispiel 1: Aufstellen einer Formel für die Geschwindigkeit

Der Versuch wurde so durchgeführt, dass SchülerInnen mit einer Stoppuhr (Handy) in konstanten Abständen entlang einer geraden Linie positioniert worden sind. Eine SchülerIn läuft an ihnen mit möglichst konstanter Geschwindigkeit vorbei. Jede SchülerIn stoppt die Zeit, sobald die LäuferIn genau an ihr vorbeiläuft. Aus diesen Stoppzeiten wird dann die Geschwindigkeit der LäuferIn ermittelt.

Aus den gemessenen Zeiten erkennt man in der Auswertung durch Eintrag in eine Tabelle, dass die Geschwindigkeit etwa konstant ist, woraus geschlossen werden kann, dass hier ein direktes Verhältnis zwischen dem zurückgelegten Weg und der Zeit mit der Geschwindigkeit als Proportionalitätsfaktor vorliegt. Daraus lässt sich dann eine Formel für den Weg in Abhängigkeit von der Zeit für eine geradlinig gleichförmige Bewegung angeben³⁴.

Gesamtweg s in m	Läufer 1		Läufer 2	
	Gemessene Zeit t in s	Geschwindigkeit in ms^{-1} : $v = s : t$	Gemessene Zeit t in s	Geschwindigkeit in ms^{-1} : $v = s : t$
2	2,0	1,0	1,2	1,7
4	3,7	1,1	2,6	1,5
6	6,1	1,0	3,5	1,7
8	9,2	0,9	4,6	1,7
10	10,1	1,0	5,7	1,8
12	12,2	1,0	6,5	1,8
14	14,0	1,0	7,7	1,8
16	16,5	1,0	8,7	1,8
18	18,2	1,0	10,2	1,8

Tabelle 1: Ergebnisse und Auswertung des Versuchs zur Geschwindigkeitsbestimmung

³⁴ In der 2. Klasse wird natürlich noch kein Unterschied zwischen Durchschnitts- und Momentangeschwindigkeit gemacht, die ermittelte Geschwindigkeit ist hier eigentlich eine Durchschnittsgeschwindigkeit jeweils über die Wegstrecken vom Ausgangspunkt $s = 0 \text{ m}$ bis zum Punkt, an dem die Zeit bestimmt wurde.

In der Tabelle 1 sind die gemessenen Zeiten protokolliert, wie sie im Versuch tatsächlich ermittelt worden sind. Es werden die Zeiten zweier LäuferInnen gemessen und daraus die Geschwindigkeiten ermittelt. In der Spalte „Geschwindigkeit“ ist sodann ersichtlich, dass die Geschwindigkeit ungefähr konstant ist und dass LäuferIn 2 eine höhere Geschwindigkeit als LäuferIn 1 hatte und somit die schnellere gewesen ist. In einem weiteren Schritt kann diese Tabelle dann natürlich auch grafisch ausgewertet werden, wobei sich dann herausstellt, dass eine gleichförmige Bewegung in einem Weg-Zeit-Diagramm durch eine Gerade dargestellt wird.

Für LäuferIn 1 erhält man als mittlere Geschwindigkeit $v = 1,0 \frac{\text{m}}{\text{s}}$, für LäuferIn 2 $v = 1,8 \frac{\text{m}}{\text{s}}$. Als Formel für die Geschwindigkeit ergibt sich dann $s = v \cdot t$ und daraus schließlich durch eine einfache Umformung

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

Beispiel 2: Gewichtskraft

Nach der Einführung der beiden Größen Masse und Kraft wird üblicherweise die Gewichtskraft als Spezialfall einer Kraft behandelt³⁵. Mit einer Federwaage oder einem elektronischen Kraftmesser kann der Zusammenhang zwischen der Masse eines Körpers und der Gewichtskraft, die auf den Körper wirkt, bestimmt werden. Wird die Gewichtskraft mehrerer Massen verschiedener Größe gemessen, ergibt sich wiederum ein direktes Verhältnis zwischen der Masse eines Körpers und der Gewichtskraft, wobei der konstante Faktor die Erdbeschleunigung g darstellt³⁶:

³⁵ Siehe [15] das Kapitel „9 Gewichtskraft“ ab Seite 20.

³⁶ g ist für unterschiedliche Punkte auf der Erde keine Konstante mehr. g hängt einerseits von der Höhe über der Erdoberfläche und andererseits von der Position des Messpunktes ab. So ist g an den Polen kleiner als am Äquator, hängt aber auch davon ab, ob man g etwa auf dem Festland oder über dem Meer bestimmt, da auch die Zusammensetzung der Erde Auswirkungen auf die Erdbeschleunigung g hat. Die Erdbeschleunigung g hat etwa an den Polen den Wert $9,83220 \frac{\text{m}}{\text{s}^2}$ und am Äquator den Wert $9,78050 \frac{\text{m}}{\text{s}^2}$, ist also an den Polen um etwa

0,53% größer als am Äquator!

$$F_G = g \cdot m$$

F_G ... Gewichtskraft des Körpers, $[F_G] = \text{N}$

m ... Masse des Körpers, $[m] = \text{kg}$

g ... Erdbeschleunigung, $[g] = \frac{\text{m}}{\text{s}^2}$

Beispiel 3: Newton'sche Bewegungsgleichung

Die Newton'sche Bewegungsgleichung stellt eine der wichtigsten Grundgleichungen der Physik dar. Sie drückt formal die direkte Proportionalität zwischen der Kraft, die auf einen Körper wirkt, und der dadurch hervorgerufenen Beschleunigung des Körpers aus. Der Proportionalitätsfaktor entspricht dabei der Masse des Körpers.

$$F = m \cdot a \quad ^{37}$$

F ... Kraft, die auf den Körper wirkt, $[F] = \text{N}$

a ... Beschleunigung des Körpers, $[a] = \frac{\text{m}}{\text{s}^2}$

m ... Masse des Körpers, $[m] = \text{kg}$

³⁷ Die Newton'sche Bewegungsgleichung ist korrekterweise eine Vektorformel, $\vec{F} = m \cdot \vec{a}$ die aber erst in der Oberstufe als solche geschrieben wird. In der Unterstufe wird nur verbal beziehungsweise als grafische Pfeildarstellung über die Geschwindigkeit, Beschleunigung und Kraft gesprochen, mathematisch wird der Begriff des Vektors aber natürlich noch nicht verwendet, dies bleibt der fünften und sechsten Klasse vorbehalten.

3.1.2 Maßeinheiten

Maßeinheiten werden im Mathematikunterricht in der ersten Klasse eingeführt und sind natürlich bereits aus der Volksschule den SchülerInnen bekannt. Es ist für viele SchülerInnen oft nicht einsichtig, dass eine Größenangabe ohne Einheit im Grunde vollkommen sinnlos ist, da nur die Maßzahl alleine keinerlei Informationswert besitzt. Im Mathematikunterricht werden folgende Längenmaße³⁸ und Massenmaße³⁹ im Kapitel Dezimalzahlen eingeführt:

- Längenmaße: km, m, dm, cm, mm
- Massenmaße: t, kg, dag, g, dg, cg, mg

In einem eigenen Kapitel Zeitmessung werden dann die Zeitmaße⁴⁰ eingeführt:

- Zeitmaße: s, min, h, d, Monat, Jahr

Sodann gegen Ende des Schuljahres im Kapitel Rechteck und Quadrat die Flächenmaße⁴¹ und im Kapitel Quader und Würfel die Raummaße⁴²:

- Flächenmaße: mm², cm², dm², m², a, ha, km²
- Raummaße: m³, dm³, cm³, mm³ sowie die Hohlmaße hl, l, dl, cl, ml

In der ersten Klasse werden diese Maßeinheiten eingeübt und sollten ab dem ersten Unterrichtsjahr Physik genügend beherrscht werden, da dann sukzessive in Physik eine immer größere Zahl von Maßeinheiten entsprechend den neu hinzukommenden physikalischen Größen dazukommen. Im Physikunterricht erfahren die SchülerInnen dann in der sechsten Klasse die verschiedenen Einteilungen der Maßeinheiten und deren Zusammenhang.

▪ *SI-Basiseinheiten*⁴³

Im SI-System gibt es 7 Basisgrößen, denen jeweils eine Basiseinheit entspricht. Alle anderen Einheiten können aus diesen sieben Basiseinheiten abgeleitet werden.

³⁸ Siehe [7] das Kapitel „Längenmaße“ ab Seite 102.

³⁹ Siehe [7] das Kapitel „Massenmaße“ ab Seite 104.

⁴⁰ Siehe [7] das Kapitel „Zeitmaße, Zeitdauer und Zeitpunkt“ ab Seite 142.

⁴¹ Siehe [7] das Kapitel „Flächenmaße“ ab Seite 222.

⁴² Siehe [7] das Kapitel „Raummaße“ ab Seite 254.

⁴³ Frz. *Système international d'unités*; dieses Einheitensystem ist das verbreitetste System für physikalische Größen. Das SI-System ist (1) ein metrisches System, das heißt eine Basiseinheit ist das Meter, (2) ein dezimales System, das heißt die Einheiten können nur ganzzahlige Exponenten besitzen und (3) ein kohärentes System, das heißt, jede abgeleitete Einheit ist ein Produkt von Potenzen von Einheiten, ohne dass zusätzliche numerische Faktoren hinzukommen.

<i>Basisgröße</i>	<i>Basiseinheit</i>
Länge (l, x, s, r)	Meter (m)
Zeit (t)	Sekunde (s)
Masse (m)	Kilogramm (kg)
absolute Temperatur (T)	Kelvin (K)
Stoffmenge (n)	Mol (mol)
elektrische Stromstärke (I)	Ampere (A)
Lichtstärke (I_v)	Candela (cd)

Tabelle 2: Die 7 SI-Basisgrößen und SI-Basiseinheiten

- Einheiten können mit *Vorsilben* versehen sein⁴⁴:

<i>Vorsilbe</i>	<i>Abkürzung</i>	<i>Vielfaches der Einheit</i>
Yotta	Y	10^{24} (Quadrillion)
Zetta	Z	10^{21} (Trilliarde)
Exa	E	10^{18} (Trillion)
Peta	P	10^{15} (Billiarde)
Tera	T	10^{12} (Billion)
Giga	G	10^9 (Milliarde)
Mega	M	10^6 (Million)
Kilo	k	10^3 (Tausend)
Hekto	H	10^2 (Hundert)
Deka	da	10^1 (Zehn)
Dezi	d	10^{-1} (Zehntel)
Centi	c	10^{-2} (Hundertstel)
Milli	m	10^{-3} (Tausendstel)
Mikro	μ	10^{-6} (Millionstel)
Nano	n	10^{-9} (Milliardenstel)
Pico	p	10^{-12} (Billionstel)
Femto	f	10^{-15} (Billiardenstel)
Atto	a	10^{-18} (Trillionstel)
Zepto	z	10^{-21} (Trilliardenstel)
Yokto	Y	10^{-24} (Quadrillionstel)

Tabelle 3: Vorsilben

⁴⁴ Es soll nicht unerwähnt bleiben, dass im Englischen die Zählweise der Tausender ab einer Milliarde eine andere als im Deutschen ist, man spricht in diesen Ländern von der „kurzen Leiter“ des Zählsystems. Im englischen Sprachgebrauch wird sozusagen doppelt so schnell gezählt, da es keine Namen auf -arde gibt. Deshalb sind größere Zahlen im Englischen nicht sehr gebräuchlich, man spricht dann etwa von thousand million millions statt quadrillion.

Deutsch	Milliarde	Billion	Billiarde	Trillion	Trilliarde
Englisch	billion	Trillion	quadrillion	quintillion	sextillion

- Aus den Basiseinheiten können neue Einheiten gebildet werden, sogenannte *abgeleitete Einheiten*. Einheiten dürfen formal wie Variablen behandelt werden, um neue Einheiten zu bilden beziehungsweise als Einheitengleichungen angeschrieben zu werden. Das „Rechnen“ mit Einheiten ist eine überaus effektive Methode, da auf einfache Art wichtige Schlussfolgerungen gezogen werden können, wie an den folgenden Beispielen gezeigt werden soll.

Ich möchte im Besonderen darauf hinweisen, dass von Anfang an, wenn nicht gewichtige Gründe dagegensprechen oder andere als die SI-Einheiten gebräuchlich sind, darauf Wert gelegt werden sollte, alle Rechnungen in SI-Einheiten durchzuführen. Die Fehlerquote in den physikalischen Rechnungen wird dadurch sicherlich um einiges gesenkt und es ist im Endeffekt für die SchülerInnen einfacher, wenn sie es gewohnt sind, SI-Einheiten zu verwenden. Natürlich gibt es auch einige Einheiten, die einfach so gebräuchlich sind, dass die SchülerInnen lernen müssen, mit diesen umzugehen und Umrechnungen von SI-Einheiten in diese und umgekehrt durchführen zu können. Ein Beispiel wäre etwa das Joule als die SI-Einheit für die Arbeit und die gebräuchliche Einheit Kilowattstunde. Diese Einheit (kWh) ist keine SI-Einheit, ist aber zum Gebrauch vom SI zugelassen.

Es seien noch einige wichtige abgeleitete Einheiten (mit eigenen Namen) aufgelistet:

<i>Physikalische Größe</i>	<i>Name der Einheit</i>	<i>In SI-Basiseinheiten</i>
Kraft	Newton (N)	$N = \text{kg} \cdot \text{m} \cdot \text{s}^{-2}$
Druck	Pascal (Pa)	$\text{Pa} = \text{N} / \text{m}^2 = \text{kg} \cdot \text{m}^{-1} \cdot \text{s}^{-2}$
Arbeit, Energie	Joule (J)	$J = \text{N} \cdot \text{m} = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2}$
Leistung	Watt (W)	$W = J / s = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-3}$
Frequenz	Hertz (Hz)	$\text{Hz} = \text{s}^{-1}$
Ladung	Coulomb (C)	$C = \text{A} \cdot \text{s}$
Elektrische Spannung	Volt (V)	$V = J / C = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-3} \cdot \text{A}^{-1}$
Elektrischer Widerstand	Ohm (Ω)	$\Omega = V / A = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-3} \cdot \text{A}^{-2}$
Elektrische Leitfähigkeit	Siemens (S)	$S = \Omega^{-1} = \text{kg}^{-1} \cdot \text{m}^{-2} \cdot \text{s}^3 \cdot \text{A}^2$
Kapazität	Farad (F)	$F = C / V = \text{kg}^{-1} \cdot \text{m}^{-2} \cdot \text{s}^4 \cdot \text{A}^2$
Magnetfeld	Tesla (T)	$T = \text{N} / (\text{A} \cdot \text{m}) = \text{kg} \cdot \text{s}^{-2} \cdot \text{A}^{-1}$
Magnetischer Fluss	Weber (Wb)	$\text{Wb} = T \cdot \text{m}^2 = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2} \cdot \text{A}^{-1}$
Induktivität	Henry (H)	$H = J / A^2 = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2} \cdot \text{A}^{-2}$

Tabelle 4: Abgeleitete physikalische Größen

Beispiel 1: Dichte und elektrische Ladung in SI-Basiseinheiten

Wird in der Physik eine neue Einheit eingeführt, so erhält sie, falls ihr eine entsprechende Bedeutung zuerkannt wird, einen eigenen Namen, aber sie kann natürlich auch immer in den SI-Basiseinheiten ausgedrückt werden. Dies soll an den beiden einfachen Beispielen der Einheit für die Dichte und der elektrischen Ladung vorgeführt werden.

Die Definitionsgleichung für die Dichte lautet

$$\rho = \frac{m}{V}$$

m ... Masse, [m] = kg

V ... Volumen, [V] = m³

Also ergibt sich für die Einheit der Dichte⁴⁵:

$$[\rho] = \frac{\text{kg}}{\text{m}^3} = \text{kg} \cdot \text{m}^{-3}$$

Für die elektrische Ladung gilt folgende Beziehung (für eine konstante Stromstärke):

$$Q = I \cdot t$$

I ... elektrische Stromstärke, [I] = A

t ... Zeit, [t] = s

Also ergibt sich für die Einheit der elektrischen Ladung:

$$[Q] = \text{A} \cdot \text{s} = \text{C} \quad (\text{Coulomb})$$

Die SchülerInnen sollen dadurch erkennen, dass jede neue Einheit, sofern sie keine Basiseinheit darstellt, durch einfache mathematische Umformungen auf SI-

⁴⁵ Die Dimension einer Größe wird in dieser Arbeit immer so angeschrieben, dass die Basisgröße in eckige Klammern geschrieben wird. Es gibt im SI-System auch eigene Symbole für die Dimension der Basisgrößen, die ich aber nicht verwende, stattdessen werde ich auch bei einer Dimensionsangabe die Einheiten anschreiben.

SI-Basisgröße	Länge	Zeit	Masse	Temp.	Stoffmenge	Stromstärke	Lichtstärke
Symbol für die Dimension	L	T	M	Θ	N	I	J

In dieser Schreibweise ergibt sich für die Dimension der Dichte $[\rho] = \text{M} \cdot \text{L}^{-3}$ und für die Kraft $[F] = \text{M} \cdot \text{L} \cdot \text{T}^{-2}$.

Basiseinheiten zurückgeführt werden kann und diese Umrechnung in SI-Basiseinheiten auch durchführen können!

Beispiel 2: Einheit der Gravitationskonstanten

Oft kommen in der Physik in Formeln Konstanten vor, deren Einheiten man sich nicht merkt und auch nicht merken muss, da sie zusammengesetzte Einheiten sind und recht komplex sein können. Das Merken dieser Einheiten von Konstanten ist auch nicht notwendig, da man sie aus Formeln, in der diese Konstante vorkommt, sehr leicht bestimmen kann, wenn die Einheiten der anderen Größen in der Formel bekannt sind, was ja auch der Fall sein sollte. Als Beispiel hierzu soll die Einheit der Gravitationskonstanten im Newton'schen Gravitationsgesetz bestimmt werden!

Das Newton'sche Gravitationsgesetz lautet in skalarer Schreibweise: ⁴⁶

$$F = G \cdot \frac{m_1 \cdot m_2}{r^2}$$

$m_1 \cdot m_2$... die beiden sich anziehenden Massen, $[m_1] = [m_2] = \text{kg}$

r ... Abstand der beiden Massen; $[r] = \text{m}$

F ... Kraft zwischen den beiden Massen; $[F] = \text{kg} \cdot \text{m} \cdot \text{s}^{-2}$

Somit gilt für die Einheitengleichung⁴⁷:

$$\text{kg} \cdot \text{m} \cdot \text{s}^{-2} = [G] \cdot \text{kg}^2 \cdot \text{m}^{-2}$$

$$\Rightarrow [G] = \text{kg} \cdot \text{m} \cdot \text{s}^{-2} \cdot \text{kg}^{-2} \cdot \text{m}^2 = \text{m}^3 \cdot \text{kg}^{-1} \cdot \text{s}^{-2} \quad (\text{in Basiseinheiten})$$

$$\text{oder: } \text{N} = [G] \cdot \text{kg}^2 \cdot \text{m}^{-2} \Rightarrow [G] = \text{N} \cdot \text{m}^2 \cdot \text{kg}^{-2}$$

Die Einheit der Gravitationskonstanten ist also $\text{N} \cdot \text{m}^2 \cdot \text{kg}^{-2}$ in abgeleiteten Einheiten oder $\text{m}^3 \cdot \text{kg}^{-1} \cdot \text{s}^{-2}$ in Basiseinheiten, wobei dies natürlich äquivalente Darstellungen sind! Auf diese Weise kann die unbekannte Einheit einer jeden physikalischen Größe

⁴⁶ In vektorieller Schreibweise lautet das Gravitationsgesetz: $\vec{F}_{12} = G \cdot \frac{m_1 \cdot m_2}{r^3} \cdot \vec{r}_{12}$, wobei $\vec{r}_{12}$ der Vektor vom Massenschwerpunkt der Masse 1 zum Massenschwerpunkt der Masse 2 ist. $\vec{F}_{12}$ ist die Kraft, die von der Masse 2 auf die Masse 1 ausgeübt wird. Auf die Masse 2 wird dann laut dem 3. Newton'schen Gesetz die Gegenkraft $\vec{F}_{21} = -\vec{F}_{12}$ ausgeübt. Allerdings genügt es für dieses Beispiel, die skalare Schreibweise des Gravitationsgesetzes zu verwenden.

⁴⁷ Korrekter wäre die Bezeichnung *Dimensionsgleichung* statt *Einheitengleichung*, es ist aber für die SchülerInnen sicher leichter nachzuvollziehen, hier von einer Einheitengleichung zu sprechen.

oder einer Konstanten in einer Formel, durch einfache mathematische Umformungen bestimmt werden!

Beispiel 3: Falsifizieren von Gleichungen

Es werden in der Physik natürlich sehr oft neue Formeln hergeleitet, wobei die Herleitung selbst sehr kompliziert sein kann und es nicht ausgeschlossen werden kann, dass sich in der Herleitung ein Fehler eingeschlichen hat. Man kann zwar keine Formel durch mathematische Methoden verifizieren, aber doch zumindest in gewissen Fällen falsifizieren, indem man die in der Formel auftretenden Einheiten einer Plausibilitätsprüfung unterzieht.

Es soll angenommen werden, dass ein Physiker in einer theoretischen Abhandlung folgende Formel für eine Kraft hergeleitet hat und die Formel folgendermaßen lautet:

$$F = \rho \cdot v^2 \cdot r^3.$$

Die Frage ist nun, ob diese Formel korrekt sein kann (abgesehen von ihrem physikalischen Inhalt, der hier aber nicht zur Debatte steht). Die Einheiten der in der Formel auftretenden physikalischen Größen sind bekannt:

ρ ... Dichte; $[\rho] = \text{kg} \cdot \text{m}^{-3}$

v ... Geschwindigkeit; $[v] = \text{m} \cdot \text{s}^{-1}$

r ... Abstand; $[r] = \text{m}$

F ... Kraft; $[F] = \text{N} = \text{kg} \cdot \text{m} \cdot \text{s}^{-2}$

Es müssen nun bei jeder physikalischen Formel auch die beiden Seiten der entsprechenden Einheitengleichung übereinstimmen. Dazu muss die linke und die rechte Seite in Basiseinheiten angeschrieben und dann beide Seiten auf ihre Äquivalenz verglichen werden. Stimmen die beiden Seiten überein, so kann zwar ausgesagt werden, dass die Formel in Bezug auf die Einheiten plausibel ist, ob sie physikalisch korrekt ist, bleibt natürlich weiterhin unbeantwortet. Sollten die Einheiten aber nicht übereinstimmen, ist die Formel sicher nicht korrekt und der Physiker weiß, dass ihm bei der Herleitung ein Fehler unterlaufen sein muss.

linke Seite: $[F] = N = \text{kg} \cdot \text{m} \cdot \text{s}^{-2}$

rechte Seite: $[\rho] \cdot [v]^2 \cdot [r]^3 = \text{kg} \cdot \text{m}^{-3} \cdot (\text{m} \cdot \text{s}^{-1})^2 \cdot \text{m}^3 = \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2}$

Man sieht, dass $\text{kg} \cdot \text{m} \cdot \text{s}^{-1} \neq \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2}$ ist, also kann die Formel nicht stimmen!

Wenn die SchülerInnen mathematisch mit Einheitengleichungen umzugehen lernen, verfügen sie über ein sehr effektives Instrument, um durch einfachste mathematische Umformungen eine Gleichung zumindest in gewissen Fällen falsifizieren zu können. Es sei nur noch einmal betont, dass dies umgekehrt natürlich nicht gilt, falls die Einheitengleichung stimmen sollte, heißt das keineswegs, dass die Formel richtig ist!

Beispiel 4: Vervollständigung von Formeln

Ein Physiker hat durch Messungen festgestellt, dass eine Größe, deren Dimension der einer Kraft F entspricht, nur von der Geschwindigkeit v , der Masse m und einem Abstand r abhängen kann. Er findet durch Messungen folgenden Zusammenhang heraus, wobei einzig der unbekannte Exponent x noch nicht bestimmt werden konnte:

$$F = m \cdot v^2 \cdot r^x$$

Wie kann die Formel vervollständigt werden, das heißt der Exponent x bestimmt werden, ohne weitere Messungen durchführen zu müssen?

Wieder müssen die Einheiten links und rechts in der Einheitengleichung übereinstimmen:

$$\text{kg} \cdot \text{m} \cdot \text{s}^{-2} \stackrel{\text{muss}}{=} \text{kg} \cdot \text{m}^2 \cdot \text{s}^{-2} \cdot \text{m}^x$$

Das geht nur, wenn $x = -1$ ist, also muss die Formel lauten: $F = m \cdot v^2 \cdot r^{-1} = \frac{m \cdot v^2}{r}$

Durch diese vier Beispiele sollten die SchülerInnen lernen, dass allein mit einem geschickten Umgang mit Maßeinheiten und einfachsten mathematischen Kenntnissen, interessante und tiefreichende Erkenntnisse über physikalische Formeln gewonnen werden können. Die SchülerInnen sollen dadurch aber auch begreifen lernen, wie wichtig der korrekte Umgang mit Maßeinheiten in den Naturwissenschaften ist.

3.1.3 Statistik, Mittelwerte

Statistische Kenngrößen werden in der Unterstufe sukzessive eingeführt, in der ersten Klasse ist das der Lageparameter arithmetischer Mittelwert⁴⁸. Weitere Lageparameter werden erst in der vierten Klasse eingeführt, dies sind dann der Median, der Modus (Modalwert) und die Quantile⁴⁹. Die Standardabweichung als Streuungsmaß wird ebenfalls in der vierten Klasse eingeführt⁵⁰.

Im Physikunterricht bietet sich die Chance, für die SchülerInnen die Sinnhaftigkeit der Berechnung und Verwendung des arithmetischen Mittelwertes vor Augen zu führen. Zur Anwendung kommt der arithmetische Mittelwert genau genommen bei vielen Messungen, die in Folge von Versuchen durchgeführt werden, wobei natürlich oft aus Zeitmangel in der Praxis auf eine umfangreiche Messreihe verzichtet wird und nur ein einzelner Wert herangezogen wird.

Beispiel: Aufstellen einer Messreihe

Die Messreihe der Zeiten der Läufer in Beispiel 1 aus Kapitel 3.1.1 wurde nur verkürzt dargestellt. Tatsächlich waren es bei jedem Messpunkt bis zu drei SchülerInnen, die mit ihren Handys die Zeiten der LäuferInnen gemessen haben. Aus diesen drei Zeiten wurde dann jeweils der arithmetische Mittelwert gebildet, der dann zur Berechnung der Geschwindigkeit herangezogen worden ist. Da die Zeiten, die die SchülerInnen jeweils mit ihren Handys gemessen haben, immer unterschiedlich waren, konnte so der Messfehler verkleinert werden, was man auch am Ergebnis der doch recht konstanten Geschwindigkeiten der LäuferInnen sehen kann.

Für die LäuferIn 1 haben sich etwa folgende gemessenen Werte ergeben, wobei jetzt die vollständige Tabelle angegeben sei:

Gesamtweg s in m	Gemessene Zeit in s	Mittlere Zeit t in s	Geschwindigkeit: $v = s : t$
--------------------	---------------------	------------------------	------------------------------

⁴⁸ Siehe [7] das Kapitel „1. Mittelwert“ ab Seite 133.

⁴⁹ Siehe [10] das Kapitel „Median und Modus“ ab Seite 135 und das Kapitel „Quartile, Quantile und Boxplots“ ab Seite 139.

⁵⁰ Siehe [10] das Kapitel „Streuungsmaße“ ab Seite 144.

2	1,93; 2,11	2,0	1,0
4	3,7; 3,12; 4,26	3,7	1,1
6	6,11	6,1	1,0
8	8,23	9,2	0,9
10	10,10; 10,36; 9,9	10,1	1,0
12	12,2	12,2	1,0
14	14,14; 13,77	14,0	1,0
16	16,5	16,5	1,0
18	18,21	18,2	1,0

Tabelle 5: Vollständige Messwertetabelle des Versuchs zur Geschwindigkeitsbestimmung

Für den ersten Messpunkt haben zwei SchülerInnen die Zeit gestoppt und die beiden Messwerte 1,03 s und 2,11 s erhalten. Der arithmetische Mittelwert dieser zwei Werte ergibt sich zu $\frac{1,93+2,11}{2} = 2,02$ s. Da manche der SchülerInnen die gemessenen Werte auf eine, andere auf zwei Nachkommastellen genau angegeben haben, wurde bei allen Werten nur eine Nachkommastelle herangezogen. Deswegen wurde der arithmetische Mittelwert von 2,02 s auf 2,0 s gerundet. Für den zweiten Messpunkt haben dann drei SchülerInnen die Zeit gemessen und es ergibt sich hier der arithmetische Mittelwert von $\frac{3,7+3,12+4,26}{3} = 3,69\dot{3}$ s, was dann auf den Wert von 3,7 s gerundet wurde.

3.1.4 Quader, Würfel, Raummaße

Volumsberechnungen gehören mit zu den mathematischen Fertigkeiten, die die SchülerInnen bereits ab der ersten Klasse lernen müssen. Die Bestimmung von Rauminhalten zieht sich dann durch alle acht Klassenstufen:

- 1. Klasse: Quader, Würfel
- 2. Klasse: Prisma
- 3. Klasse: Pyramide
- 4. Klasse: Zylinder, Kegel, Kugel

Mit den Werkzeugen der Trigonometrie, Vektorrechnung und Integralrechnung lernen die SchülerInnen dann ab der fünften Klasse weitere Methoden zur Volumsberechnung kennen und in der achten Klasse sind die SchülerInnen mittels Integralrechnung imstande, Volumsberechnungen allgemeinerer Art durchzuführen. Das Volumen eines Körpers spielt in der Physik natürlich in vielen Formeln eine ganz wichtige Rolle und obwohl es recht einfach erscheinen mag, bereitet es doch manchen SchülerInnen gewisse Probleme, diese Größe korrekt anzuwenden. Die Beispiele, die ich hierfür ausgewählt habe, sollten bereits in der zweiten und dritten Klasse anwendbar sein.

Beispiel 1: Der Begriff der Dichte

Die Masse eines Körpers ist einer der grundlegendsten und auch schwierigsten Begriffe in der Physik, der bereits in der zweiten Klasse den SchülerInnen vorgestellt wird. Oft verbindet man mit dem Massebegriff die Vorstellung des „Materieinhalts“, den ein Körper besitzt. Man sollte sich aber klar darüber sein, dass damit der Begriff der Masse sehr missverständlich beschrieben wird. Masse ist eine grundlegende Eigenschaft jeglicher Materie und es entzieht sich unserem Vorstellungsvermögen, was denn die Masse wirklich ausmacht. Den SchülerInnen wird beigebracht, dass Masse als eine der Materie innewohnende Eigenschaft nur durch ihre Wirkungen, die für uns erkennbar und messbar sind, erkannt werden kann. Diese Wirkungen der Masse sind (1) die Trägheit eines Körpers und (2) die Eigenschaft, dass sich alle Körper gegenseitig anziehen. Diese Wirkungen scheinen vorrangig nichts miteinander zu tun zu haben und so gibt es auch die beiden Begriffe der trägen Masse und der schweren Masse eines

Körpers, die den unterschiedlichen Eigenschaften verschiedene Massebegriffe zuordnen. Heute wissen wir natürlich, dass zwischen der trägen und schweren Masse eines Körpers kein Unterschied gemacht werden muss, da sie nur zwei unterschiedliche Auswirkungen ein und derselben Eigenschaft eines Körpers sind, sodass der Begriff der Masse genügt.

Um nun zwei Körper bezüglich ihrer Massen vergleichen zu können, muss ein neuer Begriff eingeführt werden, der Begriff der Dichte. Man beobachtet, dass ein Körper, der ein Vielfaches des Volumens eines anderen Körpers hat, wobei beide Körper aus demselben Material bestehen, auch die Masse des einen Körpers dasselbe Vielfache der Masse des anderen Körpers ist⁵¹. Masse und Volumen sind also direkt proportional zueinander. Wird die Masse auf das Einheitsvolumen bezogen, so spricht man von der Dichte ρ .

Somit ergibt sich als Definitionsgleichung für die Dichte:

$$\rho = \frac{m}{V}$$

m ... Masse, $[m] = \text{kg}$

V ... Volumen, $[V] = \text{m}^3$

ρ ... Dichte, $[\rho] = \text{kg} \cdot \text{m}^{-3}$

Zur Dichteberechnung sind also die Masse und das Volumen zu bestimmen. Für gasförmige und flüssige Körper ist die Volumsbestimmung verhältnismäßig einfach, da das Volumen jede Form annehmen kann, allerdings ist die Volumsausdehnung besonders bei Gasen sehr temperatur- und druckabhängig. Bei Festkörpern ist die rechnerische Volumsbestimmung bei unregelmäßigen Körpern (man denke an Steine!) sehr schwierig beziehungsweise gar nicht möglich. Hier kann man sich helfen, dass man das Volumen durch Flüssigkeitsverdrängung bestimmt. Wird ein Körper unbekannten Volumens in ein Messglas mit einer Flüssigkeit gegeben, so verdrängt der Körper dasselbe Volumen an Flüssigkeit wie sein eigenes Volumen ausmacht und die Flüssigkeit steigt im Messglas, sodass die Differenz der End- und der Ausgangshöhe der Flüssigkeit im Messbehälter dem Volumen des Körpers entspricht.

⁵¹ Natürlich müssen auch die äußeren Bedingungen dieselben sind, insbesondere die Temperatur beider Körper und der Umgebungsdruck. Körper dehnen sich normalerweise, wobei es auch Ausnahmen gibt, bei Temperaturerhöhung aus, wodurch das Volumen vergrößert und demzufolge die Dichte des Körpers verkleinert wird.

Beispiel 2: Auftrieb

Jedem, der schon einmal einen Stein zu heben versucht hat, der sich unter Wasser befindet, dürfte aufgefallen sein, dass dies viel leichter zu bewerkstelligen ist als außerhalb des Wassers. Der Grund dafür liegt in der Auftriebskraft begründet, die jeder Körper in einer Flüssigkeit erfährt und die sein Gewicht in der Flüssigkeit gegenüber seinem Gewicht außerhalb der Flüssigkeit verringert.

Wie viel wiegt ein Eisenquader mit den Abmessungen $l = 15\text{ cm}$, $b = 8\text{ cm}$, $h = 10\text{ cm}$ in Wasser und wie viel Prozent macht der Gewichtsverlust gegenüber seinem Gewicht an der Luft aus?

$$\rho_{\text{Wasser}} = 1000\text{ kg} \cdot \text{m}^{-3}, \quad \rho_{\text{Eisen}} = 7800\text{ kg} \cdot \text{m}^{-3}$$

Die Erdbeschleunigung g soll mit $10\text{ m} \cdot \text{s}^{-2}$ angenommen werden.

Lösung:

Das Volumen des Eisenquaders beträgt

$$V = 0,15 \cdot 0,08 \cdot 0,1\text{ m}^3 = 0,0012\text{ m}^3$$

Das Gewicht des Eisenquaders in Luft beträgt

$$G_{\text{Luft}} = \rho_{\text{Eisen}} \cdot V \cdot g = 7800 \cdot 0,0012 \cdot 10\text{ N} = 93,6\text{ N}$$

Die Auftriebskraft beträgt dann

$$F_A = \rho_{\text{Wasser}} \cdot V \cdot g = 1000 \cdot 0,0012 \cdot 10\text{ N} = 12\text{ N}$$

Somit beträgt das Gewicht des Eisenquaders in Wasser

$$G_{\text{Wasser}} = G_{\text{Luft}} - F_A = 93,6\text{ N} - 12\text{ N} = 81,6\text{ N}.$$

81,6 N von 93,6 N entspricht einem Prozentanteil von etwa 87,2% und somit beträgt der Gewichtsverlust des Eisenquaders 12,8%! ⁵²

⁵² Unberücksichtigt bleibt der Auftrieb des Eisenquaders, den er in Luft erhält. Da die Dichte von Luft nur etwa $1,3\text{ kg pro m}^3$ beträgt und damit im Bereich von einem Zehntelprozent der Dichte von Wasser liegt, ist auch die Auftriebskraft in demselben Bereich. Die Auftriebskraft in Luft beträgt $F_A = \rho_{\text{Luft}} \cdot V \cdot g \approx 0,016\text{ N}$, was im Vergleich zur Auftriebskraft von 12 N in Wasser vernachlässigbar ist.

Beispiel 3: Eisberge

Ein sehr schönes Beispiel, bei dem das Volumen eine wichtige Rolle spielt, ist die Frage, wie viel Prozent eines Eisberges sich unterhalb und wie viel Prozent sich oberhalb der Wasseroberfläche befinden. Ich selber bringe dieses Beispiel immer in der dritten Klasse in Physik, wenn ich den Auftrieb behandle, denn es zeigt den SchülerInnen, dass sie mit den bis zu diesem Zeitpunkt erworbenen Kenntnissen bereits interessante Fragestellungen aus der Natur auch quantitativ beantworten können.

Folgende Tatsachen sind bekannt: Die Dichte von Meerwasser (Salzwasser) beträgt etwa $1025 \text{ kg} \cdot \text{m}^{-3}$, die Dichte von Eis beträgt etwa $920 \text{ kg} \cdot \text{m}^{-3}$.

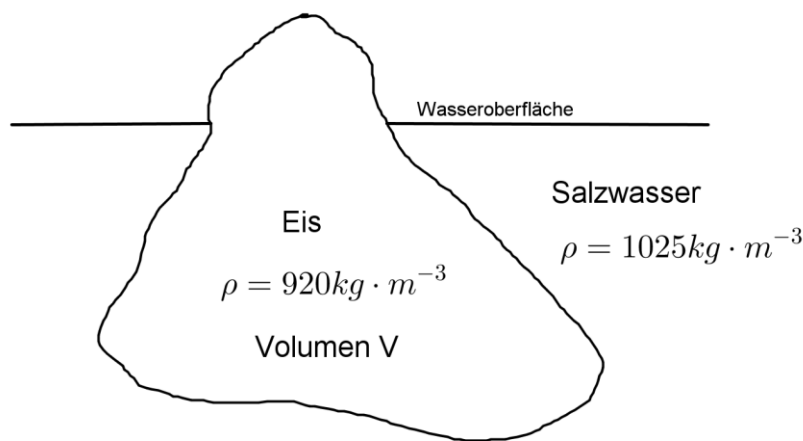


Abbildung 2: Ein Eisberg in Wasser

Es soll willkürlich angenommen werden, dass der Eisberg eine Masse von 5000 t hat.

Das Volumen des gesamten Eisberges ergibt sich dann aus der Dichte des Eises zu

$$V = \frac{m}{\rho}$$

$$V_{\text{Eisberg}} = \frac{5000000 \text{ kg}}{920 \text{ kg} \cdot \text{m}^{-3}} = 5435 \text{ m}^3$$

Das Gewicht des Eisberges beträgt

$$G = g \cdot m = 9,81 \cdot 5000000 \text{ N} = 49050000 \text{ N}$$

Gemäß dem archimedischen Prinzip gilt, dass ein Körper dann schwimmt, wenn sein Gewicht dem Gewicht des von ihm verdrängten Wassers entspricht. Also muss das

Gewicht des verdrängten Wassers 49050000 N betragen, was wiederum einer Masse von 5000000 kg entspricht.

Dieser Masse von 5000000 kg Meerwasser entspricht das Volumen

$$V_{\text{Meerwasser}} = \frac{5000000 \text{ kg}}{1025 \text{ kg} \cdot \text{m}^{-3}} = 4878 \text{ m}^3$$

Das heißt, dass sich von dem Volumen von 5435 m^3 des Eisbergs 4878 m^3 unter der Wasseroberfläche befinden müssen, damit der Eisberg schwimmt.

4878 m^3 von 5435 m^3 entsprechen einem Prozentanteil von 89,7%.

Somit ist gezeigt worden, dass sich etwa 90% des Volumens eines Eisbergs unter Wasser befinden und 10% über Wasser.

Im Übrigen ist damit auch ein Missverständnis aus der Welt geschafft, da viele Menschen glauben, der Grund, weshalb ein Eisberg schwimmt, sei, dass das Meerwasser aus Salzwasser besteht. Aber ein Eisberg würde genauso in Süßwasser schwimmen, wie man leicht nachprüfen kann, indem man Eiswürfel in ein Glas mit Leitungswasser gibt, die Eiswürfel schwimmen!

Würde man die Rechnung mit Süßwasser statt mit Salzwasser machen, müsste nur die Dichte des Wassers geändert werden, diese beträgt dann $1000 \text{ kg} \cdot \text{m}^{-3}$. Das Volumen des verdrängten Wassers beträgt dann

$$V_{\text{Süßwasser}} = \frac{5000000 \text{ kg}}{1000 \text{ kg} \cdot \text{m}^{-3}} = 5000 \text{ m}^3, \text{ was einem Prozentanteil von } 92,0\% \text{ entspricht. Das}$$

heißt es befinden sich im Süßwasser etwa 92% des Eisberges unter Wasser und etwa 8% über dem Wasser. Süßwasser statt Salzwasser macht nur einen Unterschied von 2% aus, was doch überraschend ist!

3.2 2. Klasse

3.2.1 Dreieck, Schwerpunkt, Flächeninhalt

Die SchülerInnen kennen den Begriff des Flächeninhalts bereits von der Volksschule her, zumindest den Flächeninhalt des Rechtecks und Quadrats. Die Bestimmung von Flächeninhalten zieht sich dann durch alle acht Jahre der AHS:

- 1. Klasse: Rechteck, Quadrat
- 2. Klasse: Rechtwinkeliges Dreieck
- 3. Klasse: Vierecke (Parallelogramm, Deltoid, Raute, Trapez, allgemeine Vierecke);
allgemeines Dreieck
- 4. Klasse: Kreis, Kreissektor, Kreisring

Mit den Werkzeugen der Trigonometrie, Vektorrechnung und Integralrechnung lernen die SchülerInnen dann ab der fünften Klasse weitere Methoden zur Flächenberechnung ebener Figuren kennen und in der achten Klasse sind die SchülerInnen mittels Integralrechnung imstande, jedwede Flächeninhalte von Figuren, die durch Funktionsgraphen begrenzt sind, zu berechnen und auch allgemeine Formeln für Flächeninhalte aufzustellen.

Beschränkt sich im Mathematikunterricht der Begriff des Flächeninhalts auf wirkliche geometrische „Flächen“, also Größen mit der Dimension m^2 beziehungsweise L^2 , so wird in der Physik der Begriff oft in allgemeinerer Weise verwendet, das heißt, der Flächeninhalt hat dann nicht mehr die Dimension einer Länge zum Quadrat, sondern die Dimension der Größe, die der x-Achse zugeordnet wird multipliziert mit der Dimension der Größe, die der y-Achse zugeordnet wird. Einige Beispiele mögen dies veranschaulichen und die Vorteile dieser Betrachtungsweise herausstreichen.

Beispiel 1: Flächeninhalte in Geschwindigkeits-Zeit-Diagrammen

Sei vorerst der einfachste Fall betrachtet, dass ein Körper sich mit der konstanten Geschwindigkeit v bewegt. Das Geschwindigkeits-Zeit-Diagramm sei für die zwei konstanten Geschwindigkeiten $1\text{ m}\cdot\text{s}^{-1}$ und $2\text{ m}\cdot\text{s}^{-1}$ dargestellt (Abbildung 3).

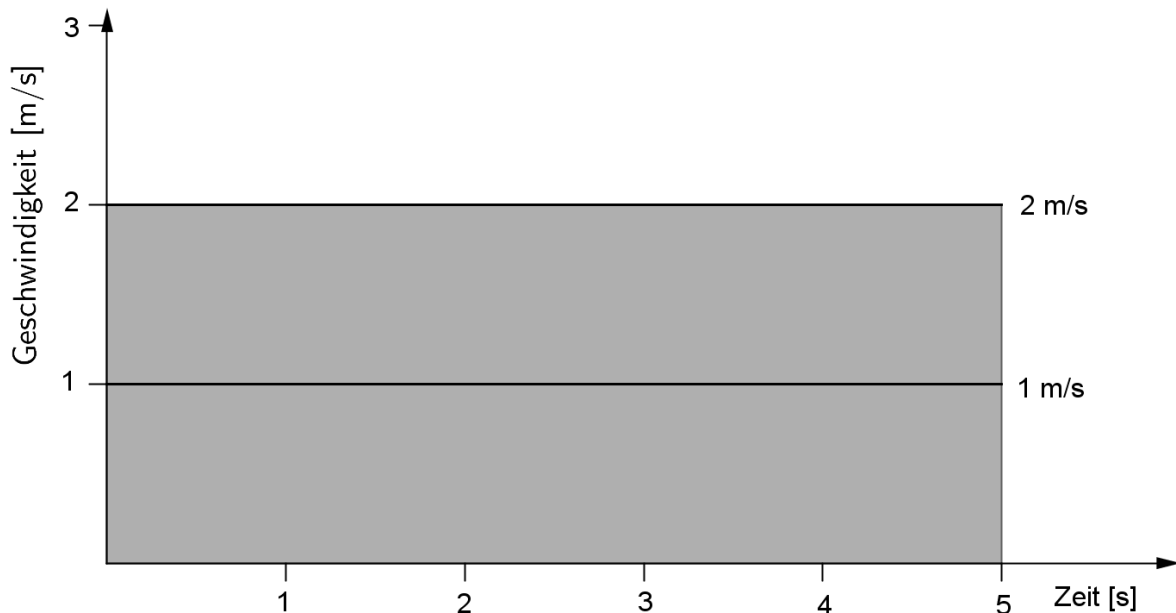


Abbildung 3: v-t-Diagramm eines gleichförmig bewegten Körpers

Für die Fragestellung, welchen Weg der Körper in 5 Sekunden zurücklegt, wenn er sich mit $1\text{ m}\cdot\text{s}^{-1}$ bzw. $2\text{ m}\cdot\text{s}^{-1}$ bewegt, ergibt sich $1\text{ m}\cdot\text{s}^{-1} \cdot 5\text{ s} = 5\text{ m}$ beziehungsweise $2\text{ m}\cdot\text{s}^{-1} \cdot 5\text{ s} = 10\text{ m}$. Man sieht natürlich sofort, dass die zurückgelegte Strecke (Geschwindigkeit mal Zeit) genau dem Flächeninhalt unter der Geraden für $1\text{ m}\cdot\text{s}^{-1}$ beziehungsweise $2\text{ m}\cdot\text{s}^{-1}$ entspricht.

Mit diesem Vorwissen ausgestattet, kann man nun mit den SchülerInnen das Problem der Fallstrecke eines Körpers im freien Fall behandeln. Unter der Annahme, dass die Geschwindigkeit eines frei fallenden Körpers (unter Vernachlässigung des Luftwiderstandes) linear mit der Zeit anwächst ($v = a \cdot t$), ergibt sich das Geschwindigkeits-Zeit-Diagramm in Abbildung 4.

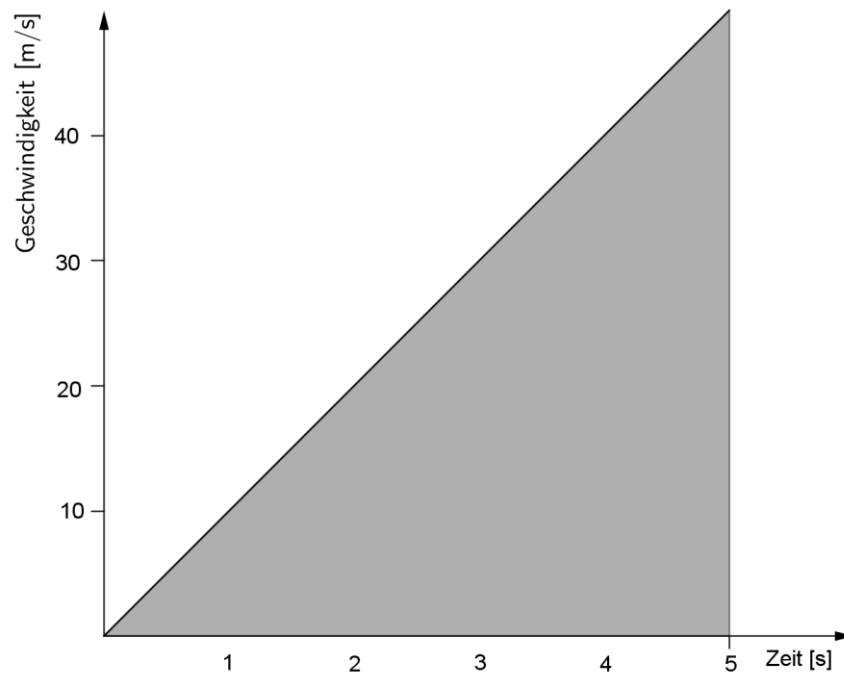


Abbildung 4: v-t-Diagramm des freien Falls

Die Fläche und somit die Falltiefe ergibt sich jetzt als einfache Dreiecksfläche:

$$s = \frac{v \cdot t}{2} = \frac{(a \cdot t) \cdot t}{2} = \frac{a}{2} \cdot t^2$$

In der achten Klasse kann dieses Problem dann allgemein mit Hilfe der Integralrechnung behandelt werden. Man kann aber bereits in der sechsten Klasse die Idee der Infinitesimalrechnung vorwegnehmen, indem man die Dreiecksfläche in Rechtecke zerlegt und so den SchülerInnen den Gedanken einer immer genaueren Annäherung durch Summation und Zerlegung und anschließender Grenzwertbildung nahebringt.

Beispiel 2: Spannungsarbeit und potenzielle Energie

Wird eine Feder gedehnt, so ist dazu eine Spannungsarbeit W_s notwendig. Diese Arbeit wird als potenzielle Energie E_p in der Feder gespeichert und kann wieder genutzt werden, etwa um eine mechanische Armbanduhr am Laufen zu halten. Um eine Feder um eine Strecke x zu dehnen, ist eine Kraft notwendig, die proportional zur Ausdehnung ist (Hooke'sches Gesetz). Der Zusammenhang lautet somit $F = k \cdot x$,

wobei k die sogenannte Federkonstante, eine charakteristische Größe der Feder ist⁵³. Der Zusammenhang zwischen der Kraft F und der Auslenkung der Feder ist in der Abbildung 5 dargestellt. Wieder kann dem Flächeninhalt eine physikalische Größe zugeordnet werden, in diesem Fall die potenzielle Energie E_p der Feder.

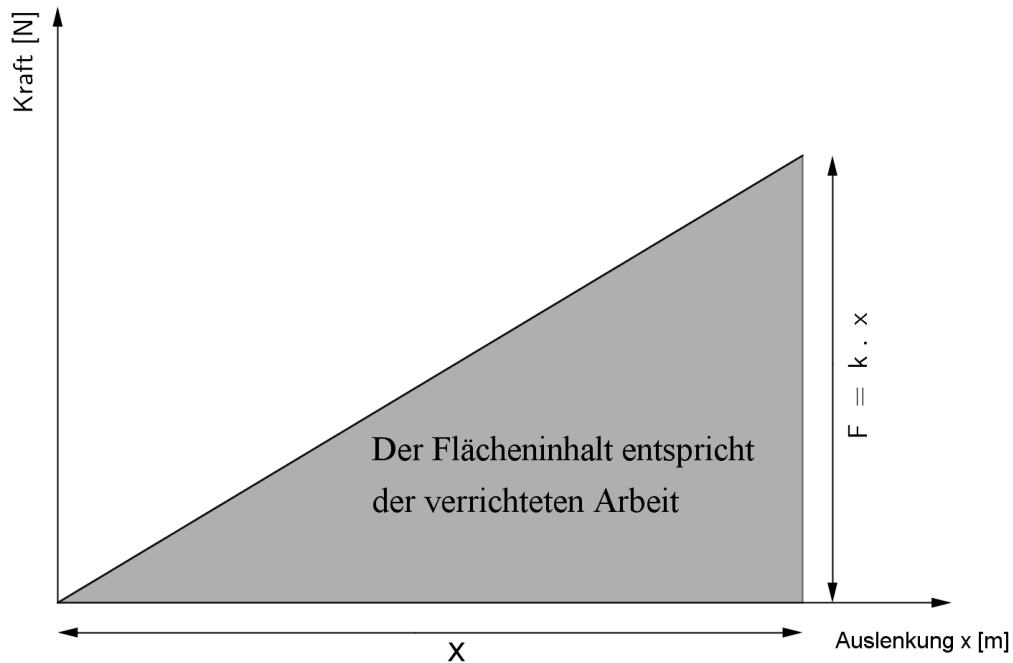


Abbildung 5: Spannungsarbeit einer gespannten Feder

$$W_s = \text{Flächeninhalt des Dreiecks} = \frac{x \cdot (k \cdot x)}{2} = \frac{k}{2} \cdot x^2$$

Die Arbeit, die zum Spannen der Feder aufgewendet werden muss, wird dann als potenzielle Energie in der Feder gespeichert, sodass man als Formel für die gespeicherte potenzielle Energie einer gespannten Feder folgenden Ausdruck erhält⁵⁴:

$$E_p = \frac{k}{2} \cdot x^2$$

⁵³ $F = k \cdot x$ ist die Kraft, die zum Spannen der Feder aufgewendet werden muss. Das Hooke'sche Gesetz lautet $F = -k \cdot x$, welches die Kraft beschreibt, die die Feder einer Auslenkung entgegensetzt. Das negative Vorzeichen bedeutet, wenn die Auslenkung in positive x -Richtung zeigt, dass die Kraft der Feder dann in die negative x -Richtung gerichtet ist.

⁵⁴ Zur Behandlung dieses Problems mittels Integralrechnung siehe das entsprechende Kapitel in der achten Klasse und die Beispiele dort.

Beispiel 3: Zweites Kepler'sches Gesetz

Ein besonders schönes Beispiel für die Anwendung von Flächeninhalten in der Physik ist das zweite Kepler'sche Gesetz. Es besagt, dass der Radiusvektor eines Planeten in gleichen Zeiten gleiche Flächen überstreicht (Abbildung 6).

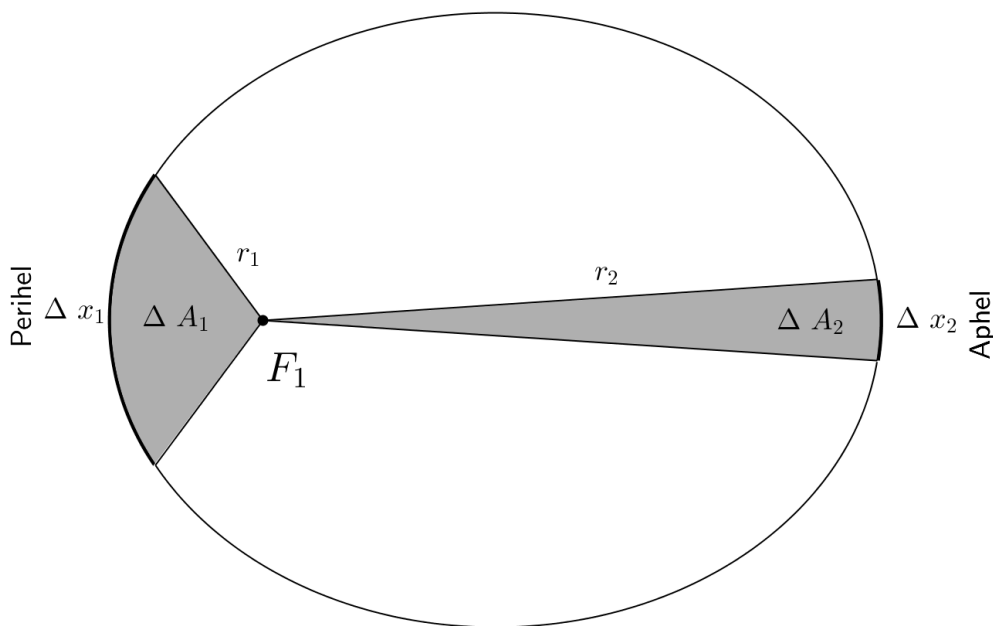


Abbildung 6: Zweites Kepler'sches Gesetz

Johannes Kepler⁵⁵ hat dieses nach ihm benannte Gesetz empirisch aus den beobachteten Bahndaten der Planeten, insbesondere des Mars, gewonnen, die der dänische Astronom Tycho Brahe⁵⁶ aufgezeichnet hat. Aufgrund dieser Bahndaten konnte Kepler zuerst für den Mars nachweisen, dass sich dieser auf einer Ellipsenbahn um die Sonne bewegt, später ist ihm das auch für die anderen damals bekannten Planeten gelungen. Aus den Bahndaten konnte Kepler sogar die Geschwindigkeiten der Planeten genau berechnen und so das nach ihm benannte 2. Kepler'sche Gesetz aufstellen.

Mit Hilfe des Drehimpulserhaltungssatzes, der besagt, dass der Drehimpuls der Planetenbewegung eine Erhaltungsgröße ist, das heißt zeitlich konstant bleibt, kann

⁵⁵ Johannes Kepler, deutscher Astronom, Mathematiker, Theologe; 1571 - 1630

⁵⁶ Tycho Brahe, dänischer Astronom; 1546 - 1601

dieses Gesetz auf einfache Weise bereits in der sechsten Klasse Oberstufe hergeleitet werden⁵⁷.

Aus der Erhaltung des Drehimpulses⁵⁸ folgt

$$L = m \cdot v \cdot r = m \cdot \frac{\Delta s}{\Delta t} \cdot r = \text{konstant}$$

also gilt auch

$$\frac{L}{m} = \frac{\Delta s}{\Delta t} \cdot r = \text{konstant} \quad (1)$$

Weiters gilt für die überstrichene Fläche

$$\Delta A = \frac{r \cdot \Delta s}{2} \Leftrightarrow \frac{\Delta A}{\Delta t} = \frac{1}{2} \cdot \frac{\Delta s}{\Delta t} \cdot r \quad (2)$$

(1) in (2) eingesetzt ergibt:

$$\frac{\Delta A}{\Delta t} = \frac{L}{2 \cdot m} = \text{konstant}$$

Somit gilt für beliebige Flächen ΔA_1 und ΔA_2 , die vom Radiusvektor des Planeten in den Zeiten Δt_1 und Δt_2 überstrichen werden:

$\frac{\Delta A_1}{\Delta t_1} = \frac{\Delta A_2}{\Delta t_2} \quad (\text{Zweites Kepler'sches Gesetz})$
--

Daraus folgt im Allgemeinen, dass die Bahngeschwindigkeit eines Planeten in Sonnennähe größer ist als in Sonnenferne, im Besonderen, dass ein Planet im Perihel (sonnennächster Punkt) die größte Bahngeschwindigkeit und im Aphel (sonnenfernster Punkt) die kleinste Bahngeschwindigkeit innerhalb eines Jahres aufweist.

⁵⁷ Es werden natürlich einige Vereinfachungen gemacht, es werden im Besonderen folgende Voraussetzungen angenommen: (1) Die Sonne wird als ruhend angenommen, ihre Bewegung um den gemeinsamen Massenmittelpunkt, der immer in der Sonne liegt, wird vernachlässigt. (2) Die gravitative Wechselwirkung des betrachteten Planeten und der restlichen Planeten des Sonnensystems wird vernachlässigt. (3) Da die Eigenrotation der Planeten nahezu unverändert bleibt, wird der Drehimpuls der Eigenrotation vernachlässigt. (4) Da die Bewegung der Sonne in Bezug auf den Massenmittelpunkt des Sonnensystems sehr klein ist, kann angenommen werden, dass der Gesamtdrehimpuls ausschließlich durch den Planeten zustande kommt.

⁵⁸ Da für die Herleitung des zweiten Kepler'schen Gesetzes die Richtung des Drehimpulses keine Rolle spielt, genügt es, nur den Betrag des Drehimpulses in der Herleitung zu betrachten. Die Erhaltung der Richtung des Drehimpulses ist dafür verantwortlich, dass sich die Planeten in einer Ebene um die Sonne bewegen.

Für die Erde etwa beträgt der Abstand von der Sonne im Perihel 147,1 Millionen Kilometer, die Entfernung im Aphel 152,1 Millionen Kilometer. Der Aphel wird am 4. Juli durchlaufen, das heißt, dass sich die Erde interessanterweise im Sommer am weitesten von der Sonne entfernt befindet und im Winter am nächsten.

3.3 3. Klasse

3.3.1 Potenzschreibweise

Die Zehnerpotenzschreibweise wird bereits in der dritten Klasse eingeführt⁵⁹, sodass die Schülerinnen eigentlich dann in der Oberstufe in Physik mit dieser Schreibweise vertraut sein sollten. Es ist aber doch erstaunlich, dass viele SchülerInnen trotzdem große Probleme haben, mit dieser Schreibweise umzugehen. Da aber in der Physik ein sicherer Umgang mit Potenzen unumgänglich ist, sollte so oft als möglich diese Schreibweise geübt und angewandt werden.

Auch haben viele SchülerInnen kein Gefühl dafür, was diese Zehnerpotenzen von der Größenordnung her für eine Bedeutung haben. Es fehlt ein intuitives Verständnis dafür, dass die Erhöhung des Exponenten um eine Einheit in einer Zehnerpotenz dasselbe bedeutet wie eine Multiplikation mit zehn. Die Beispiele in diesem Kapitel können teils schon in der dritten Klasse gebracht werden, sollten aber unbedingt in der Oberstufe am Beginn des Physikunterrichtes noch einmal geübt werden.

Beispiel 1: Lichtjahr in Meter

Ein einfaches Standardbeispiel ist die Aufgabe, wie viele Meter ein Lichtjahr ist. Hier ist natürlich unbedingt die Schreibweise mit Zehnerpotenzen heranzuziehen und es sollte auch darauf geachtet werden, dass nicht sofort die ganze Rechnung in den Taschenrechner getippt wird, sondern auch „von Hand“ mit den Zehnerpotenzen gerechnet wird.

1 Lichtjahr (1Lj) ist die Strecke, die das Licht in einem Jahr⁶⁰ zurücklegt. Es soll hier mit einem Jahr zu 365 Tagen gerechnet werden und die Lichtgeschwindigkeit mit $300000 \text{ km} \cdot \text{s}^{-1}$ angenommen werden.

$$1\text{Lj} = 3 \cdot 10^8 \cdot 3600 \cdot 24 \cdot 365 = 3 \cdot 10^8 \cdot 3,1536 \cdot 10^7 = 9,46 \cdot 10^{15} \text{ m}$$

⁵⁹ Siehe [9] das Kapitel „3.2 Potenzschreibweise“ ab Seite 73.

⁶⁰ Es gibt für die Zeiteinheit Jahr mehrere unterschiedliche Definitionen, so gibt es etwa das tropische Jahr, das siderische Jahr oder das anomalistische Jahr, siehe etwa [24] Seite 14. Um zu einer exakten Definition des Lichtjahres zu kommen, wird ein Jahr mit genau 365,25 Tagen angenommen. Da die Lichtgeschwindigkeit mit $299\,792\,458 \text{ m} \cdot \text{s}^{-1}$ als exakter Wert definiert ist, ist somit auch das Lichtjahr als Längeneinheit exakt definiert.

Beispiel 2: Anzahl der Atome im Universum

Es soll abgeschätzt werden, wie viele Atome es im Universum gibt.

Da das Universum zu 92% aus Wasserstoff besteht, soll als Vereinfachung die Anzahl der Wasserstoffatome geschätzt werden.

Die Masse der Sonne beträgt etwa $2 \cdot 10^{30}$ kg. 1 mol Wasserstoff entspricht $1 \text{ g} = 10^{-3} \text{ kg}$ und enthält etwa $6 \cdot 10^{23}$ Atome Wasserstoff. Damit enthält 1 kg Wasserstoff etwa $6 \cdot 10^{26}$ Atome Wasserstoff und die Sonne mit $2 \cdot 10^{30}$ kg Wasserstoff enthält dann $6 \cdot 10^{26} \cdot 2 \cdot 10^{30} \approx 10^{57}$ Atome.

Unsere Milchstraße enthält wiederum ungefähr 100 Milliarden Sterne, also enthält unsere Milchstraße somit $10^{57} \cdot 10^{11} = 10^{68}$ Atome. Nimmt man weiters an, dass es etwa 100 Milliarden Galaxien im Universum gibt, erhöht sich die Anzahl auf $10^{68} \cdot 10^{11} = 10^{79}$ Atome.⁶¹

Beispiel 3: Energieerzeugung der Sonne

Es sollen folgende Fragen beantwortet werden, wobei die Kenntnis der Solarkonstanten als bekannt vorausgesetzt wird:

- (1) Wie viel Energie erzeugt die Sonne mindestens pro Sekunde?
- (2) Wie viel Masse verliert die Sonne in 1 Sekunde durch die Energieabstrahlung?
- (3) Wie lange dauert es, bis die Sonne die Masse, die der Masse der Erde entspricht, in Energie umgewandelt hat?
- (4) Die Sonne hat ein recht gut bestimmtes Alter von 4,6 Milliarden Jahren. Angenommen, die Sonne erzeugt in diesen Jahren immer dieselbe Energiemenge pro Sekunde, wie viel Prozent ihrer ursprünglichen Masse hat die Sonne in dieser Zeit in Energie verwandelt? Ist dieser Massenverlust der Sonne relevant?

⁶¹ Diese Abschätzung ist zwangsläufig mit einer sehr großen Fehlerquote behaftet. In der Literatur findet man Abschätzungen von 10^{78} bis 10^{84} Atomen. Man erkennt aber deutlich, wie sehr man oft Zehnerpotenzen falsch einschätzt und das ist auch der eigentliche Sinn dieser Aufgabe, dass die SchülerInnen begreifen, welche immens großen Zahlen durch die Potenzschreibweise ausgedrückt werden können.

Lösungen:

- (1) Die Solarkonstante, das ist die Energieeinstrahlung pro m^2 , beträgt auf der Erde $1,37 \text{ kW} \cdot \text{m}^{-2}$ und wird durch Messungen bestimmt. Das bedeutet, es werden von der Sonne 1370 J in jeder Sekunde auf jeden Quadratmeter einer Kugelschale abgestrahlt, die den Radius des Abstands von der Erde zur Sonne hat. Dieser Abstand beträgt 150 Millionen Kilometer, also $R = 1,5 \cdot 10^{11} \text{ m}$. Damit ergibt sich die Energie, die die Sonne in einer Sekunde erzeugen muss zu

$$E = 1370 \cdot 4 \cdot R^2 \cdot \pi = 1370 \cdot 4 \cdot (1,5 \cdot 10^{11})^2 \pi \approx 3,87 \cdot 10^{26} \text{ J}$$

Die Sonne erzeugt also pro Sekunde mindestens $3,87 \cdot 10^{26} \text{ J}$ an Energie!⁶²

- (2) Nach Einstein sind Energie und Masse äquivalent und die Beziehung wird durch die berühmte Einsteingleichung $E = m \cdot c^2$ ausgedrückt, wobei E die Energie und m die dazu äquivalente Masse ist, c ist die Lichtgeschwindigkeit.

$$m = \frac{E}{c^2} = \frac{3,87 \cdot 10^{26}}{(3 \cdot 10^8)^2} \approx 4,3 \cdot 10^9 \text{ kg}$$

Es werden pro Sekunde mindestens 4,3 Millionen Tonnen Masse in Energie umgewandelt!⁶³

- (3) Die Masse der Erde beträgt $6 \cdot 10^{24} \text{ kg}$, also dauert es etwa $\frac{6 \cdot 10^{24}}{4,3 \cdot 10^9} \approx 1,4 \cdot 10^{15} \text{ s}$, bis eine Masse, die der Masse der Erde entspricht in Energie umgewandelt wird.

In Jahre umgerechnet sind das $\frac{1,4 \cdot 10^{15}}{3600 \cdot 24 \cdot 365} \approx 4,5 \cdot 10^7$ Jahre.

Die Sonne wandelt somit alle 45 Millionen Jahren eine Masse in Energie um, die der Masse der Erde entspricht!

- (4) Wir wissen aus (2), dass pro Sekunde $4,3 \cdot 10^9 \text{ kg}$ Masse in Energie umgewandelt wird. Der Massenverlust in 4,6 Milliarden Jahren, das sind $4,6 \cdot 10^9$ Jahre, beträgt dann $4,3 \cdot 10^9 \cdot 4,6 \cdot 10^9 \cdot 3600 \cdot 24 \cdot 365 \approx 6 \cdot 10^{26} \text{ kg}$.

⁶² Diese Größe $3,87 \cdot 10^{26} \text{ J} \cdot \text{s}^{-1} = 3,87 \cdot 10^{26} \text{ W}$ ist eine der grundlegenden charakteristischen Größen der Sonne und entsprechend natürlich aller Sterne und wird in der Astronomie als die Leuchtkraft eines Sterns bezeichnet.

⁶³ Diese Größe wird als Massenverlustrate der Sonne bezeichnet.

Die Masse der Sonne beträgt etwa $2 \cdot 10^{30}$ kg, somit wurden in der Lebenszeit der Sonne

$$\frac{6 \cdot 10^{26}}{2 \cdot 10^{30}} \cdot 100\% \approx 3 \cdot 10^{-2}\% = 0,03\%$$

Ihrer Masse in Energie umgewandelt und abgestrahlt. Man sieht, dass der Massenverlust nicht relevant ist!

3.4 4. Klasse

3.4.1 Lehrsatz des Pythagoras

Der Lehrsatz des Pythagoras⁶⁴ findet sich sowohl im Stoff der dritten Klasse⁶⁵ als auch der vierten Klasse⁶⁶. Ich habe in den Schulbüchern für Mathematik als Anwendungen des Satzes des Pythagoras nur rein geometrische Aufgaben gefunden, aber keine einzige aus der Physik. Man findet den Satz des Pythagoras in beiden Schulbüchern der dritten und vierten Klasse in derselben Formulierung wie folgt:

Satz des Pythagoras 1:⁶⁷

Für die **Seitenlängen** jedes **rechtwinkligen Dreiecks** mit den Katheten a und b und der Hypotenuse c gilt: $a^2 + b^2 = c^2$

Das ist natürlich nur die halbe Wahrheit, denn der Satz des Pythagoras gilt auch in der umgekehrten Richtung. Das heißt man darf und sollte auch formulieren:

Satz des Pythagoras 2:

Gilt in einem Dreieck mit den Seiten a, b und c die Beziehung $a^2 + b^2 = c^2$, so ist dieses Dreieck ein rechtwinkeliges (wobei dann die Seiten a und b die Katheten und die Seite c die Hypotenuse ist).

Dient der Satz des Pythagoras in der Formulierung 1 zum Beispiel zum Berechnen von Beträgen von Vektoren oder tritt in der Physik bei Berechnungen von Größen auf, die normal aufeinander stehen, so ist die Anwendung des Satzes des Pythagoras in der Formulierung 2 sicher nicht so häufig anzutreffen, obwohl sich gerade in dieser Formulierung interessante Aussagen treffen lassen. Damit können sehr interessante physikalische Erscheinungen in mathematisch einfacher Weise gezeigt werden, wie ich im folgenden Beispiel 2 aufzeigen

⁶⁴ Pythagoras von Samos, griechischer Philosoph und Mathematiker; um 570 v. Chr. – nach 510 v. Chr.

⁶⁵ Siehe [9] das Kapitel „D Lehrsatz des Pythagoras“ ab Seite 222.

⁶⁶ Siehe [10] das Kapitel „B Lehrsatz des Pythagoras“ ab Seite 156.

⁶⁷ Siehe [9] Seite 224 und [10] Seite 158

möchte. Oft ist natürlich der Einsatz der Winkelfunktionen einfacher, aber falls diese noch nicht zur Verfügung stehen oder wenn einfach einmal eine andere Lösungsmöglichkeit aufgezeigt werden soll, leistet der Satz des Pythagoras sicherlich gute Dienste. Von Vorteil für die SchülerInnen ist es auch, dass der Satz des Pythagoras von den SchülerInnen im Allgemeinen gekannt wird, was bei den Winkelfunktionen meiner Erfahrung nach nicht in demselben Ausmaß der Fall ist.

Beispiel 1: Flussüberquerung

Ein Boot möchte einen Fluss überqueren und am gegenüberliegenden Ufer ankommen. Der Fluss sei 500 m breit ist und habe eine gewisse nicht näher bekannte Fließgeschwindigkeit. Das Boot wird natürlich abgetrieben und kommt ein Stück unterhalb des anvisierten Zieles am anderen Ufer an. Die gemessene Geschwindigkeit des Bootes betrage $5\text{ m}\cdot\text{s}^{-1}$ und die Zeit, die es zum Überqueren des Flusses benötigt, betrage 2,5 Minuten.

Wie viele Meter oberhalb des Flusses muss das Boot ablegen, um das gegenüberliegende Ufer genau gegenüber der Ablegestelle zu erreichen?

Lösung:

Die Situation und die verwendeten Größen sind in der Abbildung 7 dargestellt.

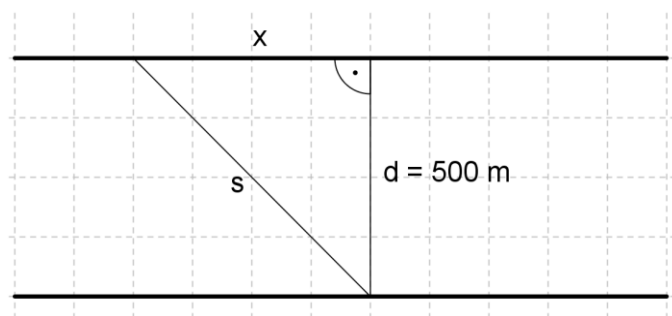


Abbildung 7: Flussüberquerung

Die vom Boot zurückgelegte Strecke ergibt sich durch $s = v \cdot t = 5 \cdot 150\text{ m} = 750\text{ m}$. Jetzt wird der Satz des Pythagoras angewandt, da wir es hier ja mit einem rechtwinkligen Dreieck zu tun haben.

$$s^2 = x^2 + d^2$$

Daraus ergibt sich:

$$x^2 = s^2 - d^2 \Rightarrow x = \sqrt{750^2 - 500^2} \approx 559 \text{ m}$$

Das Boot muss also 559 m oberhalb der Ablegestelle losfahren!

War das erste Beispiel eine Anwendung, wie mit dem Satz des Pythagoras in der Physik konkrete Zahlenwerte ausgerechnet werden können, so ist das zweite Beispiel eine Anwendung, um eine allgemeine physikalische Aussage treffen zu können.

Beispiel 2: Billardspiel

Beim Billardspiel stoße eine Kugel nichtzentral eine andere ruhende Kugel. Unter welchem Winkel bewegen sich die beiden Kugeln nach dem Stoß voneinander weg?

Lösung:

Die Situation vor und nach dem Stoß ist in der Abbildung 8 dargestellt:

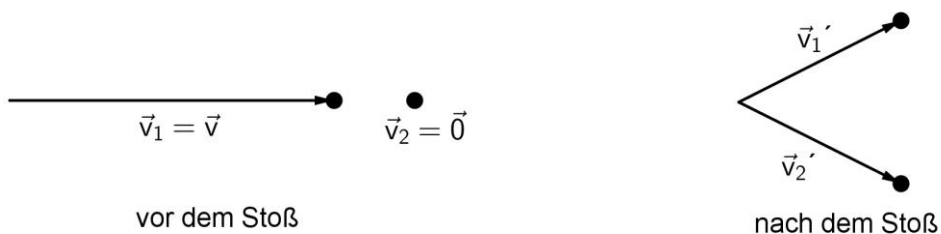


Abbildung 8: Geschwindigkeiten vor und nach dem Stoß beim Billardspiel

Aus dem Impulserhaltungssatz folgt

$$m \cdot \vec{v} = m \cdot \vec{v}_1' + m \cdot \vec{v}_2' \Rightarrow \vec{v} = \vec{v}_1' + \vec{v}_2'$$



Abbildung 9: Billardspiel, Zusammenhang der Geschwindigkeiten

Die Geschwindigkeitsvektoren werden vektoriell addiert und bilden somit ein Dreieck. Dies ist in der Abbildung 9 dargestellt, wobei das erste Bild der Abbildung die Geschwindigkeitsvektoren und das zweite Bild die Beträge der Geschwindigkeiten darstellen.

Nun ergibt sich aus dem Energieerhaltungssatz für die kinetischen Energien:

$$\frac{m \cdot v^2}{2} = \frac{m \cdot v_1'^2}{2} + \frac{m \cdot v_2'^2}{2} \Rightarrow v^2 = v_1'^2 + v_2'^2$$

Und jetzt kommt der Satz des Pythagoras in der Form 2 ins Spiel. In dem Dreieck, das aus den drei Geschwindigkeitsvektoren $\vec{v}$, $\vec{v}_1'$, $\vec{v}_2'$ gebildet wird, das heißt die Seitenlängen v , v_1' , v_2' hat, gilt die Beziehung $v^2 = v_1'^2 + v_2'^2$. Der Satz des Pythagoras besagt ja nun, dass in einem Dreieck, in dem die drei Seiten eben dieser Beziehung genügen, ein rechtwinkeliges Dreieck sein muss. Also muss der Winkel zwischen den beiden Vektoren $\vec{v}_1'$ und $\vec{v}_2'$ ein rechter Winkel sein. Da die Bewegung der Kugeln nach dem Stoß eine gerade Bewegung ist und in Richtung der Geschwindigkeitsvektoren gerichtet ist, ist somit gezeigt, dass sich die beiden Kugeln nach dem Stoß rechtwinkelig voneinander entfernen.

3.4.2 Graphen

Den ersten Kontakt mit dem Zeichnen und Lesen von Graphen machen die SchülerInnen in der zweiten Klasse im Zuge der Einführung des Koordinatensystems und bei direkten und indirekten Proportionen. In der vierten Klasse wird die Verwendung von Graphen weitergeführt bei der Einführung des Funktionenbegriffes, wobei dann im weiteren Verlauf in den folgenden Klassen mit einem immer größer werdenden Repertoire an Funktionen natürlich auch ein immer tieferes Verständnis des Funktionenverlaufs auch anhand der Darstellung der Funktion als Graph einhergeht. Mit Einführung der Differentialrechnung erreicht dann die Untersuchung des Funktionenverlaufs ihren Höhepunkt in der Schulmathematik durch Bestimmung charakteristischer Stellen wie etwa Extremstellen, Wendestellen und Bestimmung der Steigung der Funktion an beliebigen Stellen der Funktion.

Es soll hier aber um die Darstellung von Funktionen und ihre Anwendung in der Physik gehen, sodass die Differentialrechnung unberücksichtigt bleibt, aber stattdessen ganz kurz auf die nichtlineare Regression eingegangen werden soll.

Beispiel 1: s-t- und v-t-Diagramme in der Bewegungslehre

Im ersten Jahr Physikunterricht in der zweiten Klasse beginnt man am Anfang üblicherweise mit der Mechanik und innerhalb dieser mit der Geschwindigkeit und der Beschleunigung. Die Berechnung von Geschwindigkeiten und Beschleunigungen beschränkt sich auf einfachste Rechnungen, wobei es ja auch nicht das Ziel ist und sein kann, in diesem frühen Stadium, wenn die Kinder gerade mit dem neuen Fach Physik Bekanntschaft schließen, das Hauptaugenmerk auf die mathematische Behandlung dieser Größen zu legen. Vielmehr ist es wichtig, dass die SchülerInnen ein Gefühl dafür bekommen, was die verschiedenen Bewegungsarten bedeuten und wie man diese graphisch darstellen kann. Aus dem Alltagsleben sind den SchülerInnen natürlich verschiedenste Bewegungsabläufe bekannt, etwa die gleichförmige geradlinige Bewegung in einem Auto oder Flugzeug oder dem Fahrrad. Sie kennen die Effekte, die auftreten, wenn die Geschwindigkeit plötzlich abnimmt oder zunimmt aus ihren Erfahrungen in der U-Bahn oder Straßenbahn. Auf diesen Alltagserfahrungen gilt es aufzubauen und so dieses vorhandene Wissen und die Erfahrungen aus der Lebenswelt

der Kinder zu nutzen, um dadurch den SchülerInnen die Physik der Bewegungen näherzubringen.

Eine wichtige Kompetenz, die sie in der zweiten Klasse bereits erlernen müssen, ist das Lesen und Deuten von diversen Diagrammen, was doch einer nicht geringen Zahl von SchülerInnen auch noch in höheren Klassen gewisse Probleme bereitet. Ein netter Zugang, der auch aus meiner Erfahrung heraus die meisten SchülerInnen anspricht, ist es, sich zu den Diagrammen eigene kleine Geschichten auszudenken, die diesen Graphen entsprechen beziehungsweise umgekehrt, eine vorliegende Geschichte mit einem passenden Graphen zu ergänzen. Hervorragend geeignet ist auch ein Bewegungssensor, wie ich ihn selbst in meinen Klassen gerne verwende, da man mit diesem Bewegungsabläufe direkt in ein entsprechendes Diagramm abbilden kann, womit Bewegungsdiagramme quasi selbst erfahren werden können.

Bei einem Weg-Zeit-Diagramm, kurz s-t-Diagramm, ist zunächst wichtig, dass die SchülerInnen begreifen, dass in der vertikalen Richtung die Entfernung zu einem bestimmten festgesetzten Nullpunkt in Abhängigkeit von der Zeit dargestellt wird und dass es sich bei diesem Diagramm nicht um die Bahnkurve eines Körpers handelt.

Bewegt sich der Körper nicht, sondern bleibt auf ihrer Position stehen, so ist der Graph dieser Bewegung eine Gerade parallel zur x-Achse.

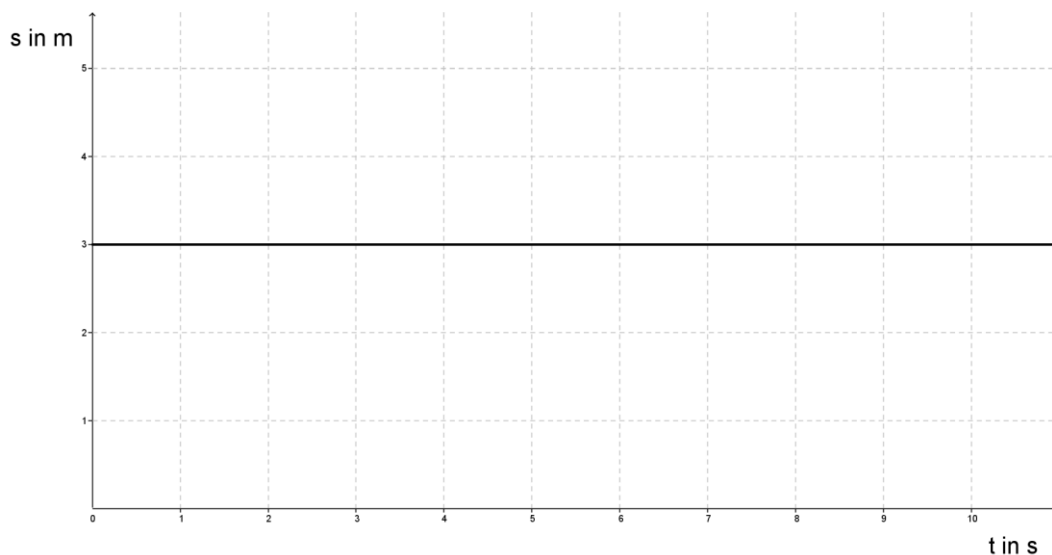


Abbildung 10: s-t-Diagramm eines ruhenden Körpers

Dies wird in der Abbildung 10 dargestellt. Der Körper ist hier in Ruhe und befindet sich 3 m vom Nullpunkt entfernt.⁶⁸

Bewegt sich der Körper mit konstanter Geschwindigkeit vom Nullpunkt weg, so ist der Graph der Bewegung ebenfalls eine Gerade, die aber nicht mehr parallel zur x-Achse verläuft, sondern so geneigt ist, dass ihre Steigung positiv ist. Bewegt sich der Körper mit gleichförmiger Geschwindigkeit wieder auf den Nullpunkt zu, ist die Steigung der Geraden negativ, die Gerade ist also anders geneigt. In der zweiten Klasse kann die Steigung natürlich noch nicht verwendet werden, man erklärt vielmehr qualitativ, dass die Geschwindigkeit umso größer ist, je „steiler“ die Gerade verläuft. Eine entsprechende Bewegung ist in Abbildung 11 dargestellt.

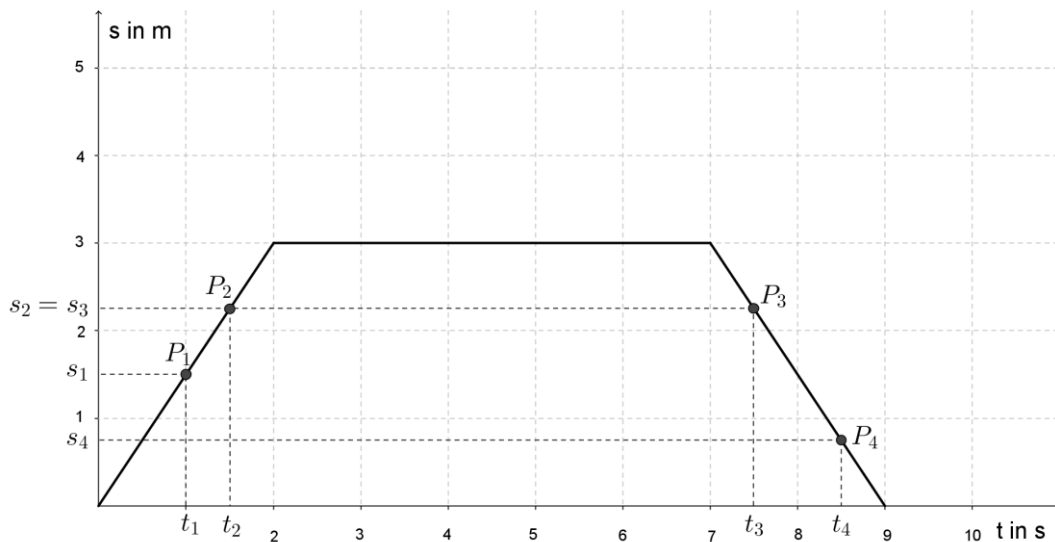


Abbildung 11: s-t-Diagramm eines Körpers mit gleichbleibenden Geschwindigkeiten

Der Körper bewegt sich mit einer gleichförmigen Geschwindigkeit vom Nullpunkt weg, wobei diese Bewegung 2 Sekunden lang andauert. Dann bleibt der Körper stehen, die Geschwindigkeit beträgt Null im Zeitintervall zwischen 2 s bis 7 s. Schließlich bewegt sich der Körper wieder mit einer gleichförmigen Geschwindigkeit auf den

⁶⁸ Genaugenommen wäre der Graph nur dann eine Gerade, wenn die Bewegung, in diesem Fall eine Bewegung mit der Geschwindigkeit Null, schon immer so stattgefunden hat und auch in Zukunft keine Änderung erfahren wird. Dies ist natürlich mathematisch denkbar, physikalisch macht dies aber keinen Sinn. Der Graph ist exakt gesprochen eine Strecke parallel zur x-Achse, die bei $t = 0$ beginnt, wenn keine Zeiten vor dem Zeitpunkt Null betrachtet werden und bei irgendeiner Zeit $t > 0$ endet, was den Endpunkt der Strecke markiert. Würde die Bewegung bei $t = 0$ beginnen und für alle Zeiten $t > 0$ andauern, wäre der Graph ein Strahl parallel zur x-Achse. Üblicherweise wird aber in der Physik allgemein mathematisch ungenau von einer Geraden gesprochen.

Nullpunkt zu, wobei der Betrag der Geschwindigkeit derselbe ist wie zu Beginn, nur die Richtung der Bewegung ist entgegengesetzt. Nach 2 Sekunden beendet der Körper seine Bewegung und seine Geschwindigkeit beträgt ab dem Zeitpunkt 9 Sekunden nach dem Start wieder Null, der Körper ruht also.

In der Abbildung 11 sind 4 ausgewählte Punkt P_1 bis P_4 eingezeichnet, um dies deutlicher zu veranschaulichen. Der Punkt P_1 auf dem Graphen entspricht einem zurückgelegten Weg von $s_1 = 1,5\text{m}$, den der Körper nach $t_1 = 1\text{s}$ vom Nullpunkt aus gemessen zurückgelegt hat. Der Punkt P_2 entspricht einem Zeitpunkt von $t_2 = 1,5\text{s}$, der Körper befindet sich jetzt in einer Entfernung von $s_2 = 2,25\text{m}$ vom Nullpunkt. Da der Zeitpunkt $t_2 > t_1$ ist und ebenfalls $s_2 > s_1$ ist, entfernt sich der Körper in diesem Zeitintervall vom Nullpunkt.

Der Körper hat sich im Zeitintervall $\Delta t = t_2 - t_1 = 1,5\text{s} - 1\text{s} = 0,5\text{s}$ um die Wegstrecke $\Delta s = s_2 - s_1 = 2,25\text{m} - 1,5\text{m} = 0,75\text{m}$ vom Nullpunkt entfernt.

Der Punkt P_3 entspricht dem späteren Zeitpunkt $t_3 = 7,5\text{s}$ und einer Entfernung vom Nullpunkt $s_3 = 2,25\text{m}$, der Punkt P_4 einem Zeitpunkt $t_4 = 8,5\text{s}$ und einer Entfernung $s_4 = 0,75\text{m}$.

Der Körper hat sich nun in diesem Zeitintervall $\Delta t = t_4 - t_3 = 8,5\text{s} - 7,5\text{s} = 1\text{s}$ um die die Wegstrecke $\Delta s = s_4 - s_3 = 0,75\text{m} - 2,25\text{m} = -1,5\text{m}$ vom Nullpunkt entfernt.

Das negative Vorzeichen ist so zu deuten, dass sich der Körper zum Nullpunkt hin zubewegt, da die Entfernung zu einem späteren Zeitpunkt kleiner ist als zu einem früheren.

Man kann durch Ablesen der zurückgelegten Wegstrecken in den gleich großen Zeitintervallen erkennen, dass diese Wegstrecken immer gleich groß sind. Das entspricht aber genau der Definition einer gleichförmigen Geschwindigkeit in der Physik, sodass gefolgert werden kann, dass eine gleichförmige Bewegung immer durch eine Gerade im s-t-Diagramm dargestellt werden kann und die Geschwindigkeit umso größer ist, je steiler die Gerade verläuft.

In der Abbildung 12 sind die Graphen zweier Bewegungen 1 und 2 dargestellt. Es wurden bei jeder Bewegung zwei Zeitintervalle $\Delta t_{1,1}$ und $\Delta t_{1,2}$ für die Bewegung 1 und $\Delta t_{2,1}$ und $\Delta t_{2,2}$ für die Bewegung 2 herausgegriffen, wobei die Zeitintervalle für die beiden Bewegungen gleich gewählt wurden:⁶⁹

$$\Delta t_{1,1} = \Delta t_{1,2} = 4,5\text{s} - 3\text{s} = 1,5\text{s}$$

$$\Delta t_{2,1} = \Delta t_{2,2} = 8,5\text{s} - 7\text{s} = 1,5\text{s}$$

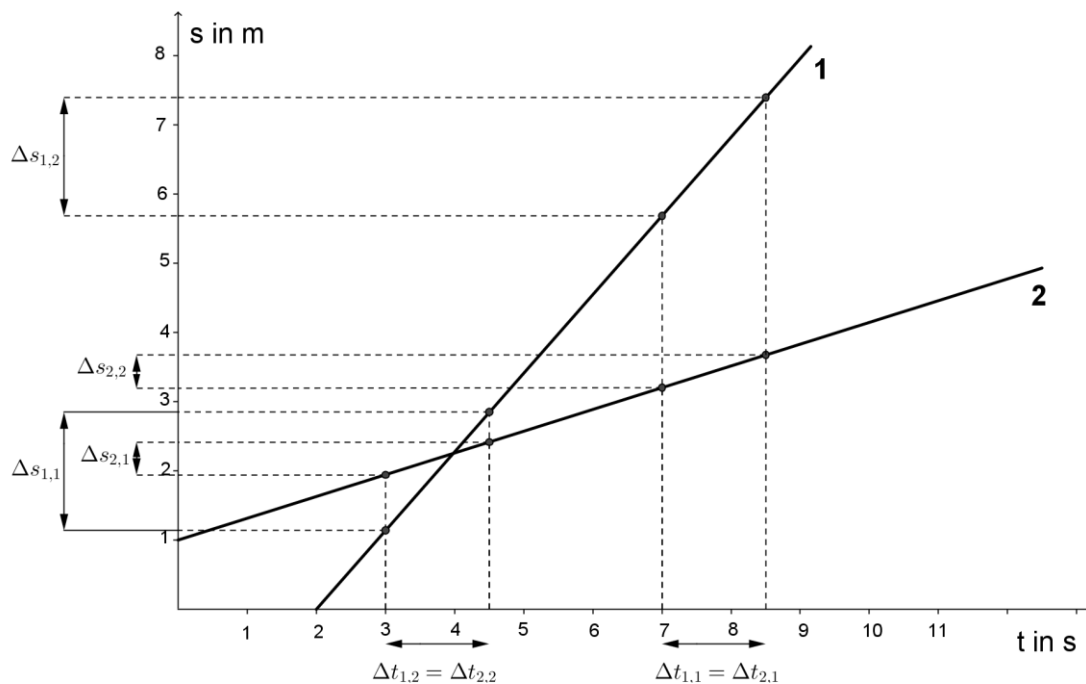


Abbildung 12: Vergleich zweier gleichförmiger Bewegungen

Die Zeitintervalle sind also tatsächlich für beide Bewegungen dieselben, die Zeitpunkte, in denen sie bestimmt wurden, allerdings verschieden. Es sollen nun die zugehörigen Wegintervalle bestimmt werden.

Bewegung 1: $\Delta s_{1,1} = 2,85\text{ m} - 1,15\text{ m} = 1,7\text{ m}$

$$\Delta s_{1,2} = 7,4\text{ m} - 5,7\text{ m} = 1,7\text{ m}$$

Bewegung 2: $\Delta s_{2,1} = 2,42\text{ m} - 1,94\text{ m} = 0,48\text{ m}$

$$\Delta s_{2,2} = 3,67\text{ m} - 3,19\text{ m} = 0,48\text{ m}$$

Daraus kann man nun die bereits genannten Eigenschaften der dargestellten Bewegungsabläufe erkennen:

⁶⁹ Der erste Index bei den Intervallen gibt die Bewegung 1 oder 2 an, der zweite Index bestimmt dann das erste und zweite Intervall jeweils der Bewegung 1 oder 2.

1. Für die Bewegung 1 ergeben sich für beide gleich großen Zeitintervalle $\Delta t_{1,1} = \Delta t_{1,2} = 1,5\text{s}$ auch gleich große Wegintervalle $\Delta s_{1,1} = \Delta s_{1,2} = 1,7\text{m}$, ebenso für die Bewegung 2 für die Zeitintervalle $\Delta t_{2,1} = \Delta t_{2,2} = 1,5\text{s}$ die gleich großen Wegintervalle $\Delta s_{2,1} = \Delta s_{2,2} = 0,48\text{m}$. Da dies für alle beliebig gewählten Zeitintervalle der beiden Bewegungen gilt, müssen die beiden Bewegungen, die durch jeweils eine Gerade dargestellt werden, gleichförmig sein.
2. Da für die beiden Bewegungen 1 und 2 dieselben Zeitintervalle $\Delta t_{1,1} = \Delta t_{1,2} = \Delta t_{2,1} = \Delta t_{2,2} = 1,5\text{s}$ gewählt wurden, ergibt sich, dass der steileren Geraden 1 die größere Geschwindigkeit entspricht, da bei dieser Bewegung die Wegintervalle, das heißt die in derselben Zeit zurückgelegten Strecken, größer sind als bei der Bewegung 2, $\Delta s_{1,1} = \Delta s_{1,2} = 1,7\text{m} > \Delta s_{2,1} = \Delta s_{2,2} = 0,48\text{m}$. Damit ist auch die zweite Behauptung gezeigt, dass eine größere Geschwindigkeit durch eine steilere Gerade dargestellt wird.

Hieraus lässt sich natürlich nun die Geschwindigkeit der beiden Bewegungen bestimmen:

$$v_1 = \frac{\Delta s_{1,1}}{\Delta t_{1,1}} = \frac{\Delta s_{1,2}}{\Delta t_{1,2}} = \frac{1,7\text{m}}{1,5\text{s}} \approx 1,13\text{m} \cdot \text{s}^{-1}$$

$$v_2 = \frac{\Delta s_{2,1}}{\Delta t_{2,1}} = \frac{\Delta s_{2,2}}{\Delta t_{2,2}} = \frac{0,48\text{m}}{1,5\text{s}} \approx 0,32\text{m} \cdot \text{s}^{-1}$$

Im Falle einer gleichförmigen Bewegung entsprechen diese Geschwindigkeiten auch tatsächlich den Momentangeschwindigkeiten, da Zeitintervalle und Wegintervalle direkt proportional zueinander sind und somit auch der Grenzwert mit $\Delta t \rightarrow 0$ denselben Wert für die jeweilige Geschwindigkeit ergibt.

Ist die Geschwindigkeit nicht gleichförmig, so wird der Körper beschleunigt, eine positive Beschleunigung bedeutet, die Geschwindigkeit des Körpers nimmt zu, eine negative, die Geschwindigkeit nimmt ab. Ein Beispiel einer nichtgleichförmigen Bewegung ist in Abbildung 13 dargestellt.

Es sind drei Zeitintervalle gleicher Größe $\Delta t_1 = \Delta t_2 = \Delta t_3 = 1\text{s}$ ausgewählt. Diesen Zeitintervallen entsprechen jeweils ein Wegintervall $\Delta s_1, \Delta s_2, \Delta s_3$, die offensichtlich

nicht mehr dieselbe Größe haben. Um dies zu quantifizieren, sollen diese Wegintervalle bestimmt werden:

$$\Delta s_1 = 0,9 \text{ m} - 0,4 \text{ m} = 0,5 \text{ m}$$

$$\Delta s_2 = 2,5 \text{ m} - 1,6 \text{ m} = 0,9 \text{ m}$$

$$\Delta s_3 = 4,9 \text{ m} - 3,6 \text{ m} = 1,3 \text{ m}$$

Man erkennt, dass die Wegintervalle in denselben Zeitintervallen immer mehr zunehmen, die Bewegung ist somit eine beschleunigte Bewegung.

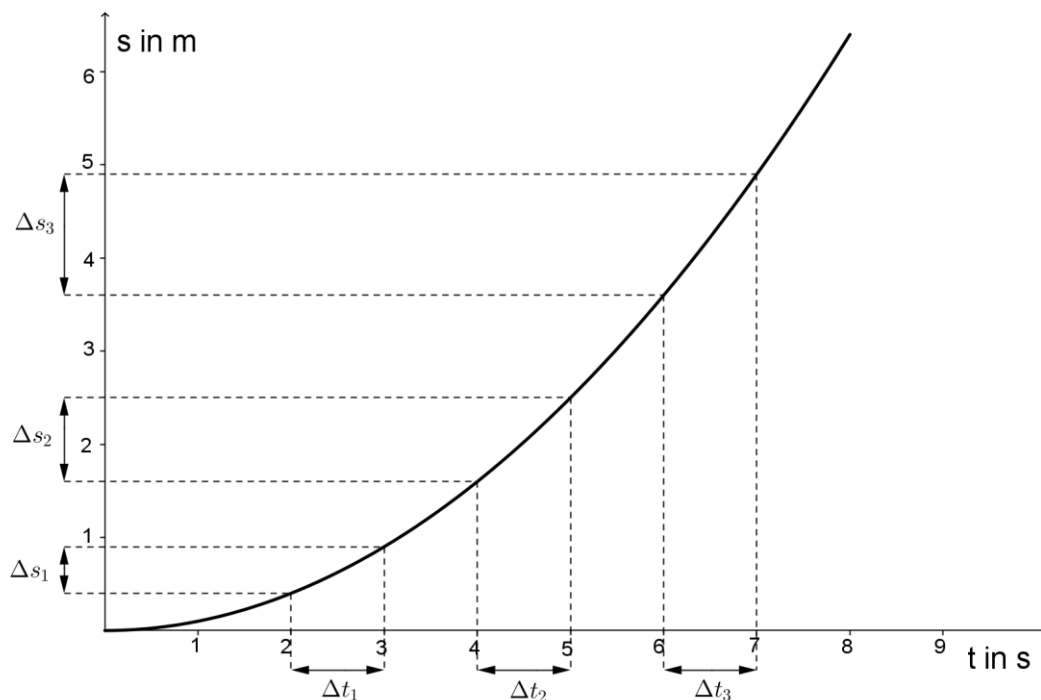


Abbildung 13: Eine nicht-gleichförmige Bewegung

Die Durchschnittsgeschwindigkeiten ergeben sich daraus als

$$\Delta v_1 = \frac{\Delta s_1}{\Delta t_1} = \frac{0,5 \text{ m}}{1 \text{ s}} \approx 0,5 \text{ m} \cdot \text{s}^{-1}$$

$$\Delta v_2 = \frac{\Delta s_2}{\Delta t_2} = \frac{0,9 \text{ m}}{1 \text{ s}} \approx 0,9 \text{ m} \cdot \text{s}^{-1}$$

$$\Delta v_3 = \frac{\Delta s_3}{\Delta t_3} = \frac{1,3 \text{ m}}{1 \text{ s}} \approx 1,3 \text{ m} \cdot \text{s}^{-1}$$

Um hieraus die Momentangeschwindigkeiten zu erhalten, muss man die Zeitintervalle immer kleiner werden lassen und schlussendlich den Grenzwert mit $\Delta t \rightarrow 0$ bilden.

Ein Anwendungsbeispiel für eine solche Bewegung wäre etwa der freie Fall, wie in der Abbildung 14 dargestellt. Zur Zeit $t=1\text{s}$ wird der Körper aus einer Höhe $s=5\text{m}$ losgelassen und erreicht den Boden, was einem $s=0\text{m}$ entspricht, nach $t=2,01\text{s}$. Der Körper befindet sich somit $t=2,01\text{s}-1\text{s}=1,01\text{s}$ lang im freien Fall.

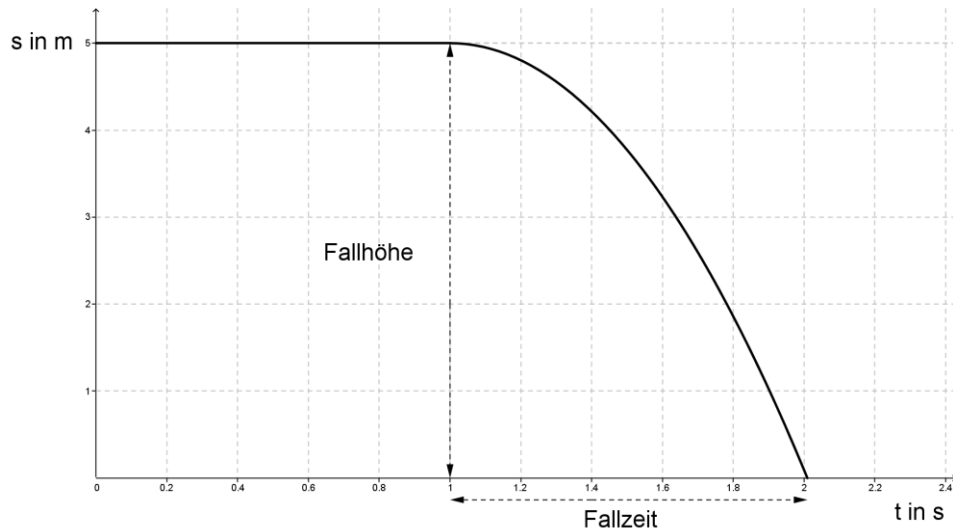


Abbildung 14: s-t-Diagramm des freien Falls

Es ist wichtig, dass die SchülerInnen mit Funktionsgraphen umzugehen lernen und fähig sind, diverse Fragestellungen anhand des vorgelegten Graphen zu beantworten. Eine mögliche Aufgabenstellung für die zweite Klasse Physik könnte etwa so aussehen:

Betrachte den Graphen aus Abbildung 15 und erfinde eine passende Geschichte dazu! Beachte die Achsenbeschriftungen!

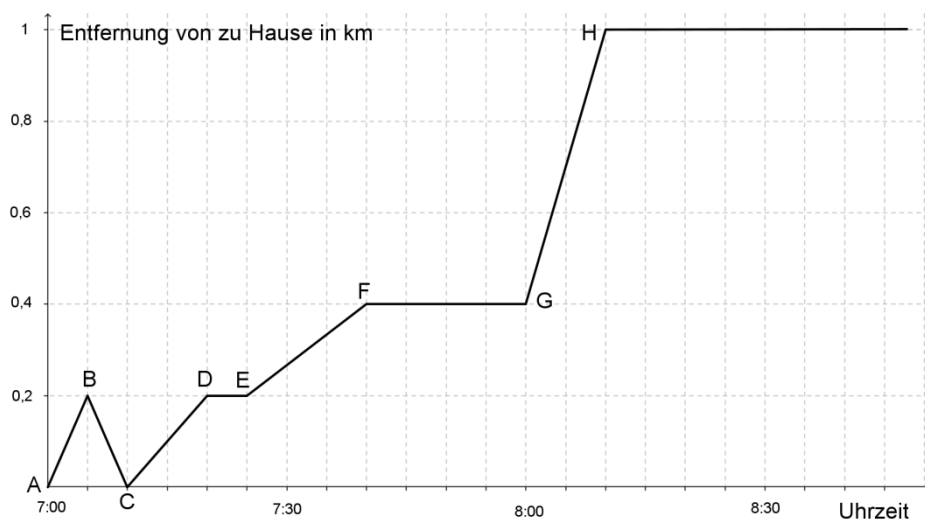


Abbildung 15: Ein s-t-Diagramm mit einer Geschichte

Eine mögliche Lösung wäre etwa die folgende:

Der Graph beschreibt den Schulweg von Sarah. Sarah geht um 7 Uhr von zu Hause weg (A). Um 7:05 Uhr merkt sie plötzlich, dass sie zu Hause ihren Schülerschein vergessen hat und geht wieder zurück. Sie kommt um 7:10 Uhr zu Hause an (B), nimmt von ihrer Mutter den Schülerschein entgegen und geht gleich wieder in die Schule (C). Sie geht jetzt langsamer als vorher und trifft um 7:20 Uhr ihre Freundin Clara (D). Sie begrüßen sich und reden 5 Minuten miteinander und gehen dann weiter in die Schule (E). Sie spazieren jetzt langsamer, weil sie miteinander plaudern und kommen dann nach 15 Minuten um 7:40 Uhr an einem Kleidergeschäft vorbei (F). Dort bleiben sie stehen und sehen sich die Kleider an. Dabei vergessen sie ganz auf die Zeit und merken plötzlich, dass es schon 8 Uhr geworden ist. Jetzt beginnen sie zu rennen (G) und kommen um 10 Minuten um 8:10 Uhr in der Schule an (H). Sie sind 10 Minuten zu spät zum Unterricht gekommen!

Wenn man ein entsprechend interessantes und eventuell auch komplizierteres s-t-Diagramm zeichnet, könnte man diese Aufgabe auch durchaus für einen fächerübergreifenden Unterricht Deutsch – Physik verwenden, was sicher einen interessanten Aspekt mit hineinbringen würde. Man könnte aber auch umgekehrt etwa einen passenden Bericht verwenden und die SchülerInnen daraus das entsprechende s-t-Diagramm zeichnen lassen.

Kurz soll noch auf die Darstellung der Bewegung in einem v-t-Diagramm eingegangen werden. Dazu soll das s-t-Diagramm aus der vorigen Aufgabe (Abbildung 15) als Grundlage verwendet werden und daraus das entsprechende v-t-Diagramm abgeleitet werden.

Die Geschwindigkeiten können leicht berechnet werden, da es sich nur um gleichförmige Bewegungen handelt, sodass keine Differentialrechnung zur Herleitung notwendig ist.

Der gesamte Bewegungsverlauf kann in acht Zeitabschnitte aufgeteilt werden, wobei in jedem einzelnen Zeitabschnitt die Geschwindigkeit konstant bleibt:

Abschnitt	Im Zeitintervall zurückgelegter Weg in km	Benötigte Zeit in h	Geschwindigkeit im Zeitintervall in km/h
A - B	0,2 km	1/12 h	$0,2 \text{ km} : (1/12 \text{ h}) = 2,4 \text{ km/h}$
B - C	-0,2 km	1/12 h	$-0,2 \text{ km} : (1/12 \text{ h}) = -2,4 \text{ km/h}$
C - D	0,2 km	1/6 h	$0,2 \text{ km} : (1/6 \text{ h}) = 1,2 \text{ km/h}$
D - E	0 km	1/12 h	$0 \text{ km} : (1/12 \text{ h}) = 0 \text{ km/h}$
E - F	0,2 km	1/4 h	$0,2 \text{ km} : (1/4 \text{ h}) = 0,8 \text{ km/h}$
F - G	0 km	1/3 h	$0 \text{ km} : (1/3 \text{ h}) = 0 \text{ km/h}$
G - H	0,6 km	1/6 h	$0,6 \text{ km} : (1/6 \text{ h}) = 3,6 \text{ km/h}$
ab H	0 km	beliebig	0 km/h

Damit kann nun das entsprechende v-t-Diagramm gezeichnet werden (Abbildung 16).

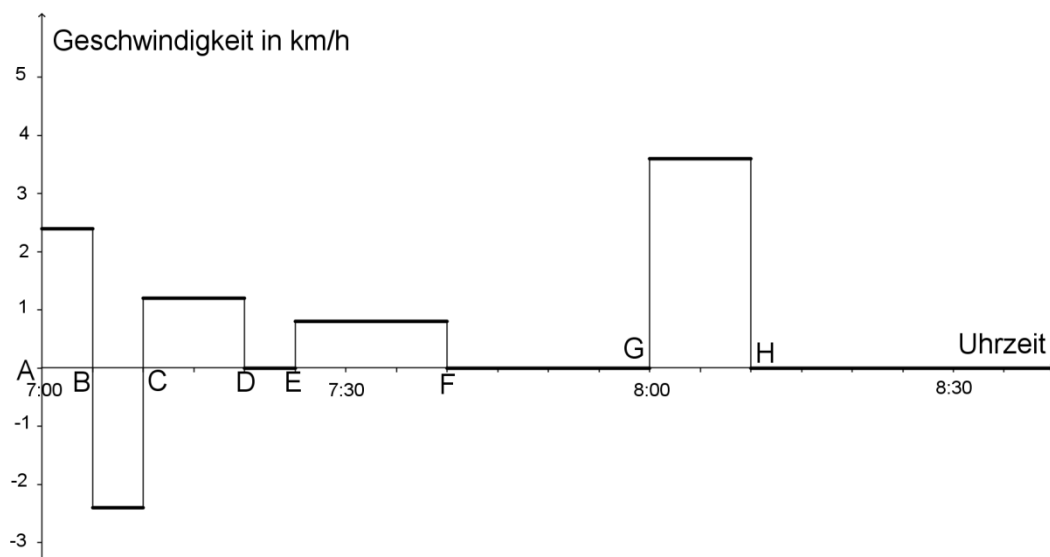


Abbildung 16: v-t-Diagramm zu Abbildung 15

Im v-t-Diagramm erkennt man, dass ein Geschwindigkeitswechsel in Nullzeit geschieht, was natürlich physikalisch unmöglich ist, insofern kann das hier gezeigte v-t-Diagramm, und damit auch das zugehörige s-t-Diagramm nur ein mathematisches Konstrukt sein, das physikalisch so nicht realisierbar ist.

Die Durchschnittsbeschleunigung ist als Geschwindigkeitsänderung geteilt durch die benötigte Zeit definiert:

$$\Delta a = \frac{v(t_2) - v(t_1)}{t_2 - t_1}$$

Für die Durchschnittsbeschleunigung ergibt sich auch anhand des Graphen in Abbildung 16 kein Problem, die Durchschnittsbeschleunigung kann immer angegeben werden, da $t_2 - t_1 > 0$ immer gilt. Probleme ergeben sich erst bei der Momentanbeschleunigung, wenn die Beschleunigung genau an einem der Unstetigkeitspunkte bestimmt werden soll.

Die Momentanbeschleunigung a ist definiert als

$$a = \lim_{t_2 \rightarrow t_1} \frac{v(t_2) - v(t_1)}{t_2 - t_1} = \lim_{\Delta t \rightarrow 0} \frac{\Delta v}{\Delta t} = \frac{dv}{dt}$$

also als Ableitung der Geschwindigkeit nach der Zeit. Eine Ableitung ist aber nur für Stetigkeitsstellen definiert, deswegen auch die Probleme an den besagten Stellen auftreten.

Dieses Problem sollte mit den SchülerInnen eingehend besprochen werden, da hier sehr deutlich wird, wie eine zu starke Vereinfachung zu Problemen führen kann. Auch sollte im Physikunterricht diskutiert werden, dass mathematische Aussagen natürlich immer dahingehend überprüft werden müssen, ob sie physikalisch Sinn machen oder nicht.

Beispiel 2: Bestimmung einer Regressionsgeraden

In den beiden Kapiteln 3.1.1 und 3.1.3 wurde durch Messungen der Zeiten, die eine LäuferIn für bestimmte Wegstrecken gebraucht hat, ihre Geschwindigkeit bestimmt. Nun soll durch die experimentell bestimmten Zahlenpaare (Wegstrecke, Zeitdauer) eine Gerade (Regressionsgerade, Ausgleichsgerade) approximiert werden und in ein Koordinatensystem eingezeichnet werden.

Es gibt grundsätzlich verschiedene Möglichkeiten, die „beste“ Gerade zu finden, die die gemessenen Punkte am besten annähert. Üblicherweise wird die Gerade in der Weise bestimmt, dass die Summe der quadrierten Vertikalabstände der y-Koordinaten der Punkte zu der Geraden minimal sein soll.⁷⁰

⁷⁰ Eine andere Möglichkeit wäre, dass die Summe der Beträge der Vertikalabstände der y-Koordinaten der Punkte zu der Geraden minimal sein soll.

Sei $y = a \cdot x + b$ die gesuchte Regressionsgerade, wobei a, b die zu bestimmenden Parameter der Geraden sind. Seien die Wegstrecken auf der x -Achse aufgetragen und mit $x_1, \dots, x_9$ bezeichnet, da 9 Messwerte bestimmt wurden. Die gemessenen Zeiten seien auf der y -Achse aufgetragen und mit $y_1, \dots, y_9$ bezeichnet. Die Werte sind noch einmal aufgelistet:

x (m)	2	4	6	8	10	12	14	16	18
y (s)	2,0	3,7	6,1	9,2	10,1	12,2	14,0	16,5	18,2

Die Summe der quadrierten Vertikalabstände der y -Koordinaten der Punkte von der Geraden ergibt sich zu

$$\sum_{i=1}^9 (y(x_i) - y_i)^2 = \sum_{i=1}^9 (a \cdot x_i + b - y_i)^2$$

und dieser Ausdruck soll ein Minimum werden.

Es soll kurz der Weg zur Lösung dieses Problems gezeigt werden⁷¹, wobei die Lösung gleich für beliebig viele Punkte gelten soll, die Anzahl der Punkte sei mit n bezeichnet.

Das Problem vereinfacht sich beträchtlich, wenn man das Koordinatensystem in den „Schwerpunkt“ der Datenpunkte verschiebt, wobei der „Schwerpunkt“ durch die arithmetischen Mittelwerte

$$\bar{x} = \frac{1}{n} \cdot \sum_{i=1}^n x_i \quad \text{und} \quad \bar{y} = \frac{1}{n} \cdot \sum_{i=1}^n y_i \quad \text{gegeben ist.}$$

Die neuen Koordinaten sind dann

$$X_i = x_i - \bar{x} \quad \text{und} \quad Y_i = y_i - \bar{y}$$

Das Schöne ist nun, dass die Summen der X_i und der Y_i gerade Null ergeben:

$$\sum_{i=1}^n (x_i - \bar{x}) = \sum_{i=1}^n x_i - \sum_{i=1}^n \bar{x} = n \cdot \bar{x} - n \cdot \bar{x} = 0$$

⁷¹ Ich selbst habe diese Methode in der Vorlesung „Stochastik“ von Univ.-Prof. Mag. Dr. Humenberger an der Universität Wien kennen gelernt, wofür ich sehr dankbar bin, da diese Methode auch bei anderen Problemstellungen angewendet werden kann, wenn die Differentialrechnung noch nicht zur Verfügung steht beziehungsweise ansonsten partielle Ableitungen verwendet werden müssten.

$$\sum_{i=1}^n (y_i - \bar{y}) = \sum_{i=1}^n y_i - \sum_{i=1}^n \bar{y} = n \cdot \bar{y} - n \cdot \bar{y} = 0$$

Die Regressionsgerade in den neuen Koordinaten schreibt sich somit als

$$Y = a \cdot X + c$$

Die Steigung der ursprünglichen Regressionsgeraden bleibt bei der Koordinatentransformation erhalten, die Koordinaten werden ja nur verschoben aber nicht gedreht. Allerdings ändert sich der Ordinatenabstand b , sodass im neuen System ein anderer Parameter c statt b zu bestimmen ist.

Die Minimalbedingung lässt sich nun schreiben als

$$\sum_{i=1}^n (a \cdot X_i + c - Y_i)^2.$$

Es wird also eine Regressionsgerade in den neuen Koordinaten ermittelt und diese dann wieder in die alten Koordinaten transformiert. Normalerweise würde man eine derartige Extremwertaufgabe mit Hilfe der Differentialrechnung lösen, wobei in diesem Fall, da zwei Parameter variabel sind, partielle Ableitungen benötigt würden, welche aber nicht zum Unterrichtsstoff gehören. Man kann sich aber mit einem sehr einfachen „Trick“ die Verwendung der Differentialrechnung gänzlich sparen, indem man die Methode der quadratischen Ergänzung geschickt anwendet.

Man quadriert die Minimalbedingung zuerst einmal aus

$$\sum_{i=1}^n (a \cdot X_i + c - Y_i)^2 = \sum_{i=1}^n (a^2 \cdot X_i^2 + 2 \cdot a \cdot c \cdot X_i - 2 \cdot a \cdot X_i \cdot Y_i + c^2 - 2 \cdot c \cdot Y_i + Y_i^2)$$

Jetzt kann einiges vereinfacht werden und es ergibt sich

$$\begin{aligned} \sum_{i=1}^n (a \cdot X_i + c - Y_i)^2 &= a^2 \cdot \sum_{i=1}^n X_i^2 + 2 \cdot a \cdot c \cdot \sum_{i=1}^n X_i - 2 \cdot a \cdot \sum_{i=1}^n X_i \cdot Y_i + \sum_{i=1}^n c^2 - 2 \cdot c \cdot \sum_{i=1}^n Y_i + \sum_{i=1}^n Y_i^2 \\ &\quad \quad \quad = 0 \quad \quad \quad = n \cdot c^2 \quad \quad \quad = 0 \\ &= a^2 \cdot \sum_{i=1}^n X_i^2 - 2 \cdot a \cdot \sum_{i=1}^n X_i \cdot Y_i + \sum_{i=1}^n Y_i^2 + n \cdot c^2 \end{aligned}$$

Die beiden ersten Terme in der erhalten Beziehung werden nun auf ein vollständiges Quadrat ergänzt:

$$\begin{aligned}
a^2 \cdot \sum_{i=1}^n X_i^2 - 2 \cdot a \cdot \sum_{i=1}^n X_i \cdot Y_i &= \sum_{i=1}^n X_i^2 \cdot \left(a^2 - 2 \cdot a \cdot \frac{\sum_{i=1}^n X_i \cdot Y_i}{\sum_{i=1}^n X_i^2} \right) \\
&= \sum_{i=1}^n X_i^2 \cdot \left(a - \frac{\sum_{i=1}^n X_i \cdot Y_i}{\sum_{i=1}^n X_i^2} \right)^2 - \left(\frac{\sum_{i=1}^n X_i \cdot Y_i}{\sum_{i=1}^n X_i^2} \right)^2
\end{aligned}$$

Somit ergibt sich der recht komplexe Ausdruck

$$\sum_{i=1}^n (a \cdot X_i + c - Y_i)^2 = \sum_{i=1}^n X_i^2 \cdot \left(a - \frac{\sum_{i=1}^n X_i \cdot Y_i}{\sum_{i=1}^n X_i^2} \right)^2 - \left(\frac{\sum_{i=1}^n X_i \cdot Y_i}{\sum_{i=1}^n X_i^2} \right)^2 + \sum_{i=1}^n Y_i^2 + n \cdot c^2$$

Damit dieser Ausdruck minimal wird, müssen die beiden Terme, in denen die Parameter a und b vorkommen, möglichst klein werden. Da beide Ausdrücke quadratische Form haben, sind sie immer größer oder gleich Null. Ihren kleinsten Wert nehmen sie also an, wenn beide Ausdrücke verschwinden, also Null werden:

$$(1) \quad n \cdot c^2 = 0$$

$$(2) \quad a - \frac{\sum_{i=1}^n X_i \cdot Y_i}{\sum_{i=1}^n X_i^2} = 0$$

Daraus ergeben sich nun direkt die beiden gesuchten Parameter zu $c=0$ und

$$a = \frac{\sum_{i=1}^n X_i \cdot Y_i}{\sum_{i=1}^n X_i^2}, \text{ womit das Regressionsproblem in den neuen Koordinaten gelöst ist.}$$

Die Gerade hat die Form $Y = a \cdot X$, sie geht also durch den Nullpunkt des neuen Koordinatensystems. Da der Nullpunkt dem Schwerpunkt $(\bar{x}, \bar{y})$ des alten Systems entspricht, muss die Regressionsgerade im alten Koordinatensystem durch den Schwerpunkt verlaufen!

Jetzt bleibt nur mehr in die Ursprungskoordinaten zurück zu transformieren.

Die Gerade lautete ja ursprünglich $y = a \cdot x + b$. Wie oben ausgeführt, verläuft die Gerade durch den Schwerpunkt, also ist dieser ein Punkt der Regressionsgeraden und es gilt $\bar{y} = a \cdot \bar{x} + b$.

Daraus ergibt sich für den Parameter b

$$b = \bar{y} - a \cdot \bar{x}$$

Der Parameter a ergibt sich durch Einsetzen zu

$$a = \frac{\sum_{i=1}^n (x_i - \bar{x}) \cdot (y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

Somit lauten die Standardformeln für die lineare Regression:

Es seien n Datenpunkte $(x_1, y_1), \dots, (x_n, y_n)$ gegeben. Die Regressionsgerade, die die n Datenpunkte am besten annähert lautet $y = a \cdot x + b$ mit

$$a = \frac{\sum_{i=1}^n (x_i - \bar{x}) \cdot (y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2} \text{ und } b = \bar{y} - a \cdot \bar{x}.$$

Natürlich ist es nicht notwendig, wenngleich sicher lehrreich, die Herleitung im Mathematikunterricht zu zeigen. Es sollte aber zumindest die LehrerIn die Herleitung kennen, um zu verstehen und entsprechende Fragen auch fachlich begründet beantworten zu können, weshalb die Regressionsgerade gerade diese Form hat und wie dies begründet werden kann. Natürlich gibt es in den gängigen Tabellenkalkulationsprogrammen wie Excel standardmäßig diese Anpassungen implementiert, die Regressionsgerade kann etwa durch den Befehl „Trendlinie einfügen“ eingezeichnet werden. Aber um wirklich verstehen zu können, was es mit dieser Regressionsgeraden beziehungsweise Trendlinie auf sich hat, ist ein tieferer

Einstieg in die Mathematik unumgänglich. Dies ist auch meine Begründung, weshalb ich diese doch etwas mühsame Herleitung gezeigt habe.

Es soll nun aber noch die Regressionsgerade auf das Problem der Messpunkte des ursprünglichen Versuchs ermittelt werden. Für die Mittelwerte ergibt sich:

$$\bar{x} = \frac{1}{9} \cdot \sum_{i=1}^9 x_i = 10, \quad \bar{y} = \frac{1}{9} \cdot \sum_{i=1}^9 y_i \approx 10,22$$

x	2	4	6	8	10	12	14	16	18
y	2,0	3,7	6,1	9,2	10,1	12,2	14,0	16,5	18,2
$x - \bar{x}$	-8	-6	-4	-2	0	2	4	6	8
$y - \bar{y}$	-8,22	-6,52	-4,12	-1,02	-0,12	1,98	3,78	6,28	7,98
$(x_i - \bar{x}) \cdot (y_i - \bar{y})$	65,78	39,13	16,49	2,04	0	3,96	15,11	37,67	63,82
$(x_i - \bar{x})^2$	64	36	16	4	0	4	16	36	64
$(y_i - \bar{y})^2$	67,60	42,54	16,99	1,04	0,01	3,91	14,27	39,41	63,64

Daraus ergibt sich durch Summenbildung $a \approx 1,02$ und $b \approx 0,05$, sodass sich die folgende Regressionsgerade ergibt:

$$y = 1,02 \cdot x + 0,05$$

In der Abbildung 17 sind die gemessenen Datenpunkte und die berechnete Regressionsgerade eingezeichnet.

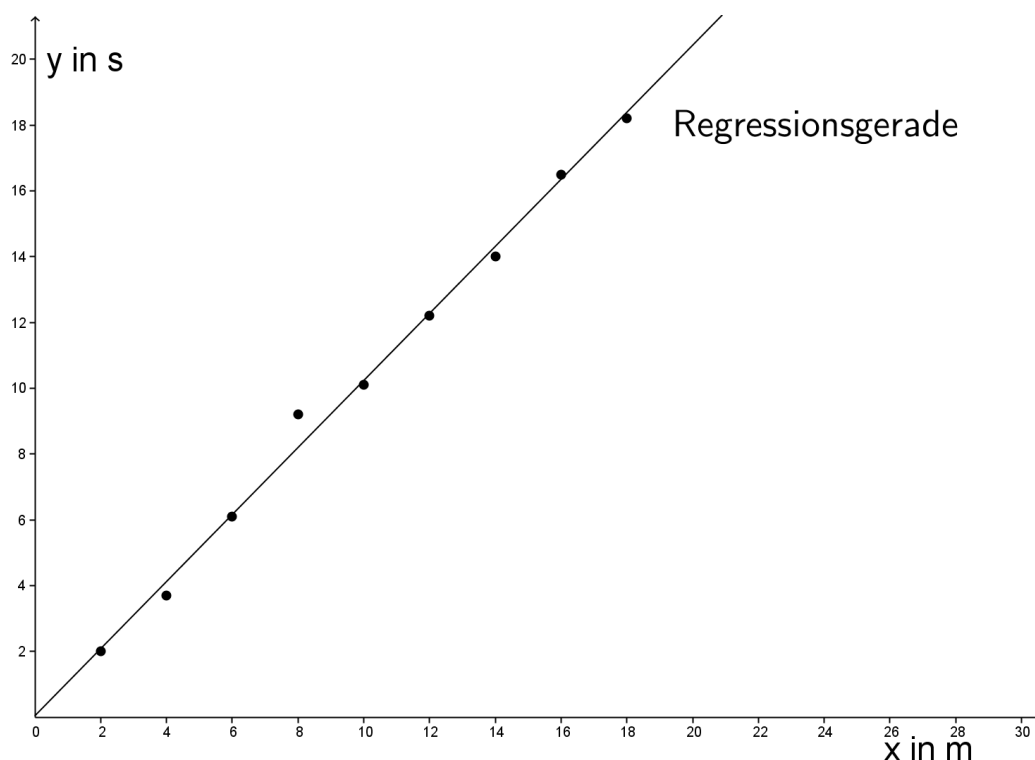


Abbildung 17: Regressionsgerade

Ohne Herleitung angegeben sei noch der zu jeder Regressionsgeraden anzugebende Korrelationskoeffizient r . Dieser kann Werte zwischen -1 und $+1$ annehmen. Ein Korrelationskoeffizient nahe ± 1 bedeutet eine sehr gute Anpassung an die Daten, ein Korrelationskoeffizient nahe Null bedeutet schlechte beziehungsweise bei exakt Null keine Korrelation mehr.

Der Korrelationskoeffizient wird nach folgender Formel berechnet:

$$r = \frac{\sum_{i=1}^n (x_i - \bar{x}) \cdot (y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \cdot \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}}$$

Für das obige Beispiel ergibt sich für den Korrelationskoeffizienten ein Wert von $r \approx 0,997$. Das bedeutet, dass die Datenpunkte hochkorreliert sind.

Für den Unterricht ist es wichtig, dass die SchülerInnen eine Vorstellung davon haben, was Regression und Korrelation bedeuten. Niemand wird händisch eine

Regressionsgerade oder den Korrelationskoeffizienten in der Praxis berechnen. Die SchülerInnen können aber durchaus angehalten werden, eine Datenreihe mit dem Computer und entsprechenden Programmen zu bearbeiten und zu analysieren. Es wäre zum Beispiel möglich, das Ohm'sche Gesetz auf diese Weise zu handhaben, indem man etwa an einem konstanten Widerstand die Stromstärken zu verschiedenen Spannungen bestimmt und daraus mittels Regression das Ohm'sche Gesetz zeigt. Die Strom- und Spannungsmessung ist nie so präzise, dass alle Datenpunkte wirklich auf einer Gerade liegen, sodass sich eine Regressionsanalyse anbietet. Auch andere physikalische Untersuchungen finden sich sicher leicht, man denke nur an die verschiedenen Gasgesetze, die man untersuchen könnte.

3.5 5. Klasse

3.5.1 Koordinatensysteme

Mit dem Thema Koordinatensysteme werden die SchülerInnen ab der zweiten Klasse konfrontiert und es bleibt auch ein wichtiges Thema bis zur Matura. Dienen Koordinatensysteme in der Unterstufe dazu, einfache Geometrieaufgaben grafisch zu lösen, werden ab der fünften Klasse mit der Vektorrechnung und der Einführung in die Trigonometrie Koordinatensysteme in erweiterter Form verwendet. Mit Hilfe der Vektorrechnung geht es darum, neben Längen und Winkeln insbesondere auch Koordinaten besonderer Punkte, die bisher in der Unterstufe nur grafisch ermittelt worden waren, zu berechnen. Bei den Vektoren beschränkt man sich dabei auf kartesische Koordinatensysteme. Sobald die trigonometrischen Funktionen Sinus, Cosinus und Tangens⁷² eingeführt worden sind, werden auch (ebene) Polarkoordinaten behandelt. Diese spielen dann auch bei den komplexen Zahlen, die Lehrstoff der siebten Klasse sind, eine wichtige Rolle, da mit ihnen in der Gauß'schen Zahlenebene sehr leicht gerechnet werden kann, wenn es um die Multiplikation und Division von komplexen Zahlen geht. Sie spielen weiters eine gewisse Rolle, wenn etwa die Wurzeln von Gleichungen der Form $x^n = z$ bestimmt werden sollen, wobei n eine natürliche Zahl und z ein Element der komplexen Zahlen sein soll.

Für die Physik sind natürlich außer dem kartesischen Koordinatensystem auch andere krummlinige Koordinatensysteme sehr wichtig. Ebene Polarkoordinaten werden zwar behandelt, aber „räumliche Polarkoordinaten“, üblicherweise als Kugelkoordinaten bezeichnet, werden nicht besprochen. Da in der Physik aber sehr viele räumlich kugelsymmetrische Aufgaben vorkommen, die zweckmäßigerweise in Kugelkoordinaten behandelt werden⁷³, sollen dieses Koordinatensystem sowie die

⁷² Der Kotangens ist nicht mehr im Lehrplan enthalten!

⁷³ Es würde, um nur ein Beispiel zu benennen, die Berechnung der Elektronenaufenthaltswahrscheinlichkeit im Wasserstoffatom unnötig verkomplizieren, wenn hier statt Kugelkoordinaten kartesische Koordinaten herangezogen würden.

Zylinderkoordinaten kurz vorgestellt werden. Wichtig ist es in der Physik, zwischen den Koordinaten in den verschiedenen Koordinatensystemen umrechnen zu können.⁷⁴

Beispiel 1: Ebene Polarkoordinaten

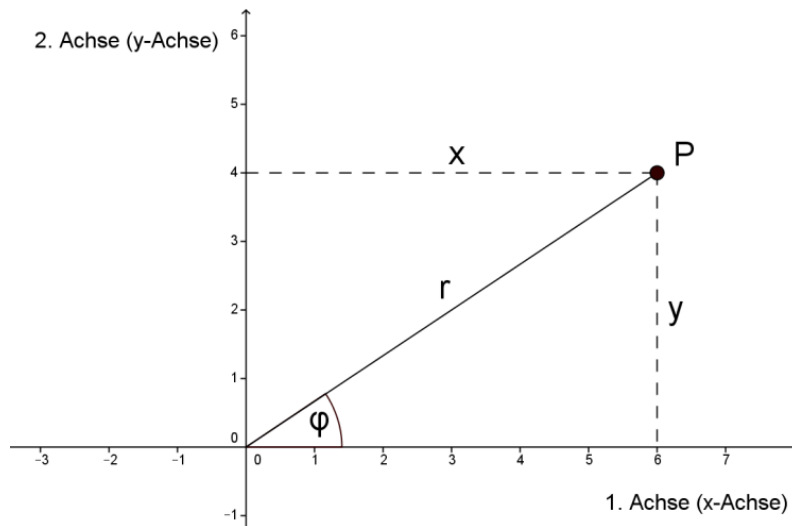


Abbildung 18: Ebene Polarkoordinaten

Der Punkt P kann sowohl durch die Angabe der kartesischen Koordinaten x und y als auch durch die Angabe der polaren Koordinaten r und φ eindeutig bestimmt werden. Die Umrechnung in Polarkoordinaten erfolgt sehr elementar und einfach unter Verwendung des pythagoreischen Lehrsatzes und der Basisdefinition des Tangens und ist direkt aus der Abbildung 18: Ebene Polarkoordinaten ersichtlich:

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

$$\tan \varphi = \frac{y}{x} \Rightarrow \varphi = \arctan\left(\frac{y}{x}\right)$$

Die Umrechnung von Polarkoordinaten in kartesische Koordinaten erfolgt genauso einfach durch die Grunddefinitionen des Sinus und des Cosinus:

⁷⁴ Es ist aber nicht nur die Umrechnung von Koordinaten nötig, sondern etwa auch in der „höheren“ Physik die „Umrechnung“ von Vektoroperatoren, wie Gradient, Divergenz oder Rotation. Diese gehören nicht zum Schulstoff und deshalb gehe ich in dieser Arbeit auch nicht weiter darauf ein.

$$x = r \cdot \cos \varphi$$

$$y = r \cdot \sin \varphi$$

Es gelten also für die Umrechnung ebener Polarkoordinaten in kartesische Koordinaten und umgekehrt folgende Beziehungen:

$$\left. \begin{array}{l} x = r \cdot \cos \varphi \\ y = r \cdot \sin \varphi \end{array} \right\} \Leftrightarrow \left\{ \begin{array}{l} r = \sqrt{x^2 + y^2} \\ \varphi = \arctan\left(\frac{y}{x}\right) \end{array} \right.$$

Es gelten dabei die Wertebereiche: $0 \leq r < \infty$

$$0 \leq \varphi \leq 2\pi$$

Beispiel 2: Räumliche Polarkoordinaten (Kugelkoordinaten)

Diese Koordinaten sind sehr häufig anzutreffen und finden bei allen kugelsymmetrischen Anwendungen ihre Verwendung. Erwähnt werden sollte noch, dass bei Verwendung dieser Koordinaten auch eigene an die Koordinaten angepasste Funktionen, welche allerdings in der Schulmathematik nicht vorkommen, eingeführt werden, die Kugelflächenfunktionen. Diese sind in vielen Fachgebieten, wie besonders etwa der Geophysik, wenn es um die Beschreibung der Eigenschaften der Erde geht oder in der Astronomie, wenn es etwa um das Schwingungsverhalten der Sterne geht, anzutreffen.

Der Winkel φ wird üblicherweise wie in der Abbildung 19: Räumliche Polarkoordinaten oder Kugelkoordinaten gezeigt, definiert, das ist aber nicht immer so der Fall und man muss sich bei Verwendung dieser Koordinaten informieren, wie dieser Winkel nun genau definiert ist. Man denke etwa an das Gradnetz der Erde, wo der Winkel φ beim Äquator 0° beträgt und dann bis zum Nordpol auf $+90^\circ$ anwächst beziehungsweise bis zum Südpol auf -90° .

Im Mathematikunterricht können aber die Umrechnungsformeln für die Kugelkoordinaten durchaus auch einfach als kleine Übung für die trigonometrischen Formeln herangezogen werden und dann ein paar Umrechnungen durchgeführt werden.

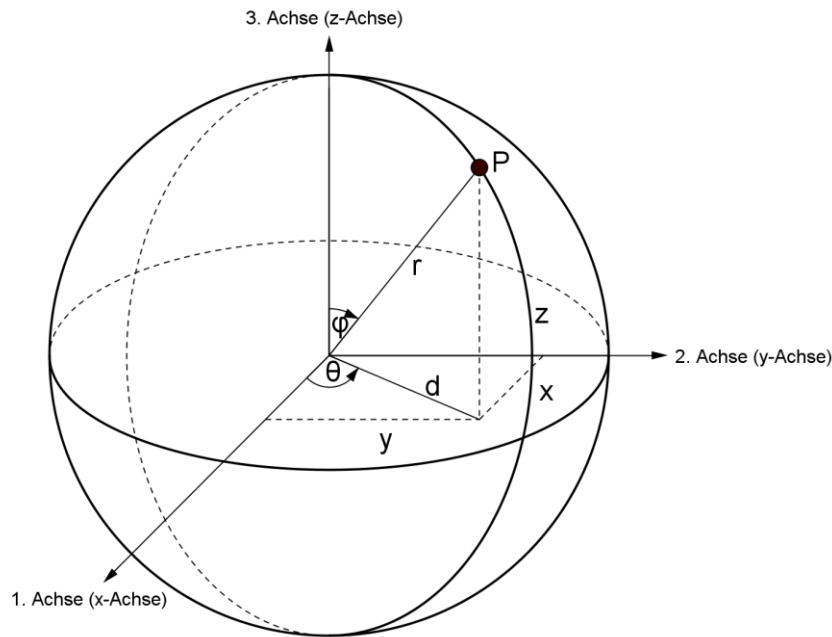


Abbildung 19: Räumliche Polarkoordinaten oder Kugelkoordinaten

Nach dem Lehrsatz des Pythagoras ergibt sich als erstes folgende Beziehung

$$d = \sqrt{x^2 + y^2}$$

und daraus weiter laut Abbildung 19

$$r = \sqrt{d^2 + z^2} = \sqrt{x^2 + y^2 + z^2} .$$

Für den

$$\tan \theta = \frac{y}{x} \Rightarrow \theta = \arctan\left(\frac{y}{x}\right)$$

$$\tan \varphi = \frac{d}{z} = \frac{\sqrt{x^2 + y^2}}{z} \Rightarrow \varphi = \arctan \frac{\sqrt{x^2 + y^2}}{z}$$

Umgekehrt gilt mit $d = r \cdot \sin \varphi$:

$$x = d \cdot \cos \theta = r \cdot \sin \varphi \cdot \cos \theta$$

$$y = d \cdot \sin \theta = r \cdot \sin \varphi \cdot \sin \theta$$

$$z = r \cdot \cos \varphi$$

Es gelten also für die Umrechnung von räumlichen Polarkoordinaten folgende Transformationsformeln:

$$\left. \begin{array}{l} x = r \cdot \sin \varphi \cdot \cos \theta \\ y = r \cdot \sin \varphi \cdot \sin \theta \\ z = r \cdot \cos \varphi \end{array} \right\} \Leftrightarrow \left\{ \begin{array}{l} r = \sqrt{x^2 + y^2 + z^2} \\ \theta = \arctan\left(\frac{y}{x}\right) \\ \varphi = \arctan \frac{\sqrt{x^2 + y^2}}{z} \end{array} \right.$$

Es gelten dabei die Wertebereiche: $0 \leq r < \infty$

$0 \leq \varphi \leq \pi$

$0 \leq \theta \leq 2\pi$ oder $-\pi \leq \theta \leq \pi$

Beispiel 3: Zylinderkoordinaten

Zylinderkoordinaten finden überall dort ihre Anwendung, wo eine zylinderförmige Symmetrie vorgegeben ist. Dies ist etwa der Fall, wenn um einen geraden Draht das Magnetfeld berechnet werden soll.

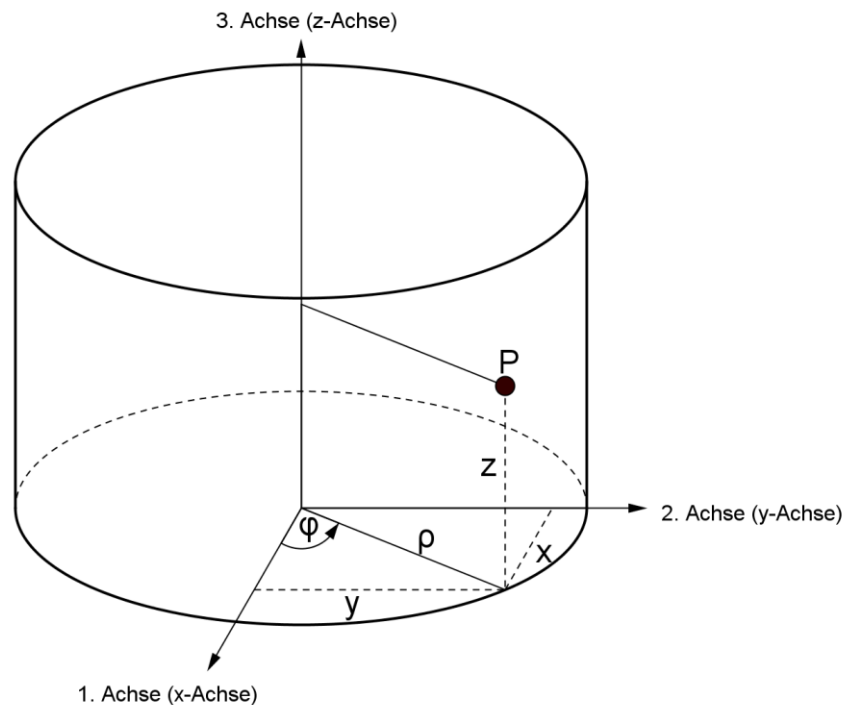


Abbildung 20: Zylinderkoordinaten

Aus der Abbildung 20 ergibt sich unmittelbar

$$\rho = \sqrt{x^2 + y^2} \quad \text{und}$$

$$\tan \varphi = \frac{y}{x} \Rightarrow \varphi = \arctan\left(\frac{y}{x}\right)$$

Umgekehrt gilt:

$$x = \rho \cdot \cos \varphi$$

$$y = \rho \cdot \sin \varphi$$

Es gelten also für die Umrechnung von Zylinderkoordinaten folgende Transformationsformeln:

$$\left. \begin{array}{l} x = \rho \cdot \cos \varphi \\ y = \rho \cdot \sin \varphi \\ z = z \end{array} \right\} \Leftrightarrow \left\{ \begin{array}{l} \rho = \sqrt{x^2 + y^2} \\ \varphi = \arctan\left(\frac{y}{x}\right) \\ z = z \end{array} \right.$$

Es gelten die Wertebereiche: $0 \leq \rho < \infty$

$$0 \leq \varphi \leq 2\pi$$

$$-\infty < z < \infty$$

3.5.2 Sinus, Cosinus, Tangens

Die trigonometrischen Funktionen Sinus, Cosinus, Tangens und Kotangens⁷⁵ gehören zu den grundlegenden Funktionen, die vielfältigste Anwendungsmöglichkeiten beinhalten. Angefangen von den rechtwinkligen Dreiecken, durch die sie zu anfangs in der Schule definiert werden⁷⁶, können mittels des Sinus- und Kosinussatzes beliebige Dreiecke und darauf aufbauend beliebige geradlinig begrenzte Figuren berechnet werden. In der Physik spielen sie natürlich in sehr unterschiedlichen Bereichen eine sehr wichtige Rolle, insbesondere sei auf die mathematische Beschreibung von Schwingungen und Wellen hingewiesen.

Beispiel 1: Entfernungsbestimmung von Sternen durch die Parallaxenmethode

Um die Entfernung von Sternen bis zu einer Entfernung von etwa 1000 Lichtjahren bestimmen zu können, bedient man sich der sogenannten *Parallaxenmethode*.⁷⁷ Diese beruht darauf, dass Sterne im Jahresverlauf eine scheinbare Bewegung ausführen, beobachtet und gemessen bezüglich des Fixsternhintergrundes⁷⁸. Die Sterne vollführen eine scheinbare Bewegung deswegen, weil nicht die Sterne selbst sich bewegen, sondern die Erde im Laufe eines Jahres durch ihre Bewegung um die Sonne sich bewegt

⁷⁵ Der Kotangens gehört aber nicht mehr zum neuen Lehrplan, in älteren Lehrplänen war er noch enthalten.

⁷⁶ Diese Definition über rechtwinklige Dreiecke als Verhältnisse der Seiten zueinander ist nur eine, wenn auch natürlich sehr anschauliche Möglichkeit der Definition. Allgemeiner werden die trigonometrischen Funktionen über ihre Reihenentwicklungen definiert, wodurch der geometrische Aspekt verloren geht. Der Vorteil dieser Definition ist es, dass etwa Beziehungen zwischen den trigonometrischen Funktionen auf einfachere Art hergeleitet werden können, dass Näherungswerte bestimmt werden können und vor allem, dass, wenn man in die komplexe Analysis übergeht, der Zusammenhang zu den Exponentialfunktionen hergestellt werden kann. Die Definitionen für Sinus und Cosinus durch Reihenentwicklungen lauten wie folgt:

$$\sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} \mp \dots = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n+1}}{(2n+1)!} \quad \cos x = 1 - \frac{x^2}{2!} + \frac{x^4}{4!} \mp \dots = \sum_{n=0}^{\infty} (-1)^n \frac{x^{2n}}{(2n)!}$$

Die Reihenentwicklungen für den Tangens und den Kotangens benötigen die sogenannten Bernoullizahlen, auf die ich hier aber nicht eingehen möchte, weshalb ich die Reihenentwicklungen nicht anführe. Diese können aber in jeder mathematischen Formelsammlung nachgeschlagen werden.

⁷⁷ Um die Entfernung von Sternen zu bestimmen, bedient man sich der sogenannten jährlichen Parallaxe, da hier der Erdbahndurchmesser als Positionsverschiebung des Beobachtungsortes herangezogen werden kann. Es gibt aber auch eine tägliche Parallaxe dadurch, dass sich der Beobachtungsort durch die Erdrotation im Laufe eines Tages ändert. Für den Mond ergibt sich so ein Parallaxenwinkel von 57', das entspricht fast zwei Vollmondbreiten, für die Sonne ein Parallaxenwinkel von 8,79". Siehe dazu etwa [24] Seite 24 ff.

⁷⁸ Der Fixsternhintergrund wird durch sehr viel weiter entfernt liegende Sterne gebildet, die sich wegen ihrer Entfernung (fast) nicht zu bewegen scheinen und somit als fix angesehen werden können.

und dadurch dieser Effekt der Sternbewegung hervorgerufen wird⁷⁹. Siehe dazu Abbildung 21.

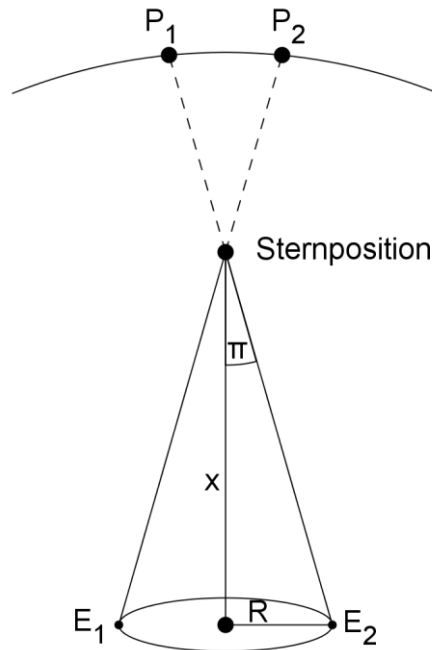


Abbildung 21: Das Zustandekommen der Sternparallaxe

E_1 und E_2 sind die Erdpositionen während eines halben Erdumlaufs um die Sonne, die Strecke R ist dabei der Erdbahnradius mit $R = 1,5 \cdot 10^{11} \text{ m}$.

Die scheinbare Position des Sterns ändert sich dabei von P_1 nach P_2 , der Winkel π ist die sogenannte (jährliche) Parallaxe des Sterns. Diese Parallaxe wird gemessen und daraus ergibt sich dann die Entfernung x des Sterns.

Aus der Abbildung 21 lässt sich unmittelbar aus der Definition des Tangens folgende Beziehung ablesen

$$\tan \pi = \frac{R}{x}$$

Und daraus folgt für die Entfernung des Sterns

⁷⁹ Es gibt natürlich auch eine tatsächliche Bewegung der Sterne zueinander, wobei diese hier aber vernachlässigt werden kann.

$$x = \frac{R}{\tan \pi}$$

Da der Winkel π sehr klein ist, kann statt dem Tangens des Winkels auch der Winkel selbst in Bogenmaß eingesetzt werden, sodass sich folgende vereinfachte Beziehung zur Bestimmung der Sternentfernung ergibt

$$x \approx \frac{R}{\pi}$$

Als Beispiel soll die Entfernung des der Sonne am nächsten gelegenen Sterns Proxima Centauri berechnet werden. Die gemessene Parallaxe beträgt $\pi \approx 0,762''$.

Umrechnung des Winkels in Bogenmaß ergibt

$$\pi_{\text{rad}} \approx \frac{0,762}{3600} \cdot \frac{2\pi}{360} \approx 3,69 \cdot 10^{-6}$$

Damit lässt sich die Entfernung des Sterns in Meter berechnen

$$x \approx \frac{1,5 \cdot 10^{11}}{3,69 \cdot 10^{-6}} \approx 4,07 \cdot 10^{16} \text{ m}$$

Da 1 Lichtjahr $9,46 \cdot 10^{15}$ Meter entspricht, ergibt sich die Entfernung von Proxima Centauri zu

$$x \approx \frac{4,07 \cdot 10^{16}}{9,46 \cdot 10^{15}} \approx 4,3 \text{ Lj}$$

Beispiel 2: Das Parsec als astronomische Entfernungsangabe

In der Astronomie werden aufgrund der großen Entfernungen der Himmelskörper eigene Längeneinheiten verwendet. Dies sind in planetarischen Bereichen die astronomische Einheit AE, die der mittleren Entfernung der Erde von der Sonne entspricht und in interstellaren und intergalaktischen Bereichen das Lichtjahr Lj, die Strecke, die das Licht in einem Jahr zurücklegt.⁸⁰

⁸⁰ Die Werte wurden [24] entnommen, Seite 535

$$1 \text{ AE} = 1,495979 \cdot 10^{11} \text{ m}$$

$$1 \text{ Lj} = 9,4608 \cdot 10^{15} \text{ m}$$

Eine weitere astronomische Längeneinheit ist das *Parsec*. 1 Parsec entspricht der Entfernung eines Sterns, der eine Parallaxe von 1" (oft gesprochen als Parallaxensekunde, deshalb auch der Name Parsec) aufweist (vergleiche Abbildung 22).

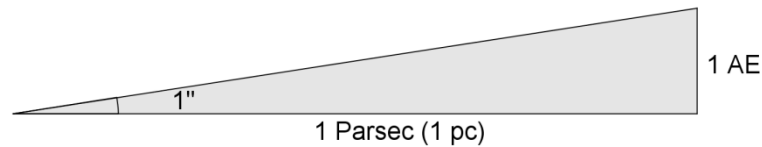


Abbildung 22: Das Parsec als astronomische Längeneinheit

1" entspricht in Bogenmaß

$$\pi_{\text{rad}} \approx \frac{1}{3600} \cdot \frac{2\pi}{360} \approx 4,85 \cdot 10^{-6}$$

Damit ergibt sich:

$$x \approx \frac{1,5 \cdot 10^{11}}{4,85 \cdot 10^{-6}} \approx 3,1 \cdot 10^{16} \text{ m} \approx \frac{3,1 \cdot 10^{16}}{9,46 \cdot 10^{15}} \approx 3,26 \text{ Lj}$$

Somit ergibt sich für die Längeneinheit Parsec folgende Umrechnung:

$$1 \text{ pc} = 3,26 \text{ Lj}$$

3.5.3 Gleichungen in zwei Variablen

Beispiel: Stoßprozesse

Stoßprozesse spielen in der gesamten Physik eine wesentliche Rolle, ob das beim Billardspiel ist oder im atomaren Bereich etwa beim Compton-Effekt. Stoßprozesse treten in der Mechanik, in der Elektrizitätslehre, in der Thermodynamik oder im Bereich der atomaren Physik auf.

Stoßprozesse sind Wechselwirkungen zwischen Teilchen, die während einer kurzen Zeitspanne Impuls und Energie austauschen können. Die Vorgänge während dem Stoßprozess selber sind überaus kompliziert und können oft im Einzelnen nicht berechnet werden, aber es können trotzdem Aussagen über die Bewegung der Teilchen vor und nach dem Stoßprozess getroffen werden. Dazu bedient man sich des Impuls- und Energieerhaltungssatzes, wonach die Gesamtimpulse und kinetischen Energien der Teilchen vor und nach dem Stoß gleich bleiben müssen, falls keine Formveränderungen stattfinden. Finden hingegen Formveränderungen statt, wird ein Teil der kinetischen Energie in innere Energie umgewandelt und die Summe der kinetischen Energien ist nicht mehr konstant. Dasselbe tritt auf, wenn kinetische Energie in Wärmeenergie umgewandelt wird, wie es etwa beim Eintritt eines Geschosses in einen Körper geschieht. Diese Stöße werden inelastische Stöße genannt und sollen in diesem Beispiel nicht näher besprochen werden.

Üblicherweise werden im Physikunterricht mehrere Spezialfälle von Stoßprozessen einzeln behandelt, da es natürlich mathematisch einfacher ist. Man kann aber auch vom allgemeinen Fall ausgehen und eine allgemeine Gleichung aufstellen und dann von dieser ausgehend Spezialfälle betrachten. Dieser Weg soll im Folgenden besprochen werden.

Es soll der elastische zentrale Stoß behandelt werden. Für einen elastischen Stoß bleiben sowohl die Summe der Impulse als auch die Summe der kinetischen Energien erhalten.

Es stoßen zwei Körper mit den Massen m_1 und m_2 und den Geschwindigkeiten v_1 und v_2 zentral aufeinander. Nach dem Stoßprozess ändern sich die Geschwindigkeiten und sollen mit v'_1 und v'_2 bezeichnet werden.

Energieerhaltung: $\frac{1}{2} m_1 v_1^2 + \frac{1}{2} m_2 v_2^2 = \frac{1}{2} m_1 v_1'^2 + \frac{1}{2} m_2 v_2'^2$

Impulserhaltung: $m_1 v_1 + m_2 v_2 = m_1 v_1' + m_2 v_2'$

Mathematisch gesehen ist das ein Gleichungssystem mit einer quadratischen und einer linearen Gleichung in 2 Variablen. Dieses Gleichungssystem soll nun gelöst werden, das heißt, die beiden Unbekannten v_1' und v_2' müssen bestimmt werden.

Die beiden Gleichungen werden zuerst etwas umgeformt:

$$\left. \begin{aligned} \frac{1}{2} m_1 (v_1^2 - v_1'^2) &= \frac{1}{2} m_2 (v_2'^2 - v_2^2) \\ m_1 (v_1 - v_1') &= m_2 (v_2' - v_2) \end{aligned} \right\} \quad (1)$$

Jetzt kann für die erste der Gleichungen eine binomische Formel angewandt werden:

$$m_1 (v_1 - v_1')(v_1 + v_1') = m_2 (v_2' - v_2)(v_2' + v_2)$$

Die zweite Gleichung kann in diese eingesetzt werden:

$$\cancel{m_2} (\cancel{v_2' - v_2}) (v_1 + v_1') = \cancel{m_2} (\cancel{v_2' - v_2}) (v_2' + v_2)$$

$$v_1 + v_1' = v_2' + v_2 \quad (2)$$

Mit der Gleichung (2) und der unveränderten 2. Gleichung aus dem Gleichungssystem (1) erhält man ein neues Gleichungssystem in 2 Unbekannten, wobei jetzt aber beide Gleichungen linear sind

$$\left. \begin{aligned} m_1 (v_1 - v_1') &= m_2 (v_2' - v_2) \\ v_1 + v_1' &= v_2' + v_2 \end{aligned} \right\} \quad (3)$$

Aus der 2. Gleichung des Gleichungssystems (3) folgt

$$v_2' = v_1 + v_1' - v_2 \quad (4)$$

Durch Einsetzen in die 1. Gleichung von (3) erhält man

$$\begin{aligned}
m_1(v_1 - v'_1) &= m_2(v_1 + v'_1 - 2v_2) \\
m_1v_1 - m_1v'_1 &= m_2v_1 + m_2v'_1 - 2m_2v_2 \\
(m_1 - m_2)v_1 + 2m_2v_2 &= (m_1 + m_2)v'_1 \\
v'_1 &= \frac{(m_1 - m_2)v_1 + 2m_2v_2}{m_1 + m_2}
\end{aligned}$$

Einsetzen der Lösung für v'_1 in (4) ergibt

$$\begin{aligned}
v'_2 &= v_1 + \frac{(m_1 - m_2)v_1 + 2m_2v_2}{m_1 + m_2} - v_2 \\
&= \frac{(m_1 + m_2)v_1 + (m_1 - m_2)v_1 + 2m_2v_2 - (m_1 + m_2)v_2}{m_1 + m_2} \\
&= \frac{2m_1v_1 + (m_2 - m_1)v_2}{m_1 + m_2}
\end{aligned}$$

Somit ist das Gleichungssystem (1) vollständig gelöst und die Lösungen können geschrieben werden als

$ \begin{aligned} v'_1 &= \frac{2m_2v_2 + (m_1 - m_2)v_1}{m_1 + m_2} \\ v'_2 &= \frac{2m_1v_1 - (m_1 - m_2)v_2}{m_1 + m_2} \end{aligned} \tag{5} $
--

Diese allgemeinen Lösungen gelten für alle Massen und alle Geschwindigkeiten und können nun auf einige interessante Spezialfälle angewandt werden.

1. Gleiche Massen: $m_1 = m_2 = m$

Aus den Lösungen (5) ergibt sich durch Einsetzen:

$$\begin{aligned}
v'_1 &= \frac{2mv_2 + (m - m)v_1}{m + m} = v_2 \\
v'_2 &= \frac{2mv_1 - (m - m)v_2}{m + m} = v_1
\end{aligned}$$

$ \begin{aligned} v'_1 &= v_2 \\ v'_2 &= v_1 \end{aligned} $
--

Bei einem elastischen Stoß, bei dem die Stoßpartner gleiche Massen haben, tauschen die Stoßpartner die Geschwindigkeiten aus, das heißt sie geben quasi ihre Geschwindigkeit an die andere Kugel weiter. Die Kugel 1 hat nach dem Stoß die Geschwindigkeit der Kugel 2 vor dem Stoß und umgekehrt.

Es können 4 Fälle⁸¹ unterschieden werden, die in der Abbildung 23⁸² dargestellt sind.

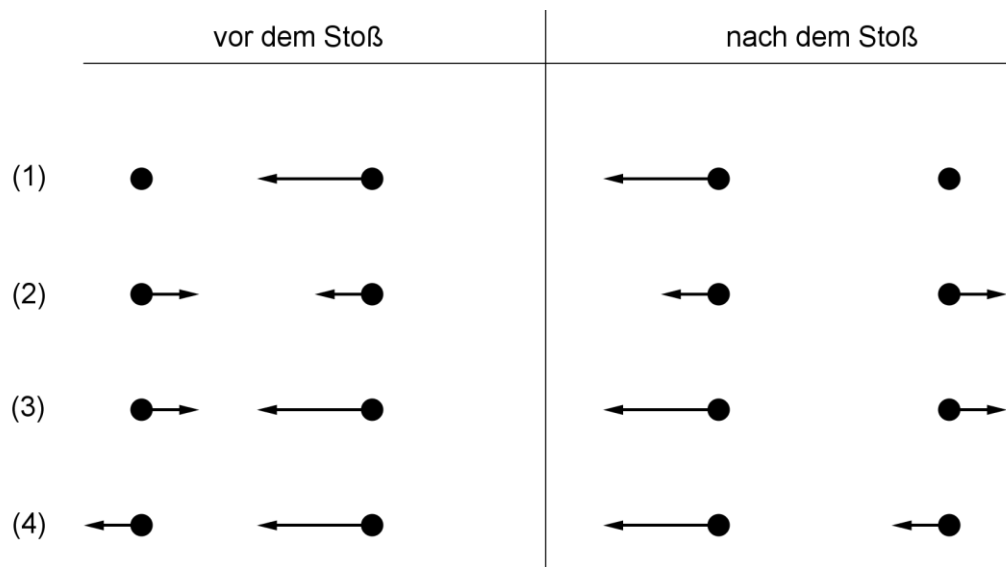


Abbildung 23: Stöße bei gleichen Massen

- (1) Die Kugel 1 ruht vor dem Stoß. Nach dem Stoß ruht die Kugel 2 und die Kugel 1 bewegt sich mit derselben Geschwindigkeit, mit der sich die Kugel 2 vor dem Stoß bewegt hat.
- (2) Die Kugeln bewegen sich mit derselben Geschwindigkeit aufeinander zu. Sie werden reflektiert und bewegen sich mit derselben Geschwindigkeit wie vorher in die andere Richtung.
- (3) Die Kugeln bewegen sich aufeinander zu, allerdings hat die Kugel 2 eine höhere Geschwindigkeit als die Kugel 1. Sie werden reflektiert und bewegen sich mit der Geschwindigkeit der jeweils anderen Kugel in die entgegengesetzte Richtung.

⁸¹ Diese 4 Fälle können sehr schön mit einem Newton'schen Kugelpendel den SchülerInnen gezeigt werden. Es sollten dabei aber nur 2 Kugeln verwendet werden.

⁸² Ein Kreis steht für die Kugel 1 und 2, die Pfeile stellen die Geschwindigkeitsvektoren der Kugeln vor und nach dem Stoß dar. Ist bei einem Kreis kein Pfeil, so bedeutet dies, dass die Kugel ruht, ihre Geschwindigkeit ist Null.

- (4) Die Kugeln bewegen sich in dieselbe Richtung. Die Kugel 2 hat eine höhere Geschwindigkeit als die Kugel 1. Nach dem Stoß bewegen sich beide Kugeln in derselben Richtung weiter, allerdings hat die Kugel 1 jetzt die höhere Geschwindigkeit und die Kugel 2 bewegt sich langsamer.

2. Ungleiche Massen: $m_1 = m$, $m_2 = 2m$

Die Massen seien so gewählt, dass die Kugel 2 die doppelte Masse der Kugel 1 hat.

Aus den Lösungen (5) ergibt sich durch Einsetzen:

$$v'_1 = \frac{4mv_2 + (-m)v_1}{3m} = \frac{4v_2 - v_1}{3}$$

$$v'_2 = \frac{2mv_1 + mv_2}{3m} = \frac{2v_1 + v_2}{3}$$

Es sollen einige spezielle Fälle betrachtet werden:

- (1) $v_1 = 0$, das heißt die leichtere Kugel ruht.

$$v'_1 = \frac{4}{3}v_2 \quad \text{und} \quad v'_2 = \frac{1}{3}v_2$$

Stößt eine doppelt so schwere Kugel 2 eine ruhende Kugel 1, so bewegen sich beide Kugeln in Richtung der Geschwindigkeit der 2. Kugel vor dem Stoß. Allerdings bewegt sich die Kugel 1 nach dem Stoß $\frac{4}{3}$ mal so schnell wie die Kugel 2 vor dem Stoß. Die Kugel 2 bewegt sich nur mehr mit einem Drittel ihrer Geschwindigkeit vor dem Stoß weiter.

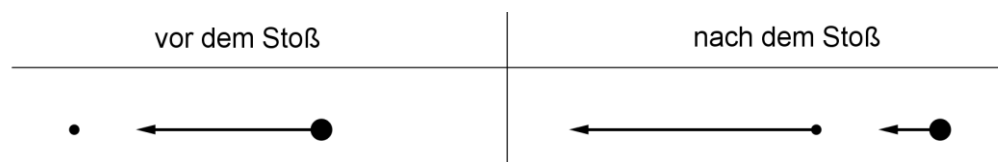


Abbildung 24: Eine doppelt so schwere Kugel stößt eine ruhende leichte Kugel

- (2) $v_2 = 0$, das heißt die schwerere Kugel ruht.

$$v'_1 = -\frac{1}{3}v_2 \text{ und } v'_2 = \frac{2}{3}v_2$$

Stößt eine halb so schwere Kugel 1 eine ruhende Kugel 2, so wird die leichtere Kugel 1 reflektiert und bewegt sich mit einem Drittel der Geschwindigkeit wie vor dem Stoß in die entgegengesetzte Richtung. Die schwerere Kugel 2 bewegt sich mit $2/3$ der Geschwindigkeit der Kugel 1 vor dem Stoß in dieselbe Richtung.



Abbildung 25: Eine halb so schwere Kugel stößt eine ruhende schwere Kugel

- (3) $v_1 = v$, $v_2 = -v$ Die Kugeln bewegen sich mit derselben Geschwindigkeit aufeinander zu.

$$v'_1 = -\frac{5}{3}v_2 \text{ und } v'_2 = \frac{1}{3}v_2$$

Bewegen sich eine Kugel 1 und eine doppelt so schwere Kugel 2 mit derselben Geschwindigkeit aufeinander zu, so werden beide Kugeln reflektiert. Die Geschwindigkeit der leichteren Kugel ist dann $5/3$ mal so groß wie zuvor, die Geschwindigkeit der schwereren Kugel nur mehr $1/3$.



Abbildung 26: Eine leichte und eine doppelt so schwere Kugel bewegen sich aufeinander zu

3. Stoß gegen eine Wand: $v_2 = 0$, $m_2 \gg m_1$

Aus den Lösungen (5) ergibt sich durch Einsetzen und Berücksichtigung, dass m_1/m_2 vernachlässigt werden kann (Nullsetzen!):

$$v_1' = \frac{2m_2 v_2 + (m_1 - m_2) v_1}{m_1 + m_2} = \frac{2v_2 + \left(\frac{m_1}{m_2} - 1\right) v_1}{\frac{m_1}{m_2} + 1} = -v_1$$

$$v_2' = \frac{2m_1 v_1 - (m_1 - m_2) v_2}{m_1 + m_2} = \frac{2\frac{m_1}{m_2} v_1 - \left(\frac{m_1}{m_2} - 1\right) v_2}{\frac{m_1}{m_2} + 1} = 0$$

$v_1' = -v_1 \text{ und } v_2' = 0$

Das kann folgendermaßen interpretiert werden:

1. Die Wand ruht auch nach dem Stoß ($v_2' = 0$), das heißt sie nimmt keine Energie auf.
2. Der Körper wird an der Wand reflektiert und hat nach dem Stoß die gleich große, aber entgegengesetzt gerichtete Geschwindigkeit. Die Richtung seines Impulses hat sich umgedreht und seine kinetische Energie ist nach dem Stoß unverändert.

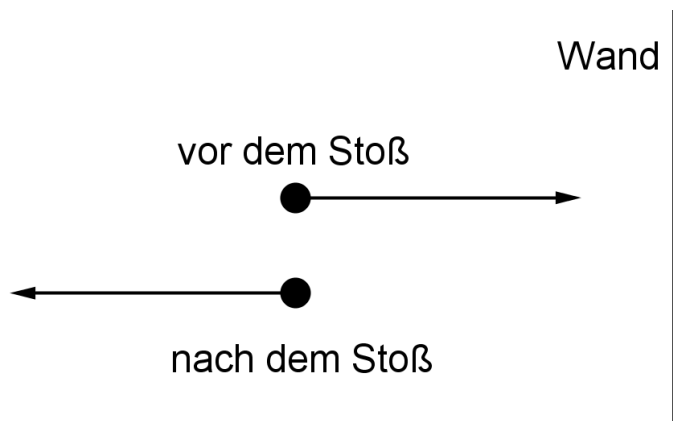


Abbildung 27: Stoß gegen eine Wand

3.5.4 Vektoren in der Ebene

Die Vektorrechnung hat ihren Ursprung in der Physik, da eine sehr große Anzahl von physikalischen Größen einen vektoriellen Charakter hat, neben den skalaren und tensoriellen Größen⁸³. Im Schulunterricht werden im Physikunterricht nur die skalaren und vektoriellen Größen behandelt, die tensoriellen bleiben unberücksichtigt. Die Anwendungsbereiche der Vektorrechnung ziehen sich in der Physik vom ersten bis zum letzten Jahr in der Oberstufe durch. Obwohl Vektoren erst in der fünften Klasse im Mathematikunterricht auftauchen, werden sie implizit oft bereits ganz am Beginn des Physikunterrichtes schon in der Unterstufe etwa bei der Behandlung der Geschwindigkeit und der Kraft erwähnt, freilich noch ohne mathematischen Hintergrund. Es wird nur qualitativ erklärt, dass bestimmte Größen wie die Geschwindigkeit und besonders natürlich Kräfte immer einen Richtungsbezug haben, der ein wichtiges Charakteristikum dieser Größen ausmacht.

In der fünften Klasse werden Vektoren im $\mathbb{R}^2$ eingeführt, während die Vektoren im $\mathbb{R}^3$ der sechsten Klasse vorbehalten bleiben. Im Grunde findet man mit den Vektoren im $\mathbb{R}^2$ im Physikunterricht das Auslangen, wenngleich es natürlich wichtige physikalische Größen gibt, die zweckmäßigerweise erst mittels des Vektorproduktes in einfacher Weise (falls der Richtungsbezug mitberücksichtigt werden soll, wie etwa beim Drehimpuls oder Drehmoment) definiert werden und deshalb auf die Vektorrechnung im $\mathbb{R}^3$ zurückgreifen müssen. Siehe hierzu das entsprechende Kapitel im 0 „Vektoren im Raum“. Besonders hier wäre es gut, wenn sich die Mathematik- und PhysiklehrerInnen einer Klasse absprechen würden, da sich besonders beim Vektorprodukt die beiden Fächer ganz wunderbar ergänzen.

Von den Rechenoperationen für Vektoren, die in der Physik wichtig sind, sind zu nennen:

- Vektoraddition
- Betrag eines Vektors
- Winkelmaß zweier Vektoren
- Orthogonalität von Vektoren
- Skalarprodukt zweier Vektoren

⁸³ Tensorielle Größen treten in der Physik häufig auf, genannt seien nur der Spannungstensor oder der Trägheitstensor als Beispiel für Tensoren 2. Stufe. Sind der Drehimpulsvektor $\vec{L}$ und der Winkelgeschwindigkeitsvektor $\vec{\omega}$ parallele Vektoren, kann die Beziehung zwischen ihnen durch die skalare Größe des Trägheitsmoments I beschrieben werden, also $\vec{L} = I \cdot \vec{\omega}$. Im allgemeinen Fall sind Drehimpulsvektor und Winkelgeschwindigkeitsvektor aber nicht mehr parallel, somit ist I auch keine skalare Größe mehr, sondern ein Tensor 2. Stufe, der als eine quadratische 3×3 Matrix geschrieben werden kann. Tensoren noch höherer Stufe haben etwa in der Allgemeinen Relativitätstheorie eine wichtige Bedeutung.

Wichtige Größen in der Physik, die zweckmäßigerweise als Vektoren definiert und geschrieben werden, sind

- Ortsvektoren
- Geschwindigkeit und Beschleunigung
- Kraft
- Impuls
- Winkelgeschwindigkeit und Winkelbeschleunigung
- Drehmoment
- Drehimpuls
- Elektrische und magnetische Feldstärke

Beispiel 1: Vektoraddition am Beispiel des senkrechten Wurfes

Ein grundlegendes Prinzip der Physik besagt, dass Größen, die vektoriellen Charakter haben, nach den Regeln der Vektoraddition addiert werden müssen. So ist etwa bei mehreren Geschwindigkeiten, denen ein Körper unterworfen ist, die resultierende Geschwindigkeit, die der Körper schlussendlich ausführt, die Vektorsumme aller auftretenden Einzelgeschwindigkeiten.

Im Falle des senkrechten Wurfes treten zwei Geschwindigkeiten auf, eine konstante Geschwindigkeit senkrecht nach oben $\vec{v}_0 = (0, v_0)$ und eine linear zunehmende Geschwindigkeit senkrecht nach unten $\vec{v}_f(t) = (0, -g \cdot t)$ ⁸⁴. Die konstante Geschwindigkeit $\vec{v}_0$ wird gemäß dem ersten Newton'schen Gesetz beibehalten und ändert sich nicht. Die resultierende Bewegung wird nur deswegen immer langsamer und kehrt sich schließlich in eine senkrechte Bewegung nach unten um, weil eine weitere Bewegung $\vec{v}_f(t)$ nach unten erfolgt, die mit der Zeit t immer schneller wird und schließlich die nach oben gerichtete Bewegung „kompensiert“, da beide Geschwindigkeiten ohne sich zu beeinflussen überlagert werden dürfen.

Als resultierende Geschwindigkeit nach der Zeit t ergibt sich schließlich:

⁸⁴ Der Vektor wäre natürlich dreidimensional zu schreiben, also als $\vec{v}_0 = (0, 0, v_0)$ beziehungsweise $\vec{v}(t) = (0, 0, -g \cdot t)$. Das ist aber hier und auch später beim schiefen Wurf nicht notwendig, da sich die Bewegungen der Körper in einer Ebene abspielen. Es ist aber sicher instruktiv, dies in einer Klasse zu thematisieren und somit den SchülerInnen den Umgang mit drei- beziehungsweise zweidimensionalen Vektoren näher zu bringen.

$$\vec{v}(t) = (0, v_0 - g \cdot t)$$

Daraus kann dann im Weiteren etwa sehr leicht der Umkehrzeitpunkt bestimmt werden, wann die nach oben gerichtete Bewegung zum Stillstand kommt und sich umkehrt, indem die y-Komponente der resultierenden Geschwindigkeit Null gesetzt wird

$$v_0 - g \cdot t = 0 \Rightarrow t = \frac{v_0}{g}.$$

Bei einer Anfangsgeschwindigkeit von $25 \text{ m} \cdot \text{s}^{-1}$ ergibt sich damit eine Zeit von

$$t_{\text{Umkehr}} = \frac{25}{9,81} \approx 2,55 \text{ s}$$

nach der der Körper seine Bewegung umkehrt beziehungsweise seinen höchsten Punkt seiner Bewegung erreicht hat. Es sei dies auch in der Abbildung 28 veranschaulicht:

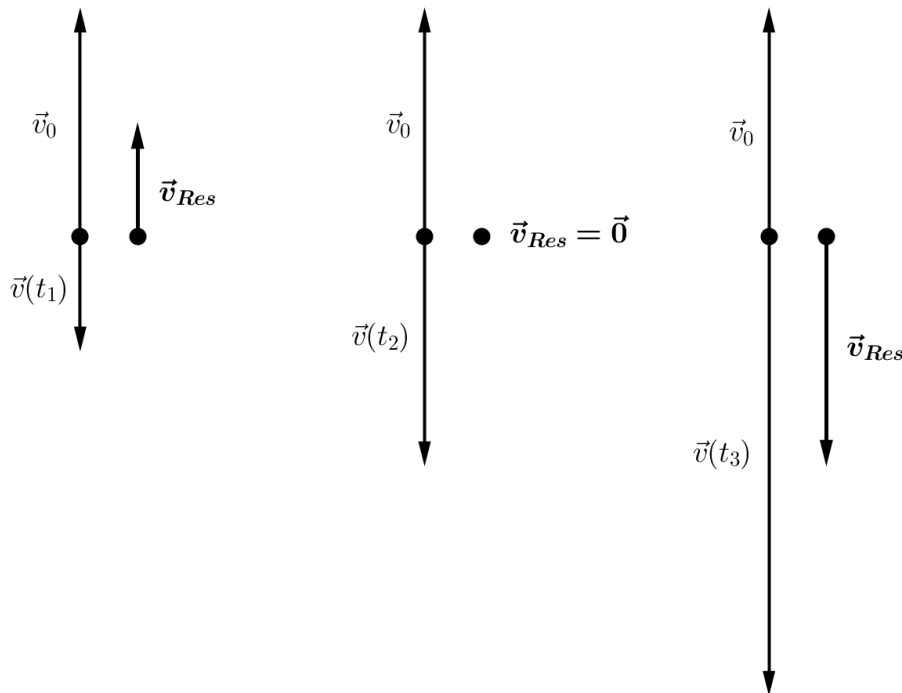


Abbildung 28: Senkrechter Wurf

Es ist jeweils die Situation zu drei verschiedenen Zeitpunkten t_1, t_2, t_3 dargestellt. Die linken Pfeile repräsentieren jeweils die Vektoren der nach oben gerichteten konstanten Geschwindigkeit $\vec{v}_0$ und der nach unten gerichteten Geschwindigkeiten $\vec{v}(t_1), \vec{v}(t_2), \vec{v}(t_3)$. Die rechten Pfeile stellen die resultierenden Geschwindigkeiten zu den Zeiten t_1, t_2, t_3 dar. Man sieht, dass zum Zeitpunkt t_1 die resultierende Geschwindigkeit durch einen Vektor, der

nach oben gerichtet ist, dargestellt wird, somit also eine Bewegung stattfindet, die zum Zeitpunkt t_1 nach oben gerichtet ist. Zum Zeitpunkt t_2 sind die beiden Vektoren $\vec{v}_0$ und $\vec{v}(t_2)$ gleich lang, aber entgegengesetzt orientiert, sodass sie sich in Summe aufheben, es gilt also $\vec{v}(t_2) = -\vec{v}_0$. Somit ist $\vec{v}(t_2) + \vec{v}_0 = \vec{0}$, es findet also keine Bewegung mehr statt, da die Geschwindigkeit in Summe der Nullvektor ist. Der Zeitpunkt t_2 ist also genau der Zeitpunkt, an dem sich die Bewegung umkehrt, somit der höchste Punkt der Bewegung erreicht wird. Zum Zeitpunkt t_3 ist der Vektor der nach unten gerichteten Geschwindigkeit länger als der nach oben gerichtete, es findet also eine Bewegung statt, die nach unten gerichtet ist, der Körper fällt.

Beispiel 2: Vektoraddition und Vektorzerlegung am Beispiel des schiefen Wurfes

Der schiefe Wurf kann wie der senkrechte Wurf behandelt werden, allerdings kommt hier noch eine Horizontalkomponente der Geschwindigkeit hinzu. Üblicherweise wird in der Physik sofort die Vektorzerlegung der Geschwindigkeiten durchgeführt und sodann nur mehr skalar gerechnet, sodass der Vektorcharakter der Bewegung für die SchülerInnen nicht wirklich erfahrbar wird und die Bewegung recht kompliziert wirkt, weil dann trigonometrische Funktionen ins Spiel kommen, die den Vektorcharakter überdecken. Natürlich ist es notwendig, die Zerlegung in eine Vertikal- und Horizontalkomponente durchzuführen, allerdings sollte zuerst allgemein darüber gesprochen werden, wie es denn zu der speziellen Bahnkurve kommt, die ein Körper beschreibt, der in einem bestimmten Winkel zur Horizontalen abgeschossen wird. Das Zustandekommen dieser Bahnkurve kann sehr schön durch eine Vektordarstellung veranschaulicht werden.

In der Abbildung 29 sind mehrere Geschwindigkeitsvektoren des Körpers zu einigen ausgewählten Zeitpunkten eingezeichnet. Die nach rechts und oben gerichteten Vektoren sind die Komponenten des konstanten Geschwindigkeitsvektors $\vec{v}_0$, der in eine horizontale und vertikale Komponente zerlegt dargestellt wird. Der nach unten gerichtete Vektor ist der Geschwindigkeitsvektor des freien Falls.

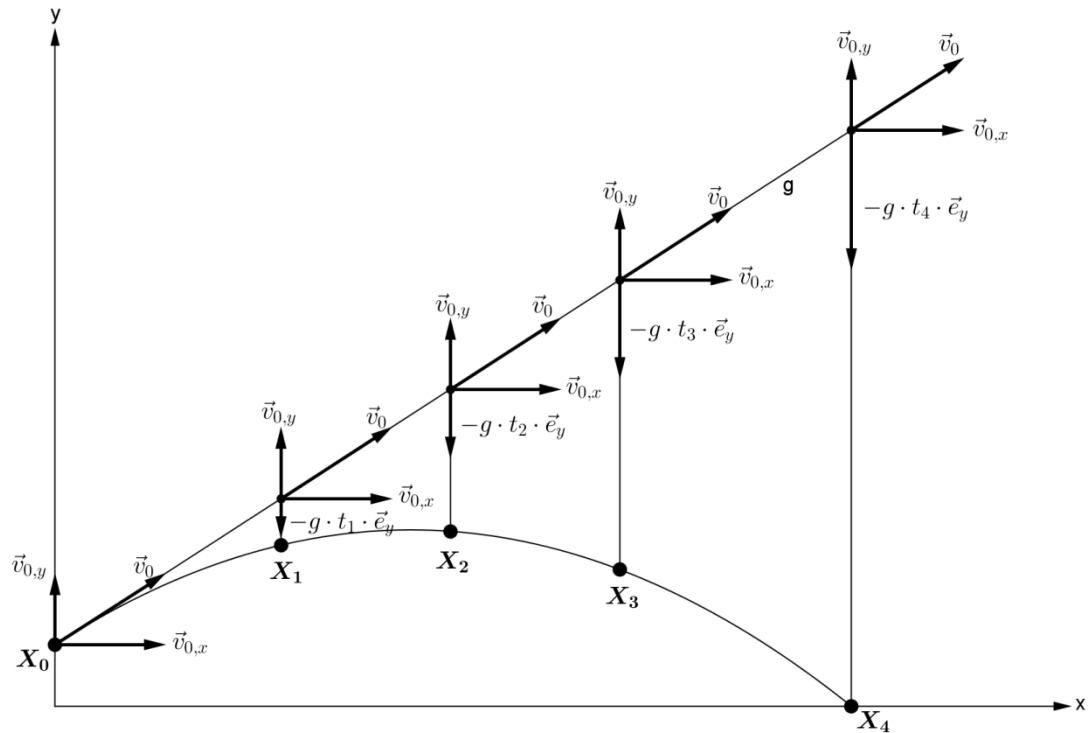


Abbildung 29: Der schiefe Wurf, Vektorzerlegung der Geschwindigkeiten

Die Gerade g stellt die Bahnkurve des Körpers dar, wenn er sich allein aufgrund des konstanten Anfangsgeschwindigkeitsvektors $\vec{v}_0$ bewegen würde⁸⁵. Im Gravitationsfeld der Erde wirkt nun aber zusätzlich die Gravitationskraft auf den Körper ein, sodass zu dieser geradlinigen Bewegung eine zweite hinzukommt, die einer in die negative y -Richtung zeigenden Geschwindigkeit $-g \cdot t_i \cdot \vec{e}_y$, die für sich genommen einen freien Fall darstellt, entspricht⁸⁶. So bewegt sich der Körper etwa bis zum Zeitpunkt t_i um die geradlinige Strecke $|\vec{v}_0 \cdot t_i|$ in Richtung des Geschwindigkeitsvektors $\vec{v}_0$ weiter und fällt dabei gleichzeitig um die Strecke $\frac{g}{2} \cdot t_i^2$. Die Punkte X_i sind dann die Punkte, die sich auf der Bahnkurve des Körpers befinden, die er tatsächlich ausführt.

Der gesamte Geschwindigkeitsvektor zu einem Zeitpunkt t lässt sich sodann schreiben als

⁸⁵ Das entspricht dem 1. Newton'schen Gesetz, nach dem sich ein Körper gleichförmig geradlinig bewegt oder im Zustand der Ruhe bleibt, solange keine äußere Kraft auf ihn einwirkt.

⁸⁶ $\vec{e}_x$ und $\vec{e}_y$ sind die Einheitsvektoren in x - und y -Richtung

$$\vec{v}(t) = \begin{pmatrix} v_0 \cdot \cos \alpha \\ v_0 \cdot \sin \alpha - g \cdot t \end{pmatrix}$$

Daraus ergeben sich die Koordinaten des Bahnpunktes X als

$$X(t) = \begin{pmatrix} v_0 \cdot \cos \alpha \cdot t \\ y_0 + v_0 \cdot \sin \alpha \cdot t - \frac{g}{2} \cdot t^2 \end{pmatrix}$$

Instruktiv für die SchülerInnen ist es sicherlich, von dieser parametrisierten Darstellung auf die Koordinatenform der Bahnkurve zu gelangen. Dazu muss lediglich der Parameter t eliminiert werden.

$$x = v_0 \cdot \cos \alpha \cdot t$$

$$y = y_0 + v_0 \cdot \sin \alpha \cdot t - \frac{g}{2} \cdot t^2$$

Aus der ersten Gleichung ergibt sich

$$t = \frac{x}{v_0 \cdot \cos \alpha} .$$

Dies wird jetzt in die zweite Gleichung eingesetzt

$$\begin{aligned} y &= y_0 + v_0 \cdot \sin \alpha \cdot \left(\frac{x}{v_0 \cdot \cos \alpha} \right) - \frac{g}{2} \cdot \left(\frac{x}{v_0 \cdot \cos \alpha} \right)^2 \\ &= y_0 + \frac{\sin \alpha}{\cos \alpha} \cdot x - \frac{g}{2 \cdot v_0^2 \cdot \cos^2 \alpha} \cdot x^2 \end{aligned}$$

Das stellt nun die Bahnkurve dar, die der Körper beschreibt. Sie kann aber noch in die Standardform einer Parabelgleichung gebracht werden:

$$\begin{aligned} y &= y_0 + \frac{\sin \alpha}{\cos \alpha} \cdot x - \frac{g}{2 \cdot v_0^2 \cdot \cos^2 \alpha} \cdot x^2 \\ y &= y_0 - \frac{g}{2 \cdot v_0^2 \cdot \cos^2 \alpha} \cdot \left(x^2 - \frac{\sin \alpha}{\cos \alpha} \cdot \frac{2 \cdot v_0^2 \cdot \cos^2 \alpha}{g} \cdot x \right) \\ &= y_0 - \frac{g}{2 \cdot v_0^2 \cdot \cos^2 \alpha} \cdot \left(x^2 - \frac{2 \cdot v_0^2 \cdot \sin \alpha \cdot \cos \alpha}{g} \cdot x \right) \end{aligned}$$

Durch quadratische Ergänzung erhält man

$$y = y_0 - \frac{g}{2 \cdot v_0^2 \cdot \cos^2 \alpha} \cdot \left(x - \frac{v_0^2 \cdot \sin \alpha \cdot \cos \alpha}{g} \right)^2 + \frac{g}{2 \cdot v_0^2 \cdot \cos^2 \alpha} \cdot \frac{v_0^4 \cdot \sin^2 \alpha \cdot \cos^2 \alpha}{g^2}$$

Und durch einige einfache Umformungen erhält man schließlich als Endergebnis:

$$y = y_0 + \frac{v_0^2 \cdot \sin^2 \alpha}{2 \cdot g} - \frac{g}{2 \cdot v_0^2 \cdot \cos^2 \alpha} \cdot \left(x - \frac{v_0^2 \cdot \sin 2\alpha}{2 \cdot g} \right)^2$$

Das entspricht einer Parabel der Form $y = a + b \cdot (x - c)^2$ mit den Koeffizienten

$$a := y_0 + \frac{v_0^2 \cdot \sin^2 \alpha}{2 \cdot g}, \quad b := -\frac{g}{2 \cdot v_0^2 \cdot \cos^2 \alpha}, \quad c := \frac{v_0^2 \cdot \sin 2\alpha}{2 \cdot g}.$$

Im Mathematikunterricht der siebten Klasse werden Parabeln der Form $y = x^2$ im Rahmen der Kegelschnitte behandelt. Aber bereits in der sechsten Klasse sollte den SchülerInnen bekannt sein, wie man eine beliebige Funktion verschiebt und streckt. Für obige Parabel bedeutet dies, dass die Standardparabel $y = x^2$ zuerst um die Konstante c in die positive x -Achse verschoben wird, sodann um den Faktor b gestreckt wird und letztendlich um die Konstante a in die positive y -Richtung verschoben wird. Durch das negative Vorzeichen des Faktors b ist die Parabel nach unten geöffnet. Damit ist gezeigt worden, dass die Bahnkurve des schiefen Wurfes einer Parabelbahn folgt!

Beispiel 3: Vektorzerlegung von Kräften am Beispiel des Pendels

Es soll die rücktreibende Kraft bei einem Pendel untersucht werden. Sei eine Pendelmasse m um den Winkel φ aus der Ruhelage ausgelenkt. Es wirkt auf die Pendelmasse in jeder Auslenkungslage sein Gewicht $\vec{G} = -m \cdot g \cdot \vec{e}_y$ in die negative y -Richtung, also senkrecht nach unten⁸⁷.

Das System sei wiederum nur zweidimensional betrachtet, was auch gerechtfertigt ist, da das Pendel in einer Ebene schwingt. Die Gewichtskraft der Pendelmasse kann nun in zwei Komponenten zerlegt werden, eine Kräftekomponente tangential zum Kreisbogen, entlang dessen sich die Pendelmasse bewegt, was wiederum bedeutet, dass diese

⁸⁷ „nach unten“ bedeutet natürlich exakter in Richtung Erdmittelpunkt gerichtet.

rücktreibende Kraft normal auf den jeweiligen Radiusvektor steht. Die zweite Komponente zeigt in Richtung des Radiusvektors und ist nach außen gerichtet. Physikalisch bedeutet diese Normalkomponente die Kräftebelastung, die auf den Stab oder die Schnur, an der die Pendelmasse befestigt ist, wirkt und damit letztlich natürlich die Kräftebelastung, die auf den Rotationspunkt beziehungsweise die Rotationsachse wirkt. Die Richtungen der Tangential- und Normalkomponente der wirkenden Kraft sind in der Abbildung 30 dargestellt.

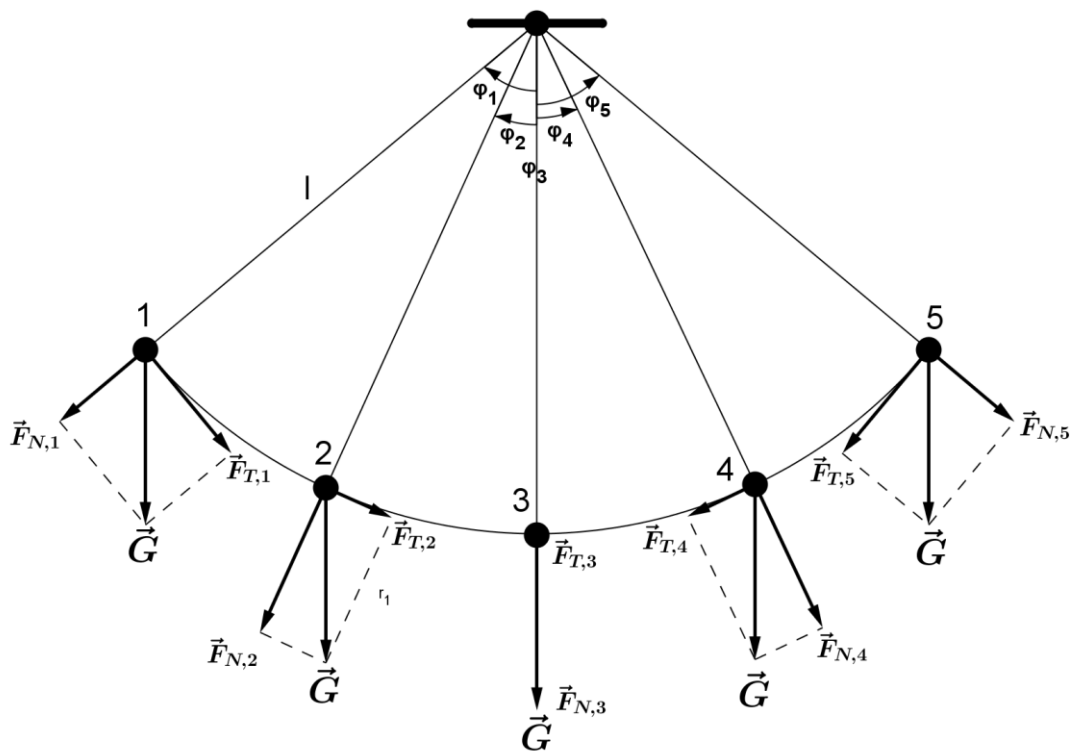


Abbildung 30: Kräftezerlegung bei einem Pendel

Man erkennt, dass die rücktreibende Kraft $\vec{F}_T$ in den beiden Umkehrpunkten am größten ist ($\vec{F}_{T,1}, \vec{F}_{T,5}$) und zur Ruhelage des Pendels ($\varphi = 0^\circ$) hin gerichtet ist. Gleichzeitig ist in diesen Punkten die Normalkomponente ($\vec{F}_{N,1}, \vec{F}_{N,5}$) am kleinsten. Die rückwirkende Kraft nimmt dann ab ($\vec{F}_{T,2}, \vec{F}_{T,4}$), bleibt aber natürlich immer in Richtung Ruhelage gerichtet und wird schließlich Null in der Ruhelage selbst. Die Normalkraft entspricht in der Ruhelage des Pendels seinem Gewicht ($\vec{F}_{N,3} = \vec{G}$).

In der Ruhelage bewegt sich die Pendelmasse gemäß dem 1. Newton'schen Gesetz aufgrund seiner Trägheit weiter, da keine Kraft entgegen der Richtung der Bewegung

wirkt, die die Bewegung der Pendelmasse behindern könnte. Die Bewegung, die die Pendelmasse ausführt, ist die einer harmonischen Schwingung mit der Schwingungsgleichung

$$\varphi(t) = \varphi_0 \cdot \cos\left(\frac{2\pi \cdot t}{T}\right)^{88}$$

A ist die Amplitude und T die Schwingungsdauer, also die Zeit, die eine volle Schwingung benötigt. Es wurde für obige Schwingungsgleichung angenommen, dass das Pendel zur Zeit $t=0$ die maximale Auslenkung $\varphi = \varphi_0$ einnimmt und keine Anfangsgeschwindigkeit hat⁸⁹.

Zur Herleitung der Periodendauer ist es notwendig, die Differentialgleichung der Pendelschwingung aufzustellen und zu lösen, was natürlich in der sechsten beziehungsweise siebten Klasse noch nicht möglich ist. Im Schulbuch Big Bang 6 im Kapitel 14.3 wird dies in Analogie zur Periodendauer des Federpendels gezeigt. Die Gleichung für die Periodendauer T des Federpendels wird durch eine Versuchsreihe bestimmt und ergibt sich zu

$$T = 2\pi \cdot \sqrt{\frac{k}{m}}$$

Dabei ist k die Federkonstante. Diese tritt immer als Proportionalitätsfaktor einer direkt proportionalen Beziehung zwischen der rücktreibenden Kraft F_r und der Auslenkung x auf.

$$F_r \sim x \Leftrightarrow F_r = k \cdot x$$

Falls also für irgendeine Bewegung die rücktreibende Kraft proportional zur Auslenkung ist, so gilt die experimentell gefundene Gleichung für die Periodendauer, es

⁸⁸ Statt der Kosinusfunktion kann immer stattdessen auch die Sinusfunktion geschrieben werden, da der Sinus dem Cosinus immer nur um $\frac{\pi}{2}$ vorläuft, also immer die Identität $\cos \alpha = \sin\left(\alpha + \frac{\pi}{2}\right)$ gilt.

⁸⁹ Im allgemeinen Fall, dass zur Zeit $t=0$ das Pendel einen beliebigen Anfangswinkel und Anfangsgeschwindigkeit einnimmt, muss ein Phasenwinkel hinzugenommen werden und die Schwingungsgleichung hat dann die Form $\varphi(t) = \varphi_0 \cdot \cos\left(\frac{2\pi \cdot t}{T} + \delta\right)$, was auch in der äquivalenten Form

$\varphi(t) = \varphi_1 \cdot \sin\left(\frac{2\pi \cdot t}{T}\right) + \varphi_2 \cdot \cos\left(\frac{2\pi \cdot t}{T}\right)$ geschrieben werden kann.

muss dann nur mehr für die jeweilige „Federkonstante“ der entsprechende physikalische Sachverhalt verwendet werden.

So ist für das Federpendel die rücktreibende Kraft, also die Tangentialkomponente der Gewichtskraft der Pendelmasse, proportional zur Auslenkung φ für kleine Winkel. Das ist folgendermaßen erkennbar, siehe hierzu die Abbildung 31.

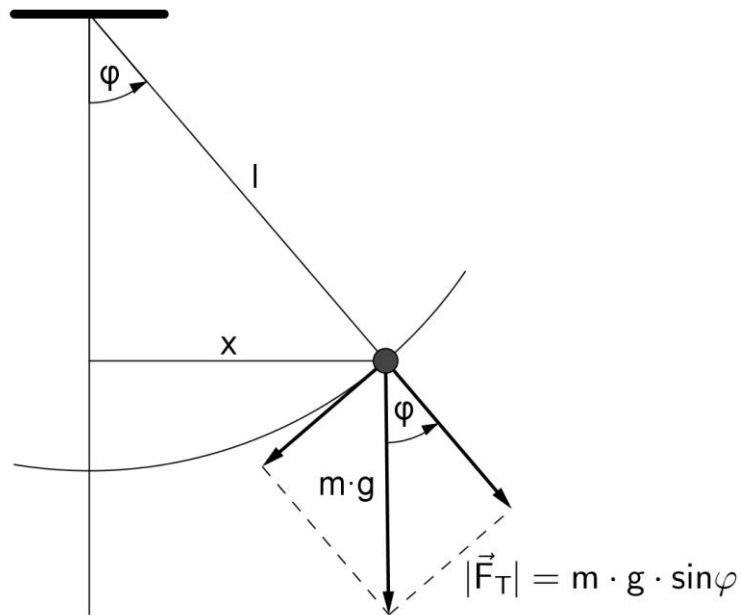


Abbildung 31: Kräftezerlegung für das Federpendel

Der Betrag der rücktreibenden Kraft ergibt sich laut Abbildung 32 als $F_T(\varphi) = m \cdot g \cdot \sin \varphi$.

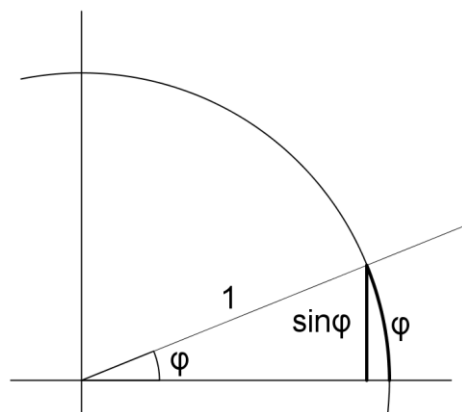


Abbildung 32: Näherung für kleine Winkel φ

Für kleine Winkel φ ist der Sinus des Winkels gleich dem Winkel, wenn φ in Bogenmaß angegeben wird. Dies erkennt man, wenn man etwa den Einheitskreis betrachtet, auf dem der Sinus des Winkel φ in Bogenmaß ungefähr der Bogenlänge des Winkels auf dem Einheitskreis entspricht, wenn der Winkel klein genug ist. siehe dazu Abbildung 32.

Also lässt sich schreiben:

$$\sin \varphi \approx \varphi \stackrel{90}{.}$$

Somit ergibt sich in dieser Näherung:

$$F_T \approx m \cdot g \cdot \varphi = m \cdot g \cdot \frac{x}{l} = k \cdot x, \text{ mit } k := \frac{m \cdot g}{l}$$

Damit ergibt sich mit $F_T = k \cdot x$ die Schwingungsdauer für das Fadenpendel als

$$T = 2\pi \cdot \sqrt{\frac{m}{k}} = 2\pi \cdot \sqrt{\frac{m}{\frac{m \cdot g}{l}}} = 2\pi \cdot \sqrt{\frac{l}{g}}.$$

$$T = 2\pi \cdot \sqrt{\frac{l}{g}}$$

⁹⁰ Man kann auch die Taylorreihe der Sinusfunktion $\sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} \mp \dots$ betrachten. Für kleine x können Reihenglieder ab dem 2. Glied vernachlässigt werden, sodass gilt: $\sin x \approx x$. Für einen Winkel $\varphi = 10^\circ$, das entspricht in Bogenmaß $\frac{\pi}{18}$ rad ergibt sich etwa: $\sin \varphi \approx 0,1736$ und $\frac{\pi}{18} \approx 0,1745$, sodass der Fehler im Promillebereich liegt und die Näherung gerechtfertigt ist.

3.6 6. Klasse

3.6.1 Exponentialfunktionen und Wachstumsprozesse

Exponentialfunktionen gehören sicherlich zu den in der Physik, aber auch in allen anderen Wissenschaften, häufigsten Funktionentypen, wenn es um Wachstums- oder Abnahmevorgänge geht, wobei ein Abnahmevorgang natürlich immer als ein negativer Wachstumsprozess angesehen werden kann.

Es ist natürlich schon interessant, warum gerade die Exponentialfunktion zur Basis e , $e \approx 2,718281828\dots$ ist dabei die Euler'sche Zahl, bei Wachstums- oder Abnahmevorgängen so häufig auftritt. Mathematisch gesehen ist das Besondere an der Exponentialfunktion zur Basis e , dass ihre Ableitung an einer Stelle x gleichzeitig ihrem Funktionswert an dieser Stelle entspricht, also gilt $f(x) = f'(x)$. Man kann auch umgekehrt zeigen, dass es genau eine Funktion gibt, die diese Bedingung, Ableitung an der Stelle x ist gleich dem Funktionswert an dieser Stelle, erfüllt.

Nun ergibt sich natürlich die Frage, weshalb diese Funktion für Wachstumsprozesse in der Natur sozusagen so sehr begünstigt ist. Es sollte meines Erachtens mit den SchülerInnen darauf eingegangen werden, weshalb dies der Fall ist, und zwar im Unterrichtsfach Mathematik genauso wie im Fach Physik, da dadurch sehr viel darüber gelernt und verstanden werden kann, wie die Natur „funktioniert“ und welcher physikalische Gehalt in der mathematischen Beschreibung von Wachstumsvorgängen enthalten ist.

In der Natur gibt es sehr viele Vorgänge, angefangen von radioaktiven Zerfallsvorgängen über die Abnahme des Luftdrucks mit der Höhe bis zur Temperaturabnahme eines Körpers mit der Zeit, die sich durch eine Funktion der Form $f(x) = c \cdot a^x$ beschreiben lassen. Mathematisch wird diese Funktion damit begründet, dass eine Zunahme des Arguments um h eine Erhöhung beziehungsweise Verminderung des Funktionswertes um den Faktor a^h bewirkt, was sich leicht zeigen lässt

$$f(x+h) = c \cdot a^{x+h} = a^h \cdot c \cdot a^x = a^h \cdot f(x) .$$

Physikalisch liegt einer Exponentialfunktion immer folgende Überlegung zugrunde. Angenommen, eine beliebige Größe G ändert sich mit der Zeit⁹¹ und die Änderung werde als ΔG bezeichnet. Das Zeitintervall, in dem die Änderung von G passiert, sei Δt . Entscheidend ist nun, dass die Änderung ΔG nun einerseits proportional zur Größe G zum Zeitpunkt t ist und andererseits auch proportional zum Zeitintervall Δt .⁹²

Dann kann man schreiben

$$\Delta G \sim G \cdot \Delta t$$

beziehungsweise, was mathematisch äquivalent ist

$$\Delta G = \pm \lambda \cdot G \cdot \Delta t$$

Dabei ist λ eine für den jeweiligen Vorgang charakteristische Konstante, wobei das positive Vorzeichen für einen Wachstumsvorgang und das negative Vorzeichen für einen Abnahmenvorgang steht.

Man formt die erhaltene Gleichung um und schreibt

$$\frac{\Delta G}{\Delta t} = \pm \lambda \cdot G$$

wobei $\frac{\Delta G}{\Delta t}$ einem Differenzenquotienten entspricht.

Mathematisch macht man nun den Grenzübergang und lässt Δt gegen Null gehen und erhält somit

$$\lim_{\Delta t \rightarrow 0} \frac{\Delta G}{\Delta t} = \frac{dG}{dt} = \pm \lambda \cdot G.$$

Das ist eine Differentialgleichung 1. Ordnung, die man durch einfaches Integrieren lösen kann.

$$\frac{dG}{G} = \pm \lambda \cdot dt \Rightarrow \int \frac{dG}{G} = \pm \lambda \cdot \int dt \Rightarrow \ln G = \pm \lambda \cdot t + c$$

Dies kann nun durch Exponentieren umgeschrieben werden als

$$e^{\ln G} = G = e^{\pm \lambda \cdot t + c} = e^c \cdot e^{\pm \lambda \cdot t} = G_0 \cdot e^{\pm \lambda \cdot t}$$

⁹¹ Statt der Zeit kann die Änderung auch mit einer Länge, einem Winkel etc. erfolgen.

⁹² Siehe hierzu das Beispiel „radioaktiver Zerfall“ im Kapitel Differentialgleichungen.

wobei k eine neue Konstante ist und $G_0 = e^c$ gesetzt wurde.

Man erhält also aus der Bedingung der Proportionalität zwischen der Änderung einer Größe und dem Produkt der Größe und der Zeit immer eine Exponentialfunktion als Lösung der Form

$$G(t) = G_0 \cdot e^{\pm \lambda \cdot t}$$

Die Konstante G_0 ist dabei der Wert der Größe G zur Zeit $t = 0$.

Beispiel: Radioaktiver Zerfall

Es finden sich in den Mathematikschulbüchern zur Exponentialfunktion einige Beispiele, insbesondere auch zum radioaktiven Zerfall. Es sollte also für den Physikunterricht kein Problem darstellen, auch einige Beispiele und Anwendungen zu den Exponentialfunktionen zu behandeln. Folgende Beispiele sind aus Mathematik verstehen entnommen.⁹³

Aufgabe 4.66

Im Jahr 1991 wurde im Gletschereis die Mumie des Mannes vom Hauslabjoch – genannt „Ötzi“ – gefunden. Eine erste Analyse von Gewebeproben ergab, dass vom ursprünglichen C-14-Anteil nur mehr 57% vorhanden waren.

- 1) Vor wie vielen Jahren ist der „Ötzi“ ungefähr gestorben?
- 2) Die Halbwertszeit von C-14 wird oft mit 5730 Jahren $\pm$ 30 Jahre angegeben. Wie wirkt sich diese Unsicherheit auf die Berechnung des Alters von „Ötzi“ aus?

- 1) Bei jeder exponentiellen Abnahme- oder Wachstumsfunktion kann die Basis beliebig gewählt werden. Es sollen zur Illustration die Basen so gewählt

⁹³ Siehe [12]

werden, dass einmal im Exponenten nur die Zeit ohne Koeffizient und einmal als Basis die Euler'sche Zahl gewählt wird.

Das Abnahmegesetz kann geschrieben werden als $N(t) = N_0 \cdot a^t$.

Die Halbwertszeit gibt an, nach welchem (beliebigen) Zeitintervall die Hälfte der zur Anfangszeit des Zeitintervalls vorhandenen Menge des radioaktiven Materials zerfallen ist, insbesondere, nach welcher Zeit die Ausgangsmenge zur Zeit $t = 0$ auf die Hälfte abgesunken ist.

$$\frac{N_0}{2} = N_0 \cdot a^{5730} \Leftrightarrow \frac{1}{2} = a^{5730}$$

Hieraus sieht man, dass die Halbwertszeit unabhängig von der Menge des Ausgangsmaterials ist.

Die Basis lässt sich nun sehr einfach berechnen zu

$$0,5 = a^{5730} \Rightarrow a = 0,5^{\frac{1}{5730}} \approx 0,99987904.$$

Damit ergibt sich das Zerfallsgesetz

$$N(t) = N_0 \cdot 0,99987904^t.$$

Laut Angabe sind nach einer gewissen, noch unbekannten Zeit t , $57\% = 0,57$ des Ausgangs-C-14 vorhanden, also $N(t) = 0,57 \cdot N_0$. Damit ergibt sich

$$0,57 \cdot N_0 = N_0 \cdot 0,99987904^t \Rightarrow 0,57 = 0,99987904^t$$

Durch Logarithmieren erhält man schließlich die Zeit, die nach dem Tod des „Ötzi“ vergangen ist:

$$\ln 0,57 = t \cdot \ln 0,99987904$$

$$t = \frac{\ln 0,57}{\ln 0,99987904} \approx 4647$$

Ötzi ist somit vor etwa 4647 Jahren gestorben.

- 2) Für die „obere“ Halbwertszeit 5730 Jahren + 30 Jahre ergibt sich

$$0,5 = a^{5760} \Rightarrow a = 0,5^{\frac{1}{5760}} \approx 0,99987967$$

und für die „untere“ Halbwertszeit

$$0,5 = a^{5700} \Rightarrow a = 0,5^{\frac{1}{5700}} \approx 0,99987840.$$

Das „obere“ und „untere“ Alter von Ötzi ergibt sich zu

$$t = \frac{\ln 0,57}{\ln 0,99987967} \approx 4671$$

und das „untere“ Alter zu

$$t = \frac{\ln 0,57}{\ln 0,99987840} \approx 4622$$

Als Ergebnis ergibt sich, dass das Alter von „Ötzi“ zwischen 4622 und 4671 Jahren liegt, also bis auf eine Unsicherheit von 49 Jahren genau bestimmt werden kann.

Man kann auch alternativ folgendes Zerfallsgesetz mit der Euler'schen Zahl als Basis annehmen:

$$N(t) = N_0 \cdot e^{\lambda \cdot t}$$

Die Bestimmung der Zerfallskonstanten λ ergibt mit der Halbwertszeit

$$\frac{N_0}{2} = N_0 \cdot e^{\lambda \cdot 5730}, \text{ also } 0,5 = e^{\lambda \cdot 5730}.$$

Durch Logarithmieren erhält man

$$\ln 0,5 = \lambda \cdot 5730 \cdot \ln e = \lambda \cdot 5730$$

und damit

$$\lambda = \frac{\ln 0,5}{5730} \approx -0,000120968$$

Das Zerfallsgesetz lässt sich dann schreiben als

$$N(t) = N_0 \cdot e^{-0,000120968 \cdot t}$$

Für das Alter erhält man somit

$$0,57 = e^{-0,000120968 \cdot t} \Rightarrow \ln 0,57 = -0,000120968 \cdot t$$

$$t = \frac{\ln 0,57}{-0,000120968} \approx 4647$$

Wie zuvor erhält man für das Alter von Ötzi etwa 4647 Jahre.⁹⁴

Der Verlauf der erhaltenen Zerfallskurve ist in Abbildung 33 gezeichnet.

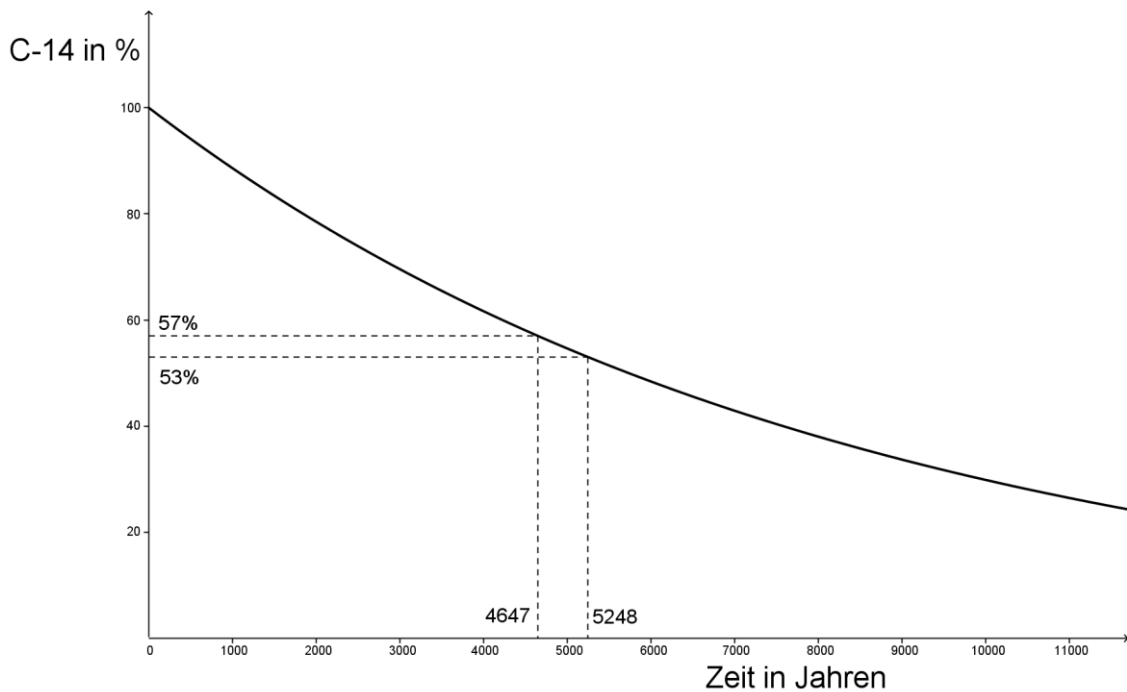


Abbildung 33: Zerfallsgesetz für die Altersbestimmung von „Ötzi“

⁹⁴ Laut neueren Daten sind noch etwa 53% vom ursprünglichen C-14-Anteil vorhanden. Damit erhöht sich das Alter von Ötzi auf etwa $t = \frac{\ln 0,53}{-0,000120968} \approx 5248$ Jahre.

3.6.2 Trigonometrische Funktionen

Beispiel 1: Schwingungen und Schwingungsgleichung

Die Physik der Schwingungen ist grundlegend für das Verständnis sehr viele physikalischer Vorgänge. Das reicht von mechanischen Schwingungen wie die Schwingungen eines Pendels bis zur Akustik, wo es Druckunterschiede in der Luft sind, die periodische Schwingungsvorgänge ausführen. Das Licht ist natürlich in der klassischen Optik das Paradebeispiel für Schwingungsvorgänge. Aber Schwingungen treten auch auf, wenn etwa Erdbebenwellen untersucht werden oder Moleküle Schwingungen ausführen. Allein aus diesen wenigen Beispielen sollte ersichtlich sein, dass die Schwingungs- und Wellenlehre in der Physik nicht wichtig genug eingeschätzt werden kann.

In diesem Beispiel sollen einfache Schwingungsformen vorgestellt werden und die Mathematik dazu erklärt werden. Periodische Schwingungen mit sinusförmiger Charakteristik treten immer dann auf, wenn eine die rücktreibende Kraft eines Vorganges proportional zu seiner Auslenkung ist.⁹⁵

Sei $y(t)$ die Auslenkung eines beliebigen periodischen Vorganges⁹⁶ und t die Zeit. Die Schwingungsgleichung ist dann allgemein von der Form

$$y(t) = A \cdot \sin(\omega \cdot t + \varphi)$$

Dabei sind A die sogenannte *Amplitude*, ω die sogenannte *Kreisfrequenz* und φ der Phasenwinkel oder kurz *Phase* der Schwingung. Es soll kurz die Bedeutung dieser beiden Größen besprochen werden.

Da ein Vorgang, der periodisch schwingt, sich nach einer gewissen Zeit T , diese Zeitdauer wird Schwingungsdauer oder Periodendauer der Schwingung genannt, wiederholt, muss gelten

⁹⁵ Siehe hierzu auch Kapitel 3.5.4, Beispiel 3

⁹⁶ $y(t)$ kann etwa eine Länge sein wie beim Federpendel, ein Winkel wie beim Fadenpendel, ein Druckunterschied bei einer Schallschwingung, eine Spannung oder eine elektrische Stromstärke in einem Stromkreis, eine elektrische und magnetische Feldstärke bei Licht oder sonst eine beliebige Größe, die die Auslenkung bei einem periodischen Vorgang beschreibt. Die Schwingungslehre ist dann für all diese Vorgänge anwendbar und somit sehr allgemein in ihren Anwendungsmöglichkeiten.

$$y(t) = y(t + T) = A \cdot \sin(\omega \cdot (t + T) + \varphi)$$

Da die Sinusfunktion die Periode 2π hat, muss auch gelten

$$\sin(\omega t + \varphi) = \sin(\omega t + \varphi + 2\pi)$$

Wenn man diese beiden Bedingungen verknüpft, erhält man folgenden Zusammenhang zwischen der Kreisfrequenz und der Periodendauer einer Schwingung:

$$\omega \cdot (t + T) + \varphi = \omega \cdot t + \varphi + 2\pi$$

$$\omega \cdot t + \omega \cdot T + \varphi = \omega \cdot t + \varphi + 2\pi$$

$$\omega \cdot T = 2\pi$$

$$\omega = \frac{2\pi}{T}$$

Benötigt eine vollständige Schwingung die Zeit T Sekunden, so werden dann in einer Sekunde $\frac{1}{T}$ Schwingungen ausgeführt. Diese Größe f wird die Frequenz einer Schwingung genannt, ihre Einheit ist s^{-1} oder Hertz (Hz).

Somit hängen bei einer harmonischen Schwingung die Kreisfrequenz, die Periodendauer und die Frequenz der Schwingung wie folgt zusammen:

$$\omega = \frac{2\pi}{T} = 2\pi \cdot f$$

Die Schwingungsgleichung lautet dann

$$y(t) = A \cdot \sin\left(\frac{2\pi \cdot t}{T} + \varphi\right) = A \cdot \sin(2\pi f \cdot t + \varphi)$$

Die Bedeutung der Phase erkennt man, wenn man die Auslenkung zur Zeit $t = 0$ betrachtet:

$$y(t = 0) = A \cdot \sin \varphi$$

Die Phase bestimmt also die Auslenkung zur Zeit $t = 0$, andererseits ist die Phase eine additive Konstante des Arguments des Sinus. Geometrisch bedeutet das, dass die Funktion um den Phasenwinkel φ in der x-Richtung verschoben wird, im Falle eines positiven Phasenwinkels nach links, ansonsten nach rechts.

Da der Wert des Sinus im Intervall $[-1, +1]$ liegt, gibt die Amplitude den Maximal- beziehungsweise Minimalwert der Auslenkung an.

$$\begin{aligned}\max(y(t)) &= A \\ \min(y(t)) &= -A\end{aligned}$$

Die Schwingungsgleichung kann auch in der folgenden äquivalenten Form geschrieben werden:

$$y(t) = A_1 \cdot \sin(\omega \cdot t) + A_2 \cdot \cos(\omega \cdot t)$$

Um zu zeigen, dass diese beiden Darstellungen der Schwingungsgleichung wirklich äquivalent sind, kann man folgendermaßen vorgehen. Für den Sinus gilt das folgende Additionstheorem:

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

Angewandt auf die Schwingungsgleichung ergibt dies

$$\begin{aligned}y(t) &= A \cdot \sin(\omega \cdot t + \varphi) \\ &= \underbrace{A \cdot \cos \varphi}_{A_1} \cdot \sin(\omega \cdot t) + \underbrace{A \cdot \sin \varphi}_{A_2} \cdot \cos(\omega \cdot t) \\ &= A_1 \cdot \sin(\omega \cdot t) + A_2 \cdot \cos(\omega \cdot t)\end{aligned}$$

Umgekehrt gilt

$$A_1^2 + A_2^2 = A^2 \cdot \cos^2 \varphi + A^2 \cdot \sin^2 \varphi = A^2 \cdot (\cos^2 \varphi + \sin^2 \varphi) = A^2, \text{ also } A = \sqrt{A_1^2 + A_2^2}$$

und

$$\frac{A_2}{A_1} = \frac{A \cdot \sin \varphi}{A \cdot \cos \varphi} = \tan \varphi, \text{ also } \varphi = \arctan\left(\frac{A_2}{A_1}\right)$$

Somit ergibt sich zusammengefasst

$$A \cdot \sin(\omega \cdot t + \varphi) \Rightarrow A_1 \cdot \sin(\omega \cdot t) + A_2 \cdot \cos(\omega \cdot t)$$

$$\text{mit } A_1 = A \cdot \cos \varphi \text{ und } A_2 = A \cdot \sin \varphi$$

$$A_1 \cdot \sin(\omega \cdot t) + A_2 \cdot \cos(\omega \cdot t) \Rightarrow A \cdot \sin(\omega \cdot t + \varphi)$$

$$\text{mit } A = \sqrt{A_1^2 + A_2^2} \text{ und } \varphi = \arctan\left(\frac{A_2}{A_1}\right)$$

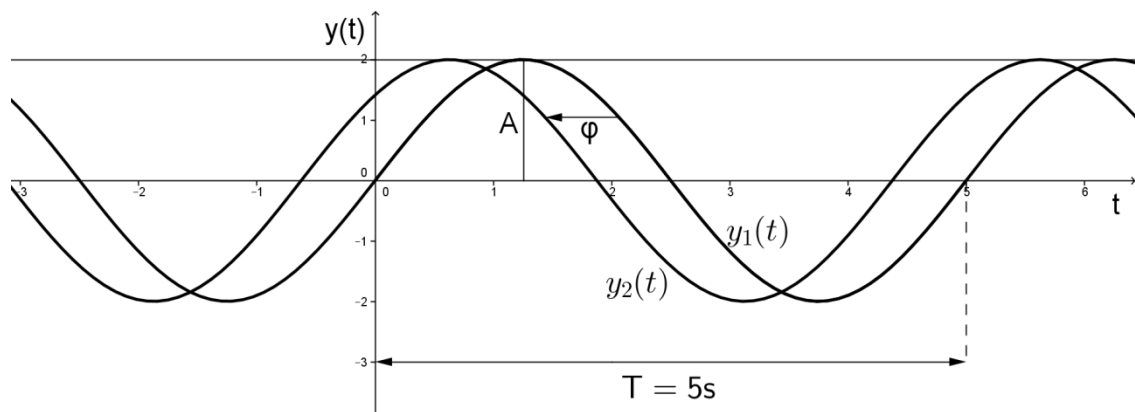


Abbildung 34: Periodendauer, Amplitude und Phase zweier harmonischen Schwingungen

In der Abbildung 34 sind die beiden Schwingungen $y_1(t) = 2 \cdot \sin\left(\frac{2\pi \cdot t}{5}\right)$ und $y_2(t) = 2 \cdot \sin\left(\frac{2\pi \cdot t}{5} + \frac{\pi}{4}\right)$ gezeichnet. Die Phase der Schwingung $y_1(t)$ ist Null und somit auch die Auslenkung $y_1(t=0)$, deshalb geht die Schwingung durch den Nullpunkt. Die Periodendauer und damit die Dauer einer vollständigen Schwingung beträgt $T = 5\text{s}$. In der Abbildung ist eine Periode der Schwingung zur besseren Kenntlichmachung farblich hervorgehoben. Es ist auch die Amplitude $A = 2$ als Maxima und Minima der Schwingung zu sehen.

Die zweite Schwingung $y_2(t)$ hat dieselbe Periodendauer und Amplitude wie die Schwingung $y_1(t)$. Allerdings hat sie einen Phasenwinkel $\varphi = \frac{\pi}{4}$. Da der Phasenwinkel positiv ist, ist die Schwingung $y_2(t)$ um $\frac{\pi}{4}$ in die negative x-Richtung verschoben, sie eilt also der Schwingung $y_1(t)$ hinterher.

Es soll auch graphisch veranschaulicht werden, wie die Schwingung $y(t) = 3 \cdot \sin\left(2\pi \cdot t - \frac{\pi}{3}\right)$ aus zwei Schwingungen mit dem Phasenwinkel $\varphi = 0^\circ$ superponiert werden kann. Siehe dazu die Abbildung 35.

Aus den hergeleiteten Formeln ergeben sich folgende Amplituden:

$$A_1 = A \cdot \cos \varphi = 3 \cdot \cos \frac{\pi}{3} = \frac{3}{2} = 1,5$$

$$A_2 = A \cdot \sin \varphi = 3 \cdot \sin \frac{\pi}{3} = \frac{3\sqrt{3}}{2} \approx 2,60$$

Somit lässt sich schreiben

$$y(t) = 3 \cdot \sin\left(2\pi \cdot t + \frac{\pi}{3}\right) = \frac{3}{2} \cdot \sin(2\pi \cdot t) + \frac{3\sqrt{3}}{2} \cdot \cos(2\pi \cdot t)$$

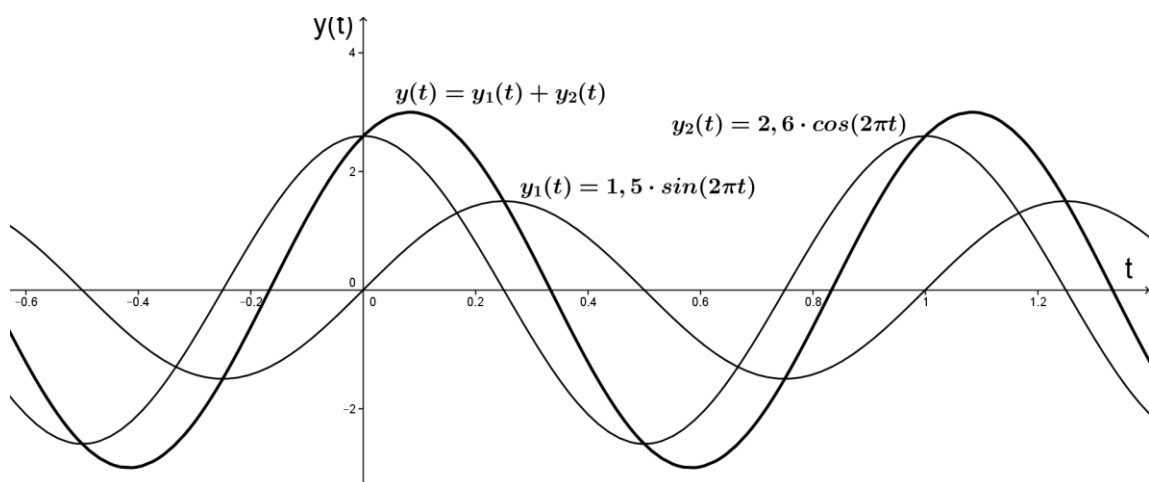


Abbildung 35: Superposition einer Schwingung

Beispiel 2: Interferenzen

Im vorigen Beispiel wurde gezeigt, dass die Überlagerung von zwei Schwingungen, in diesem Fall waren das eine reine Sinus- und Kosinusschwingung mit derselben Periodendauer und dem Phasenwinkel $\varphi = 0$, zu einer harmonischen Schwingung geführt haben. Wir wollen das in diesem Beispiel verallgemeinern auf den Fall von zwei Schwingungen mit derselben Periodendauer, aber verschiedenen Phasenwinkeln. Das

wird uns dann zu dem interessanten und wichtigen Phänomen der konstruktiven und destruktiven Interferenz führen.

Nehmen wir zwei Schwingungen mit derselben Kreisfrequenz ω und den Amplituden A und B und den verschiedenen Phasenwinkeln φ_1 und φ_2 . Die beiden Schwingungen lassen sich dann schreiben als

$$y_1(t) = A \cdot \sin(\omega \cdot t + \varphi_1) \text{ und } y_2(t) = B \cdot \sin(\omega \cdot t + \varphi_2)$$

Für die Überlagerung beider Schwingungen addieren wir diese einfach

$$y(t) = y_1(t) + y_2(t) = A \cdot \sin(\omega \cdot t + \varphi_1) + B \cdot \sin(\omega \cdot t + \varphi_2)$$

Wir wollen die resultierende Schwingung $y(t)$ in einer anderen Form schreiben und verwenden dazu wieder das Additionstheorem für den Sinus:

$$\begin{aligned} y(t) &= A \cdot \sin(\omega \cdot t) \cdot \cos \varphi_1 + A \cdot \cos(\omega \cdot t) \cdot \sin \varphi_1 + B \cdot \sin(\omega \cdot t) \cdot \cos \varphi_2 + B \cdot \cos(\omega \cdot t) \cdot \sin \varphi_2 \\ &= \sin(\omega \cdot t) \cdot (A \cdot \cos \varphi_1 + B \cdot \cos \varphi_2) + \cos(\omega \cdot t) \cdot (A \cdot \sin \varphi_1 + B \cdot \sin \varphi_2) \end{aligned}$$

Dieses Ergebnis können wir jetzt wieder zusammenfassen und erhalten eine harmonische Schwingung, die eine veränderte Amplitude A_s und einen Phasenwinkel φ_s hat, die Kreisfrequenz sich aber gegenüber den ursprünglichen Schwingungen nicht verändert hat.

$$y(t) = A_s \cdot \sin(\omega \cdot t + \varphi_s)$$

Die Amplitude A_s und der Phasenwinkel φ_s ergeben sich wieder aus den Formeln aus Beispiel 1 zu:

$$A_s = \sqrt{(A \cdot \cos \varphi_1 + B \cdot \cos \varphi_2)^2 + (A \cdot \sin \varphi_1 + B \cdot \sin \varphi_2)^2}$$

Durch Ausquadrieren ergibt sich

$$A_s = \sqrt{A^2 \cos^2 \varphi_1 + 2AB \cos \varphi_1 \cdot \cos \varphi_2 + B^2 \cos^2 \varphi_2 + A^2 \sin^2 \varphi_1 + 2AB \sin \varphi_1 \cdot \sin \varphi_2 + B^2 \sin^2 \varphi_2}$$

Dieser Ausdruck vereinfacht sich durch Verwendung von $\sin^2 \alpha + \cos^2 \alpha = 1$ zu

$$A_s = \sqrt{A^2 + B^2 + 2AB \cdot \cos \varphi_1 \cdot \cos \varphi_2 + 2AB \cdot \sin \varphi_1 \cdot \sin \varphi_2}$$

Verwendet man noch das Additionstheorem

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

so ergibt sich für die Amplitude der Summenschwingung

$$A_s = \sqrt{A^2 + B^2 + 2AB \cdot \cos(\varphi_1 - \varphi_2)}$$

Der Phasenwinkel ergibt sich einfach durch Einsetzen als

$$\varphi_s = \arctan\left(\frac{A \cdot \sin \varphi_1 + B \cdot \sin \varphi_2}{A \cdot \cos \varphi_1 + B \cdot \cos \varphi_2}\right)$$

Zusammenfassend lässt sich schreiben:

Werden zwei harmonische Schwingungen der Form $y_1(t) = A \cdot \sin(\omega \cdot t + \varphi_1)$ und $y_2(t) = A \cdot \sin(\omega \cdot t + \varphi_2)$ überlagert, so erhält man als überlagerte Schwingung wieder eine harmonische Schwingung $y(t) = A_s \cdot \sin(\omega \cdot t + \varphi_s)$ mit der neuen Amplitude

$$A_s = \sqrt{A^2 + B^2 + 2AB \cdot \cos(\varphi_1 - \varphi_2)} \quad \text{und der neuen Phase}$$

$$\varphi_s = \arctan\left(\frac{A \cdot \sin \varphi_1 + B \cdot \sin \varphi_2}{A \cdot \cos \varphi_1 + B \cdot \cos \varphi_2}\right).$$

Ein Zahlenbeispiel mag dies erläutern (siehe Abbildung 36):

$$y_1(t) = \sin(2\pi \cdot t) \quad \text{und} \quad y_2(t) = \sin\left(2\pi \cdot t + \frac{\pi}{3}\right)$$

Für die Überlagerung der Schwingungen erhält man die Amplitude

$$A_s = \sqrt{2 \cdot \left(1 + \cos\left(-\frac{\pi}{3}\right)\right)} = \sqrt{2 \cdot \left(1 + \frac{1}{2}\right)} = \sqrt{3} \quad \text{und den Phasenwinkel}$$

$$\varphi_s = \arctan \left(\frac{\sin 0 + \sin \frac{\pi}{3}}{\cos 0 + \cos \frac{\pi}{3}} \right) = \arctan \left(\frac{\frac{\sqrt{3}}{2}}{1 + \frac{1}{2}} \right) = \arctan \left(\frac{\sqrt{3}}{3} \right) \approx 0,52$$

Somit die resultierende Schwingung $y(t) \approx 1,93 \cdot \sin(2\pi \cdot t + 1,31)$

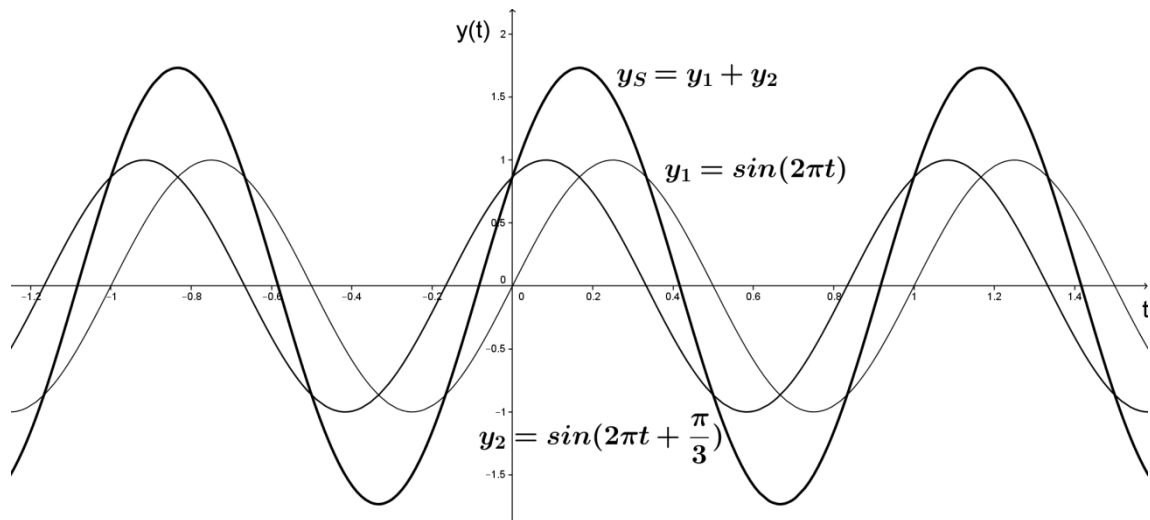


Abbildung 36: Überlagerung zweier Sinusschwingungen mit verschiedener Phase

Meist interessiert man sich weniger für den Phasenwinkel der Summenschwingung als vielmehr für die Amplitude $A_s = \sqrt{A^2 + B^2 + 2AB \cdot \cos(\varphi_1 - \varphi_2)}$.

Der Cosinus in der Formel für die Amplitude kann den Maximalwert $+1$ und den Minimalwert -1 annehmen. Hieraus ergeben sich zwei überaus interessante Spezialfälle:

1. Konstruktive Interferenz

Der Maximalwert $+1$ bedeutet

$$\cos(\varphi_1 - \varphi_2) = 1 \Leftrightarrow \varphi_1 = \varphi_2.$$

Die beiden Phasenwinkel müssen dann gleich groß sein, man spricht in diesem Fall von *gleichphasig*.

Die Amplitude hat dann den Wert

$$A_s = \sqrt{A^2 + B^2 + 2AB} = \sqrt{(A+B)^2} = A+B,$$

das heißt, die Amplituden addieren sich dann einfach.

Auch hierzu ein Zahlenbeispiel (siehe Abbildung 37):

$$y_1(t) = \sin\left(2\pi \cdot t + \frac{\pi}{2}\right) \text{ und } y_2(t) = 2 \cdot \sin\left(2\pi \cdot t + \frac{\pi}{3}\right)$$

Für die Überlagerung der Schwingungen erhält man die Amplitude

$$A_s = A+B=3$$

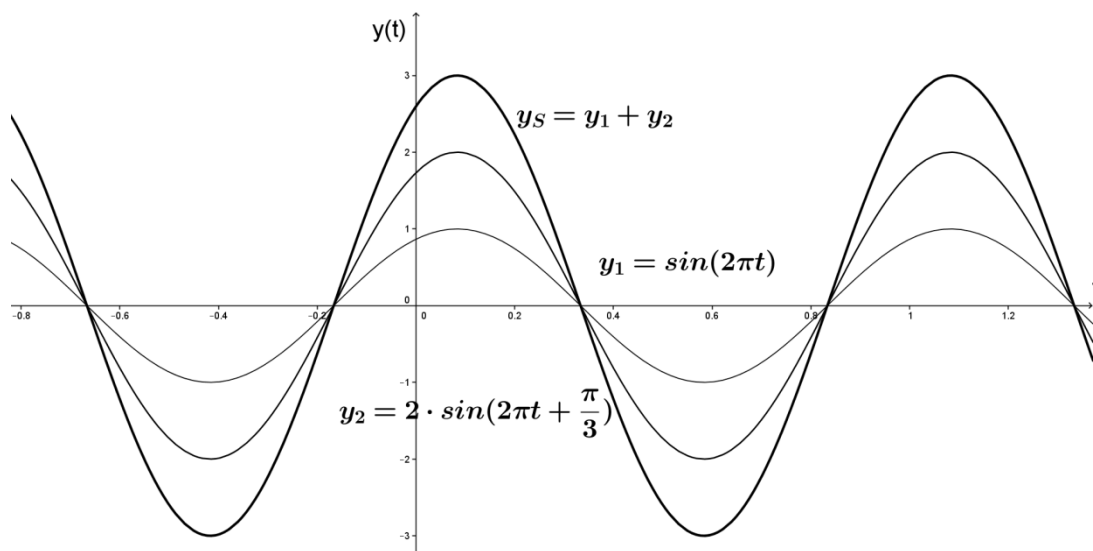


Abbildung 37: Konstruktive Interferenz

2. Destruktive Interferenz

Der Minimalwert -1 bedeutet

$$\cos(\varphi_1 - \varphi_2) = -1 \Leftrightarrow \varphi_1 - \varphi_2 = \pi.$$

Die beiden Phasenwinkel müssen sich um den Winkel π unterscheiden, man spricht in diesem Fall von *gegenphasig*.

Die Amplitude hat dann den Wert $A_s = \sqrt{A^2 + B^2 - 2AB} = \sqrt{(A-B)^2} = A-B$, die

Amplituden subtrahieren sich einfach. Im Falle gleich großer Amplituden erfolgt sogar völlige Auslöschung der Summenschwingung.

Auch hierzu ein Zahlenbeispiel (siehe Abbildung 38):

$$y_1(t) = 3 \cdot \sin\left(2\pi \cdot t + \frac{\pi}{2}\right) \text{ und } y_2(t) = 2,8 \cdot \sin\left(2\pi \cdot t - \frac{\pi}{6}\right)$$

Man sieht die Phasendifferenz ist $\varphi_1 - \varphi_2 = \frac{\pi}{3} - \left(-\frac{2\pi}{3}\right) = \pi$

Die Amplitude ist dann $A - B = 0,2$

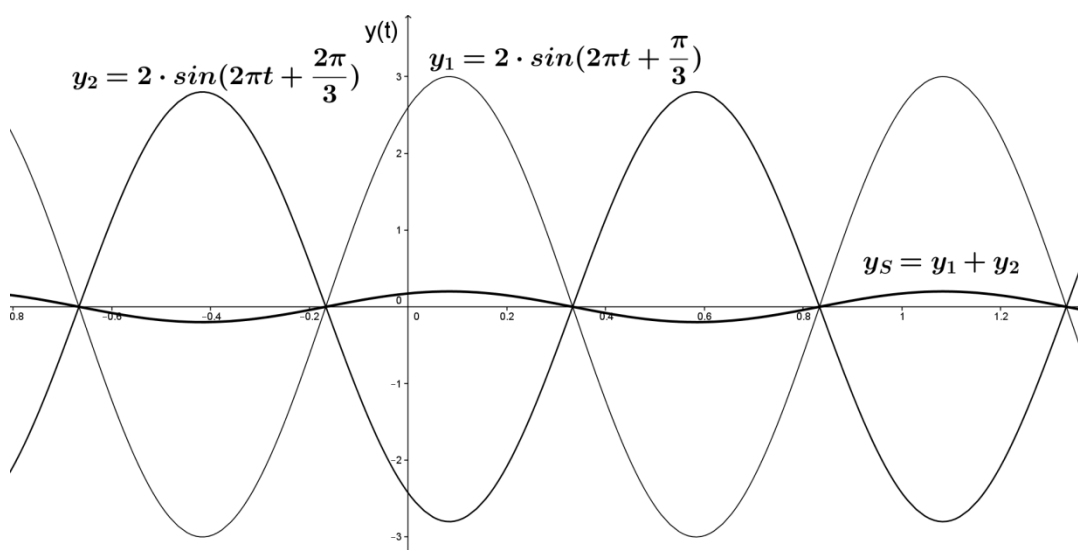


Abbildung 38: Destruktive Interferenz

Beispiel 3: Schwebungen

Unter Schwebungen versteht man die Überlagerung von Schwingungen unterschiedlicher Frequenzen, wobei die Frequenzen nur leicht unterschiedlich sein sollen. Es seien folgende beiden Schwingungen angenommen, zur Vereinfachung seien die Amplituden gleich groß und die Phasenwinkel Null.

$$y_1(t) = A \cdot \sin(2\pi f_1 \cdot t) \text{ und } y_2(t) = A \cdot \sin(2\pi f_2 \cdot t)$$

Für die Überlagerung beider Schwingungen addieren wir diese einfach:

$$y(t) = y_1(t) + y_2(t) = A \cdot \sin(2\pi f_1 \cdot t) + A \cdot \sin(2\pi f_2 \cdot t)$$

Wir wollen die resultierende Schwingung $y(t)$ in einer etwas anderen Form schreiben und verwenden dazu folgendes Additionstheorem:

$$\sin \alpha + \sin \beta = 2 \cdot \sin \frac{\alpha + \beta}{2} \cdot \cos \frac{\alpha - \beta}{2}$$

Damit lässt sich die Summenschwingung folgendermaßen schreiben:

$$y(t) = 2A \cdot \sin \left(2\pi \cdot \frac{f_1 + f_2}{2} \cdot t \right) \cdot \cos \left(2\pi \cdot \frac{f_1 - f_2}{2} \cdot t \right)$$

Die Summenschwingung lässt sich somit als das Produkt einer Überlagerungsschwingung f_R und einer Frequenz f_S der Einhüllenden ausdrücken.

$$f_R = \frac{|f_1 + f_2|}{2} \text{ und } f_S = \frac{|f_1 - f_2|}{2}$$

Die Schwebungsfrequenz ist dabei doppelt so groß wie die Frequenz f_S :

$$f_{\text{Schwebung}} = |f_1 - f_2|$$

Die Dauer einer Schwebung beträgt dann

$$T_{\text{Schwebung}} = \frac{1}{f_{\text{Schwebung}}} = \frac{1}{|f_1 - f_2|}$$

Ein Zahlenbeispiel hierzu: (Siehe Abbildung 39)

Ein reiner Sinuston in der Akustik habe die Frequenz 440 Hz. Ihm sei ein zweiter Sinuston der Frequenz 470 Hz überlagert.

$$y_1(t) = \sin(880\pi \cdot t) \text{ und } y_2(t) = \sin(940\pi \cdot t)$$

Damit ergibt sich

$$f_R = \frac{|440 + 470|}{2} = 455 \text{ Hz und}$$

$$f_S = \frac{|440 - 470|}{2} = 15 \text{ Hz}$$

Die Schwebungsfrequenz beträgt $f_{\text{Schwebung}} = |440 - 470| = 30 \text{ Hz}$

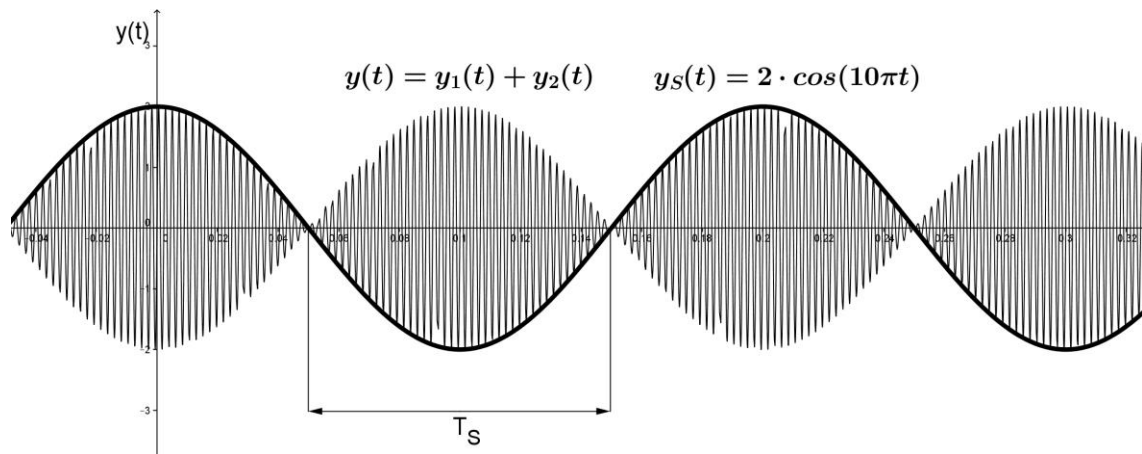


Abbildung 39: Eine Schwebung

3.6.3 Zahlenfolgen und Reihen

Beispiel 1: Bücherüberhänge

Ein Interessantes Beispiel aus dem Bereich der Mechanik ist folgendes, welches im Big Bang 5 von Apolin zu finden ist⁹⁷. Die gestellte Aufgabe lautet folgendermaßen:

Ist es möglich, einen Stapel Bücher so zu bauen, dass das oberste Buch ganz über der Tischkante hängt? Wie würdest du vorgehen? Und wie viele Bücher Überhang kann man mit einem Stapel erzeugen?⁹⁸

Es sollen zuvor ein paar Vorüberlegungen zum Schwerpunkt eines Systems getroffen werden. Der (Körper)Schwerpunkt⁹⁹ ist jener (gedachte) Punkt in einem Körper, an dem der Körper unterstützt werden könnte, sodass er im Gleichgewicht ist. Gleichgewicht bedeutet, dass kein Drehmoment auf ihn wirkt.

Betrachten wir einen masselosen Stab, an dem zwei Massen m_1 und m_2 befestigt sind. Das Gewicht der Massen beträgt dann $m_1 \cdot g$ und $m_2 \cdot g$, das sind auch die beiden einzigen Kräfte, die ein Drehmoment bewirken (vergleiche dazu Abbildung 40).

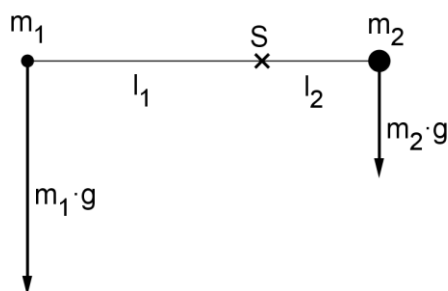


Abbildung 40: Drehmoment zweier Massen 1

⁹⁷ Siehe [18] Seite 41

⁹⁸ Siehe [18] Seite 40 Aufgabe F4

⁹⁹ Es soll angenommen werden, dass die Schwerkraft konstant sei, das heißt, dass in jedem Punkt des Körpers dieselbe Schwerkraft wirkt. In diesem Fall sind der (Körper)Schwerpunkt und der Massenmittelpunkt des Körpers ident.

Es sollen als erstes die beiden Strecken l_1 und l_2 bestimmt werden, die die Lage des Schwerpunktes zwischen den beiden Massen festlegen. Wir betrachten die Drehmomente der beiden Kräfte und erhalten

$$m_1 \cdot g \cdot l_1 = m_2 \cdot g \cdot l_2 \ .$$

Daraus folgt unmittelbar

$$l_1 = \frac{m_2}{m_1} \cdot l_2$$

Anders angeschrieben erkennt man, dass sich die Strecken l_1 und l_2 zueinander so verhalten wie die Massen m_2 und m_1 zueinander. Das ist auch verständlich, da eine größere Masse eine größere Kraft und damit ein größeres Drehmoment bewirkt und daher der Schwerpunkt näher bei dieser Masse liegen muss, damit die Drehmomente auf beiden Seiten gleich groß sind. Es gilt also immer

$$\frac{l_1}{l_2} = \frac{m_2}{m_1}$$

Betrachten wir noch den Fall, dass die beiden Massen m_1 und m_2 immer noch denselben Horizontalabstand haben, aber in der y-Richtung verschoben sind (Abbildung 41).

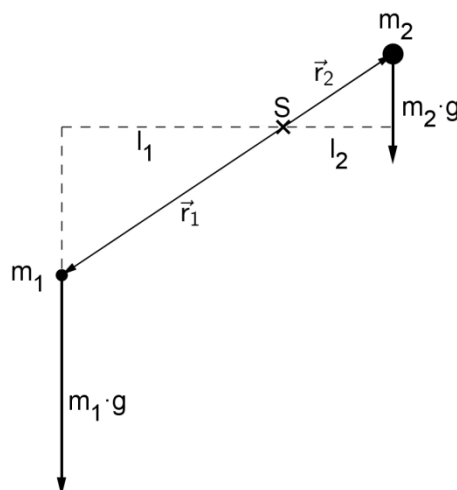


Abbildung 41: Drehmoment zweier Massen 2

Das Drehmoment muss jetzt vektoriell angeschrieben werden:

$$\vec{M}_1 = \vec{r}_1 \times (-m_1 g \cdot \hat{e}_z) \quad \text{und} \quad \vec{M}_2 = \vec{r}_2 \times (-m_2 g \cdot \hat{e}_z)$$

Da $\vec{r}_1$, $\vec{r}_2$ und $\hat{e}_z$ in derselben Ebene liegen, zeigt das Vektorprodukt dieser Vektoren, also $\vec{M}_1$ und $\vec{M}_2$ in dieselbe Richtung, sind aber entgegengesetzt orientiert, da $\vec{r}_1$ und $\vec{r}_2$ in entgegengesetzte Richtungen zeigen. Damit der Gesamtdrehimpuls auf den Schwerpunkt bezogen Null ergibt, müssen noch die Beträge übereinstimmen. Der Betrag eines Vektorprodukts ist gleich der Fläche des Parallelogramms, das von den beiden Vektoren aufgespannt wird. Also kann man schreiben

$$m_1 \cdot g \cdot l_1 = m_2 \cdot g \cdot l_2$$

Und daraus erhält man wiederum wie zuvor für die Längen l_1 und l_2 die Beziehung

$$\frac{l_1}{l_2} = \frac{m_2}{m_1}.$$

Mit diesen Vorüberlegungen kann man nun das Problem des Bücherüberhangs angehen.

Für ein Buch gilt natürlich, dass der Schwerpunkt genau in der Mitte des Buches liegt. Sei die Länge eines Buches gleich L , dann hat der Schwerpunkt den Abstand $\frac{L}{2}$ vom Rand des Buches und genau so weit kann das Buch auch über die Tischkante geschoben werden (Abbildung 42).

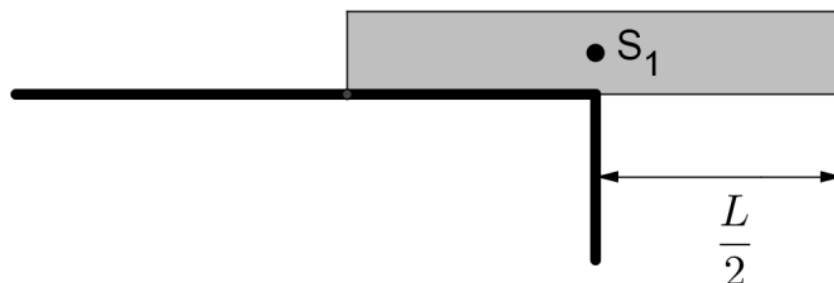


Abbildung 42: Überhang eines Buches

Für zwei Bücher stellt man sich vor, dass das neu hinzugekommene Buch unter den bereits vorhandenen Bücherstapel gelegt wird (welcher in diesem Fall nur aus dem Buch 1 besteht). Das Buch 1 wird so auf das Buch 2 gelegt, dass der Schwerpunkt von Buch 1 auf der Kante von Buch 2 zu liegen kommt, denn damit erreicht man den maximal möglichen Überhang (Abbildung 43).

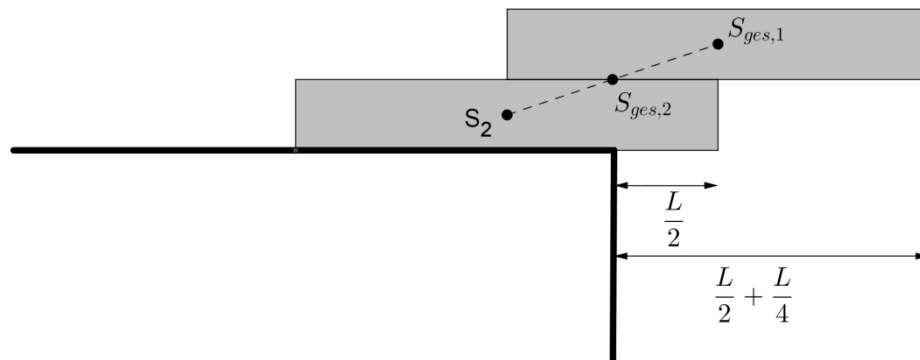


Abbildung 43: Überhang zweier Bücher

Da die Massen beider Bücher gleich sind, liegt der Gesamtschwerpunkt $S_{ges,2}$ in der Mitte der Schwerpunkte S_1 und $S_{ges,1}$, also die Länge $\frac{L}{4}$ von S_1 entfernt. Damit ist der gesamte Überhang $\frac{L}{2}$ (Überhang des Buches 1) + $\frac{L}{4}$ (neuer Überhang), also $\frac{L}{2} + \frac{L}{4}$.

Erweitern wir den Stapel wieder um ein Buch, das wir so unter den Stapel legen, dass der Schwerpunkt $S_{ges,2}$ des vorigen Stapels genau über der Kante des neuen Buches 3 zu liegen kommt und berechnen wir den neuen Gesamtschwerpunkt $S_{ges,3}$ (Abbildung 44).

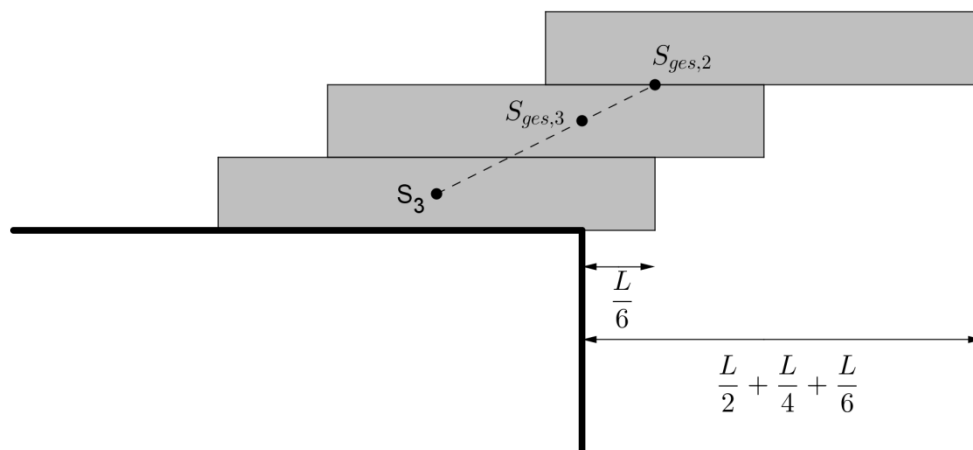


Abbildung 44: Überhang dreier Bücher

Der Gesamtschwerpunkt $S_{\text{ges},3}$ teilt den Abstand der beiden Schwerpunkte $S_{\text{ges},2}$ (Schwerpunkt des bisherigen Stapels aus 2 Büchern) und dem Schwerpunkt S_3 des neu hinzugekommenen Buches 3. Da aber der bisherige Stapel die Masse von 2 Büchern hat und damit die doppelte Masse, liegt der neue Schwerpunkt näher beim Schwerpunkt $S_{\text{ges},2}$. Der Abstand $S_{\text{ges},3}$ zu $S_{\text{ges},2}$ beträgt ein Drittel des Abstandes von S_3 zu $S_{\text{ges},2}$, also

$$\frac{1}{3} \cdot \frac{L}{2} = \frac{L}{6}$$

Damit ergibt sich der gesamte Überhang als $\frac{L}{2} + \frac{L}{4}$ (Überhang des bisherigen Stapels) + $\frac{L}{6}$ (neuer Überhang), also $\frac{L}{2} + \frac{L}{4} + \frac{L}{6}$.

Diese Überlegungen kann man nun bei jedem neu hinzugekommenen Buch wieder anwenden. Jedes neue Buch wird so unter den Stapel gelegt, dass seine Kante genau unter dem Gesamtschwerpunkt des bisherigen Stapels zu liegen kommt. Der Abstand des Gesamtschwerpunkts des bisherigen Stapels beträgt immer $\frac{L}{2}$. Der Überhang kann dann immer um $\frac{L}{2 \cdot \text{Anzahl der Bücher}}$ vergrößert werden. Das kann man mathematisch auch genauer fassen und beweisen durch vollständige Induktion.¹⁰⁰

Behauptung: Bei n gestapelten Büchern ist der maximale Gesamtüberhang

$$D_n = \frac{L}{2} + \frac{L}{4} + \dots + \frac{L}{2 \cdot n}.$$

Für $n=1,2,3$ Bücher ist die Behauptung durch die obigen Überlegungen gezeigt worden. Es komme jetzt ein neues Buch $n+1$ hinzu. Dann ist die Masse des bisherigen Stapels n mal die Masse des neu hinzugekommenen Buches. Also liegt der Schwerpunkt des Stapels aus $n+1$ Büchern näher beim Gesamtschwerpunkt des bisherigen Stapels. Der Abstand des neuen Gesamtschwerpunktes $S_{\text{ges},n+1}$ zum Gesamtschwerpunkt $S_{\text{ges},n}$ des bisherigen Stapels beträgt $\frac{1}{n+1} \cdot \frac{L}{2} = \frac{L}{2 \cdot (n+1)}$. Um diesen Abstand vergrößert sich auch der Überhang gegenüber dem alten Überhang. Dann ist der neue Überhang

¹⁰⁰ Im Physikunterricht wird man sich sicher mit den bisherigen Überlegungen zufriedengeben, in der Mathematik sollte dies aber exakt bewiesen werden, was auch hier geschehen soll.

$$D_{n+1} = \frac{L}{2} + \frac{L}{4} + \dots + \frac{L}{2 \cdot n} + \frac{L}{2 \cdot (n+1)}$$

Der Induktionsbeweis selbst ist in diesem Falle trivial.

$$D_{n+1} = \underbrace{\frac{L}{2} + \frac{L}{4} + \dots + \frac{L}{2 \cdot n}}_{D_n} + \frac{L}{2 \cdot (n+1)}$$

Durch Einsetzen für D_n erhält man natürlich offensichtlich die Formel für D_{n+1} und unter Verwendung des Induktionsanfanges für D_1 ist die Formel für alle n gültig. Also kann man jetzt allgemein für den maximalen Überhang, der bei n Büchern erreicht werden kann, schreiben

$$D_n = \frac{L}{2} \cdot \left(1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n}\right)$$

Um einen Überhang zu erreichen, sodass mehr als eine Buchlänge über der Kante liegt, muss n so bestimmt werden, dass gilt $\frac{L}{2} \cdot \left(1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n}\right) \geq L$ beziehungsweise $1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n} \geq 2$ erfüllt ist. Das ist für $n \geq 4$ der Fall, da $1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} = \frac{25}{12} > 2$ ist.

Um einen Überhang zu erreichen, der mehr als zwei Buchlängen über der Kante liegt, sind bereits 31 Bücher notwendig, da $1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{31} \approx 4,027 > 4$.

Es soll noch die Frage geklärt werden, ob es möglich ist, jeden Überhang mit einer ausreichend großen Anzahl von Büchern zu erreichen. Diese Frage kann man beantworten, wenn man sich die Reihe $1 + \frac{1}{2} + \frac{1}{3} + \dots + \frac{1}{n}$ ansieht. Die Reihe wird in der Mathematik als die harmonische Reihe bezeichnet und man kann zeigen, dass diese Reihe divergiert. Das bedeutet, dass zu jeder vorgegebenen Zahl immer ein n angegeben werden kann, sodass der Wert der Summe der Reihe diese Zahl überschreitet. Also kann jeder beliebige Überhang erzielt werden, zumindest theoretisch, da die Anzahl der nötigen Bücher sehr schnell anwächst.¹⁰¹

¹⁰¹ Um einen Überhang von 10 Büchern zu erreichen, wären bereits mindestens etwa $1,5 \cdot 10^{43}$ Bücher notwendig!

3.6.4 Vektoren im Raum

Es gibt in der Physik neben den skalaren Größen die vektoriellen, die einen Richtungsbezug aufweisen. Sind nun gewisse Größen in der Physik als vektorielle Größen definiert, so können aus diesen wiederum weitere vektorielle Größen definiert werden, genauso wie natürlich aus vektoriellen Größen skalare Größen definiert werden können, was aber in diesem Abschnitt nicht von Interesse ist.

Nehmen wir an, es seien im Raum zwei beliebige Vektoren vorgegeben. Es können nun mit den drei Möglichkeiten der Multiplikation von Vektoren wieder verschiedene weitere Größen abgeleitet werden, die sich insbesondere durch ihren Richtungsbezug zu den Ursprungsvektoren auszeichnen. Folgende Möglichkeiten kommen in Betracht:

1. Multiplikation mit einem Skalar

Wird ein Vektor mit einem Skalar multipliziert, so ergibt sich ein Vektor, der eine andere Länge und eine andere Orientierung haben kann als der Ursprungsvektor.

Ein Beispiel hierzu sind etwa Impuls und Geschwindigkeit. Die Geschwindigkeit $\vec{v}$ ist ein Vektor, multipliziert mit der Masse m , die ein Skalar ist, ergibt sich wieder eine vektorielle Größe, der Impuls $\vec{p}$, also $\vec{p} = m \cdot \vec{v}$. Impuls und Geschwindigkeit sind immer parallel und gleich orientiert, da die Masse nur positiv sein kann.

2. Skalare Multiplikation mit einem Vektor

Werden zwei Vektoren skalar miteinander multipliziert, so ergibt sich immer eine skalare Größe. Ein Beispiel hierzu wäre ein Vektor $\vec{s}$, der den zurückgelegten Weg eines Teilchens angibt und ein Kraftvektor $\vec{F}$, der die Kraft angibt, die verrichtet wird. Skalar multipliziert ergibt sich eine neue physikalische Größe, die Arbeit W , die am oder vom Teilchen verrichtet wird, also $W = \vec{F} \cdot \vec{s}$. Hier zeigt sich auch die Vereinfachung, die die Vektorschreibweise mit sich bringt, denn es ist nur die Kraftkomponente entscheidend, die in Richtung des zurückgelegten Weges wirkt. Ohne Vektoren müsste man die Arbeit dann mit Hilfe trigonometrischer Funktionen definieren, was natürlich auch nicht so elegant wie mit Vektoren formuliert werden kann.

3. Vektorielle Multiplikation zweier Vektoren

Bei der vektoriellen Multiplikation zweier Vektoren $\vec{a}$, $\vec{b}$ erhält man als Ergebnis wieder einen Vektor $\vec{c} = \vec{a} \times \vec{b}$. Dieser Vektor hat drei besondere Eigenschaften, die ihn charakterisieren:

- i) $\vec{c}$ steht normal auf jeden der beiden Vektoren $\vec{a}$ und $\vec{b}$, was bedeutet, dass das skalare Produkt jeweils Null ergibt: $\vec{a} \cdot \vec{c} = 0$ und $\vec{b} \cdot \vec{c} = 0$.
- ii) Die Länge beziehungsweise der Betrag von $\vec{c}$ entspricht dem Flächeninhalt des von den Vektoren $\vec{a}$ und $\vec{b}$ aufgespannten Parallelogramms
 $|\vec{c}| = |\vec{a}| \cdot |\vec{b}| \cdot \sin(\angle(\vec{a}, \vec{b}))$
- iii) Die Orientierung des Vektors $\vec{c}$ ist so, dass $\vec{c}$ mit den Vektoren $\vec{a}$ und $\vec{b}$ ein Rechtssystem bildet.

4. Es soll hier nur erwähnt werden, dass es mittels einer Vektormultiplikation nicht möglich ist, zwei Vektoren zu verknüpfen, die nicht parallel sind und nicht normal aufeinander stehen. Um dies zu bewerkstelligen, ist eine Matrixmultiplikation nötig, welche aber in einer Allgemeinbildenden Höheren Schule nicht zum Lehrplan gehört. Formal angeschrieben lautet eine solche Multiplikation wie folgt:

$$\begin{pmatrix} a_1 \\ a_2 \\ a_3 \end{pmatrix} = \begin{pmatrix} c_{11} & c_{12} & c_{13} \\ c_{21} & c_{22} & c_{23} \\ c_{31} & c_{32} & c_{33} \end{pmatrix} \cdot \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix}$$

Zur Anwendung kommt diese Matrix, eigentlich ein sogenannter Tensor 2. Stufe, beim Zusammenhang zwischen der Winkelgeschwindigkeit und dem Drehimpuls, wobei der Winkelgeschwindigkeitsvektor $\vec{\omega}$ und der Drehimpulsvektor $\vec{L}$ nicht parallel sind. Die Matrix, die die Verknüpfung von $\vec{\omega}$ und $\vec{L}$ vermittelt, ist die sogenannte Trägheitsmomentenmatrix I , es gilt also $\vec{L} = I \cdot \vec{\omega}$.

Im Folgenden sollen einige Größen vorgestellt werden, die räumliche Vektoren sind und in der Physik grundlegende Bedeutung haben. Einige der Herleitungen sind für

die Schulphysik sicherlich nur für besonders motivierte Klassen sinnvoll, begabte SchülerInnen können die Herleitungen aber sicherlich nachvollziehen.

Beispiel 1: Vektorprodukte in der Mechanik

1. Die Winkelgeschwindigkeit $\vec{\omega}$

Die grundlegende vektorielle Größe bei Drehbewegungen ist die Winkelgeschwindigkeit, aus der dann durch Vektoroperationen, insbesondere mit Hilfe des Vektorproduktes, weitere wichtige Vektorgrößen folgen.

Betrachten wir eine Scheibe, die senkrecht um eine Achse durch ihren Mittelpunkt rotiert (Abbildung 45).

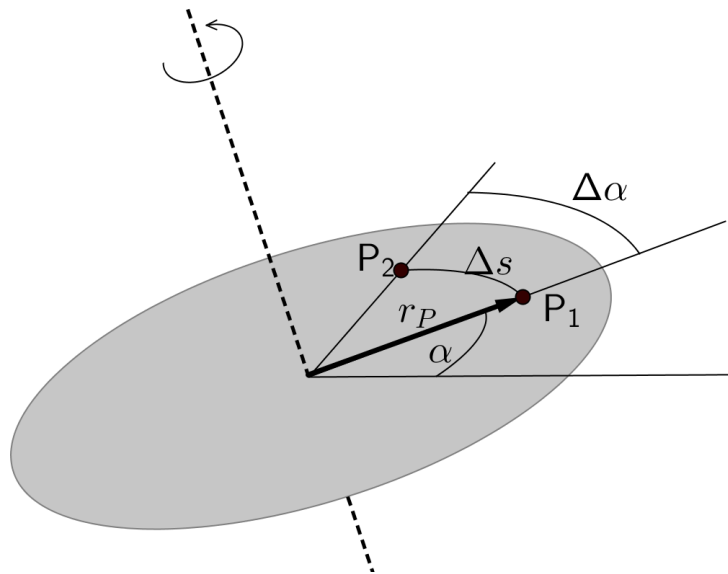


Abbildung 45: Zur Herleitung der Winkelgeschwindigkeit

Ein Punkt bewege sich im Zeitintervall Δt auf einem Kreisbogen um die Bogenstrecke Δs von P_1 nach P_2 . Die Strecke Δs ergibt sich dann als $\Delta s = r_p \cdot \Delta \alpha$, wobei α in Bogengrad angegeben sein muss. Der Quotient $\frac{\Delta s}{r_p}$ entspricht dem

Drehwinkel $\Delta \alpha$ und ist für alle Punkte der Scheibe gleich. Deshalb ist die zeitliche Änderung des Drehwinkels ein geeignetes Maß für die Geschwindigkeit der Punkte der Scheibe, da diese Geschwindigkeit für jeden Punkt der Scheibe gleichermaßen

gilt. Lässt man den Grenzübergang $\Delta t \rightarrow 0$ erfolgen, gelangt man zur Winkelgeschwindigkeit ω .

$$\omega = \lim_{\Delta t \rightarrow 0} \frac{\alpha(t + \Delta t) - \alpha(t)}{\Delta t} = \frac{d\alpha}{dt} = \dot{\alpha}$$

Man fasst nun die Winkelgeschwindigkeit zweckmäßigerweise als Vektor auf, um die räumliche Lage der Drehachse mit einzubeziehen. Der Betrag der Winkelgeschwindigkeit ist die zeitliche Ableitung des Drehwinkels α und die Richtung entspricht der Richtung der Drehachse, da diese Richtung für die Drehung charakteristisch ist. Die Orientierung von ω wird so festgelegt, dass ω positiv ist, wenn die Winkeländerung $\Delta\alpha$ positiv ist, das heißt, wenn die Drehung gegen den Uhrzeigersinn erfolgt (Abbildung 46). Daraus ergibt sich die praktische Rechte-Hand-Regel, nach der der Daumen der rechten Hand in Richtung der Winkelgeschwindigkeit zeigt, wenn die Drehachse mit den restlichen Fingern so umschlossen wird, dass die Finger in Richtung der Drehung zeigen.

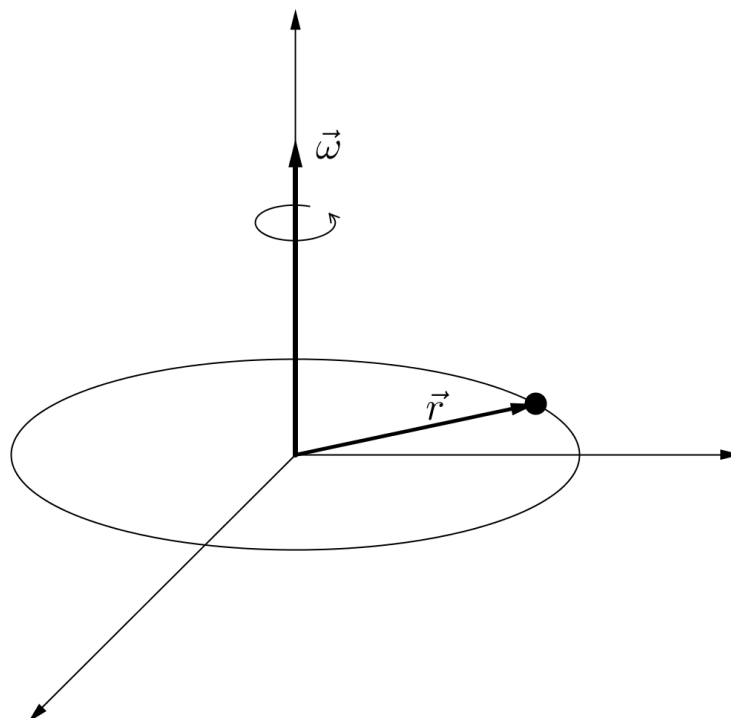


Abbildung 46: Drehgeschwindigkeitsvektor

2. Der Bahngeschwindigkeitsvektor $\vec{v}$

Der Massenpunkt legt in der Zeit Δt den Winkel $\Delta\alpha$ zurück. Für kleine Zeitintervalle Δt sind Δr und Δs fast gleich groß (siehe Abbildung 47), also gilt

$$\frac{\Delta r}{\Delta t} \approx \frac{\Delta s}{\Delta t} = \frac{r \cdot \Delta\alpha}{\Delta t}.$$

Die Bahngeschwindigkeit ist

$$\mathbf{v} = \lim_{\Delta t \rightarrow 0} \frac{\Delta \mathbf{r}}{\Delta t}.$$

Im Grenzfall $\Delta t \rightarrow 0$ können Δr und Δs gleichgesetzt werden, sodass für die Bahngeschwindigkeit gilt

$$\mathbf{v} = \lim_{\Delta t \rightarrow 0} \frac{\Delta \mathbf{r}}{\Delta t} = \lim_{\Delta t \rightarrow 0} \frac{\Delta s}{\Delta t} = \lim_{\Delta t \rightarrow 0} \frac{r \cdot \Delta\alpha}{\Delta t} = r \cdot \lim_{\Delta t \rightarrow 0} \frac{\Delta\alpha}{\Delta t} = r \cdot \omega.$$

Da $\vec{v}$ auf $\vec{r}$ und $\vec{\omega}$ normal steht, lässt sich die Bahngeschwindigkeit als Vektorprodukt schreiben:

$$\vec{v} = \vec{\omega} \times \vec{r}$$

Daraus ergibt sich sofort der Impulsvektor

$$\vec{p} = m \cdot \vec{\omega} \times \vec{r}$$

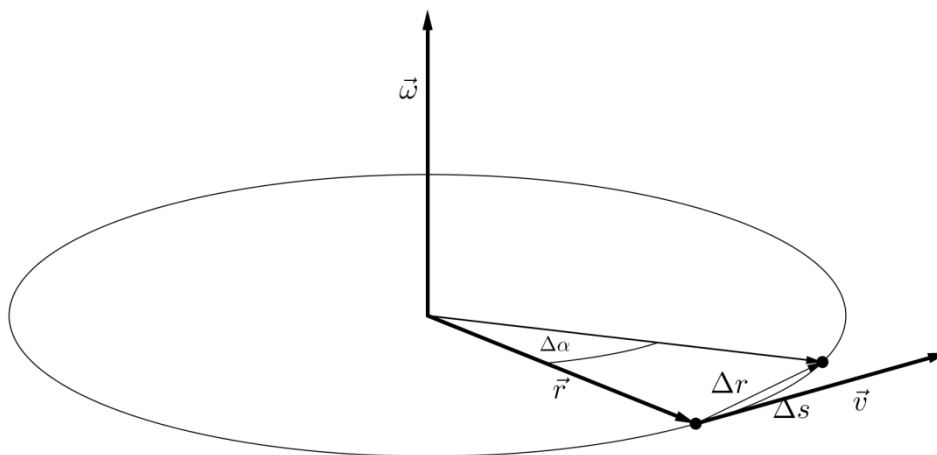


Abbildung 47: Der Bahndrehimpulsvektor $\vec{v}$

3. Der Drehimpulsvektor $\vec{L}$

Aus dem Radiusvektor $\vec{r}$ und dem Impulsvektor $\vec{p}$ lässt sich ein besonders wichtiger neuer Vektor bilden, der Drehimpulsvektor $\vec{L}$, der lautet

$$\vec{L} = m \cdot \vec{r} \times \vec{v}$$

Man sieht, dass der Drehimpulsvektor normal auf den Radiusvektor $\vec{r}$ und den Bahngeschwindigkeitsvektor $\vec{v}$ beziehungsweise den Impulsvektor $\vec{p} = m \cdot \vec{v}$ steht.

Setzt man für den Bahngeschwindigkeitsvektor $\vec{v} = \vec{\omega} \times \vec{r}$, so lautet der Drehimpulsvektor

$$\vec{L} = m \cdot \vec{r} \times (\vec{\omega} \times \vec{r})$$

4. Der Drehmomentvektor $\vec{M}$

Das Drehmoment, das durch eine äußere Kraft $\vec{F}$ hervorgerufen wird, ist ebenfalls eine physikalische Größe, die als Vektorprodukt ausgedrückt werden kann. Es sei, wie in Abbildung 48 ersichtlich, ein rotierender Massenpunkt im Abstand r von der Drehachse betrachtet, auf den eine äußere Kraft $\vec{F}$ einwirken soll.

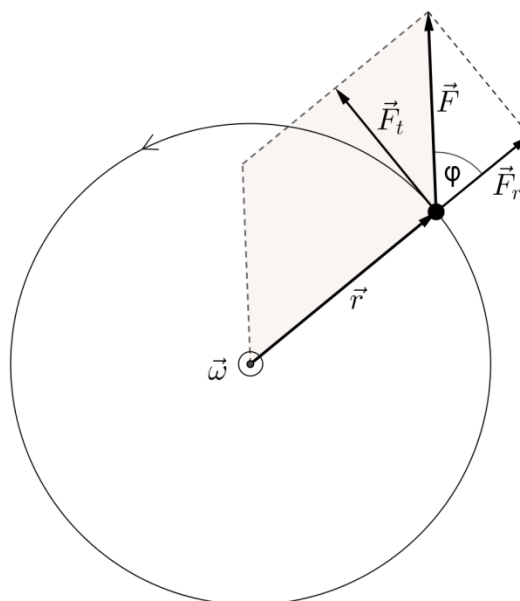


Abbildung 48: Der Drehmomentvektor $\vec{M}$

Die Ebene, in der der Massenpunkt rotiert, soll in der Seitenebene liegen und er rotiere entgegen dem Uhrzeigersinn, sodass der Vektor der Winkelgeschwindigkeit $\vec{\omega}$ aus der Seitenebene herausragt. Die Kraft $\vec{F}$ sei so gerichtet, dass sie den Winkel φ zum Radiusvektor $\vec{r}$ einnimmt.

Die äußere Kraft $\vec{F}$ wird in einen Vektor $\vec{F}_t$ tangential zur Bahnkurve und einen Vektor $\vec{F}_r$ in radialer Richtung zerlegt, wobei nur der tangential Anteil von $\vec{F}$ eine Änderung der Winkelgeschwindigkeit bewirkt.

Da $\vec{F}_t$ normal auf den Radiusvektor $\vec{r}$ steht, ist der Betrag des Drehmoments das Produkt aus den Beträgen der Tangentialkomponente des äußeren Kraftvektors und des Radiusvektors, somit $M = |\vec{r}| \cdot |\vec{F}_t|$. Aus der Definition des Sinus folgt

$$\sin \varphi = \frac{|\vec{F}_t|}{|\vec{F}|}.$$

Somit kann für den Betrag des Drehmoments auch geschrieben werden

$$M = |\vec{r}| \cdot |\vec{F}| \cdot \sin \varphi,$$

was wiederum dem Betrag des Vektorproduktes beider Vektoren entspricht. Da die Bewegungsänderung in der Ebene stattfindet, in der die Vektoren $\vec{r}$ und $\vec{F}$ liegen, steht der Vektor $\vec{M}$ normal sowohl auf den Radiusvektor $\vec{r}$ als auch auf den Kraftvektor $\vec{F}$. Deshalb lässt sich das Drehmoment sehr elegant als Vektorprodukt schreiben.

$$\vec{M} = \vec{r} \times \vec{F}$$

Beispiel 2: Der Euler'sche Drehimpulssatz und die Drehimpulserhaltung

Interessanterweise können die Bewegungsgesetze der Translation und der Rotation in völliger Analogie formuliert werden. Dadurch erhält man einander zugeordnete Größen und Gleichungen. Diese einander zugeordneten Größen und Beziehungen seien in der Tabelle 6 aufgelistet.

Translation		Rotation	
Weg	$\vec{r}$	Drehwinkel	$\vec{\phi}$
Geschwindigkeit	$\vec{v}$	Winkelgeschwindigkeit	$\vec{\omega}$
Beschleunigung ¹⁰²	$\vec{a} = \dot{\vec{v}} = \ddot{\vec{r}}$	Winkelbeschleunigung	$\vec{\alpha} = \dot{\vec{\omega}} = \ddot{\vec{\phi}}$
Kraft	$\vec{F}$	Drehmoment	$\vec{M} = \vec{r} \times \vec{F}$
Arbeit	$W = \int \vec{F} \cdot d\vec{s}$	Dreharbeit	$W = \int \vec{M} \cdot \vec{\omega} dt$
Masse	m	Trägheitsmoment	I ¹⁰³
Impuls	$\vec{p} = m \cdot \vec{v}$	Drehimpuls	$\vec{L} = I \cdot \vec{\omega}$
Impulssatz	$\dot{\vec{p}} = \vec{F}$	Drehimpulssatz	$\dot{\vec{L}} = \vec{M}$
Impulserhaltungssatz	$\dot{\vec{p}} = \vec{0}$	Drehimpulserhaltungssatz	$\dot{\vec{L}} = \vec{0}$
Kinetische Energie	$E_{\text{kin}} = \frac{1}{2} \cdot m \vec{v}^2$	Rotationsenergie	$E_{\text{rot}} = \frac{1}{2} \cdot I \vec{\omega}^2$
Bewegungsgleichung			
$\vec{F} = m \cdot \vec{a}$		$\vec{M} = I \cdot \vec{\alpha}$	

Tabelle 6: Gegenüberstellung Translation - Rotation

So gilt etwa für translatorische Bewegungen das zweite Newton'sche Gesetz $\vec{F} = m \cdot \vec{a}$. Das kann auch in folgender Form geschrieben werden, wenn man berücksichtigt, dass die Beschleunigung die zeitliche Ableitung der Geschwindigkeit ist

$$\vec{F} = m \cdot \frac{d\vec{v}}{dt} = \frac{d}{dt}(m \cdot \vec{v}) = \frac{d\vec{p}}{dt}.$$

Das ergibt bereits die wichtige Erkenntnis, dass eine Kraft immer Änderung des Impulses bewirkt.¹⁰⁴

¹⁰² Ein Punkt über einer Größe bedeutet die zeitliche Ableitung der Größe, zwei Punkte die zweite Ableitung nach der Zeit.

¹⁰³ Das Trägheitsmoment beträgt für eine Punktmasse m, die im Abstand r rotiert, $I = m \cdot r^2$, für einen Körper, der aus n Massenpunkten besteht, berechnet sich das Trägheitsmoment zu $I = \sum_{i=0}^n m_i \cdot r_i^2$. Ist der Körper durch eine Massenverteilung gegeben, so ergibt sich $I = \int_V r^2 dm = \int_V r^2 \rho(\vec{r}) dV$. Stehen der Winkelgeschwindigkeitsvektor und der Drehimpulsvektor nicht parallel zueinander, so muss das skalare Trägheitsmoment verallgemeinert werden zum Trägheitstensor.

1754 hat Leonhard Euler¹⁰⁵ den nach ihm benannten Euler'schen Drehimpulssatz formuliert, nach dem ein Drehmoment eine zeitliche Änderung des Drehimpulses bewirkt, das also gilt

$$\vec{M} = \frac{d\vec{L}}{dt}$$

Die mathematische Herleitung dieser wichtigen Gleichung wird im Schulunterricht üblicherweise nicht gebracht, da Ableitungen von Vektoroperationen, wie dem Vektorprodukt, nicht zum Lehrstoff gehören. Unter der Annahme, dass für das Vektorprodukt für die Ableitung genauso die Produktregel angewandt werden muss, kann der Drehimpulssatz sehr einfach auch im Physikunterricht gezeigt werden.

Der Drehimpuls ist definiert als $\vec{L} = \vec{r} \times m \cdot \vec{v}$ und somit ergibt sich für die erste Ableitung

$$\frac{d}{dt} \vec{L} = \frac{d\vec{r}}{dt} \times m \cdot \vec{v} + \vec{r} \times m \cdot \frac{d\vec{v}}{dt}$$

Mit $\frac{d\vec{r}}{dt} = \vec{v}$ und $m \cdot \frac{d\vec{v}}{dt} = \frac{d\vec{p}}{dt} = \vec{F}$ folgt

$$\frac{d}{dt} \vec{L} = \vec{v} \times m \cdot \vec{v} + \vec{r} \times \vec{F}$$

Unter Berücksichtigung, dass das Vektorprodukt eines Vektors mit sich selbst immer der Nullvektor ist, also $\vec{v} \times \vec{v} = \vec{0}$, ergibt sich somit direkt der Drehimpulssatz zu

$$\frac{d}{dt} \vec{L} = \vec{r} \times \vec{F} = \vec{M}.$$

Wenn auf einen rotierenden Körper beziehungsweise ein rotierendes System kein äußeres Drehmoment wirkt, so ergibt sich, dass die zeitliche Ableitung des Drehimpulses der Nullvektor ist

¹⁰⁴ Im allgemeinen Fall muss auch berücksichtigt werden, dass die Masse nicht konstant bleibt, sondern sich ändern kann, etwa bei einer Rakete, die kontinuierlich ihren Treibstoff verbrennt, sodass ihre Masse immer mehr abnimmt. In diesem Fall ergibt sich für die Bewegungsgleichung:

$$\vec{F} = \frac{d}{dt} (m \cdot \vec{v}) = \dot{m} \cdot \vec{v} + m \cdot \dot{\vec{v}} = \dot{m} \cdot \vec{v} + m \cdot \vec{a}$$

¹⁰⁵ Leonhard Euler, schweizer Mathematiker, 1707-1783; arbeitete neben der Mathematik auch auf den Gebieten der Strömungsmechanik, der Kreiseltheorie und der Optik.

$$\frac{d}{dt} \vec{L} = \vec{0}$$

Daraus folgt unmittelbar der Drehimpulserhaltungssatz, der besagt, dass der Drehimpuls eines Systems erhalten bleibt, solange kein äußeres Drehmoment einwirkt.

$\vec{L} = \text{konstant, falls } \vec{M}_{\text{extern}} = \vec{0}$

Dieser Satz gehört zu den wichtigsten der Physik, der Anwendungen von der klassischen Mechanik bis zur Quantenmechanik hat, der besonders natürlich in der Astronomie eine immense Bedeutung hat. Es sei nur daran erinnert, dass die Erhaltung der Richtung des Gesamtdrehimpulses des Sonnensystems zur Folge hat, dass sich die Planeten in einer Ebene um die Sonne bewegen beziehungsweise die Erhaltung des Betrages des Drehimpulses zum zweiten Kepler'schen Gesetz führt.

3.6.5 Statistik

Im Grunde spielt die Statistik in der Physik immer eine Rolle, da man es in Experimenten immer mit, mathematisch gesprochen, Stichproben, zu tun hat. Die Stichproben entsprechen den Ausgängen von Experimenten, von denen dann etwa Mittelwerte berechnet werden. Des Weiteren müssen dann Standardabweichungen bestimmt werden, um die Qualität der Ergebnisse abzuschätzen. Im Falle, dass Regressionsgeraden bestimmt werden müssen, sollte auch der Korrelationskoeffizient berechnet werden, um die Güte der Ausgleichsgeraden abzuschätzen. Das alles sind Standardverfahren, die jeder Physiker in der Praxis immer wieder anwenden muss und die er heute natürlich einfach mit dem Computer lösen kann.

Die sicherlich beeindruckendste Anwendung der Statistik in der Physik ist sicher auf dem Gebiet der statistischen Mechanik beziehungsweise kinetischen Gastheorie gegeben. Das Interessante dabei ist, dass hier Aussagen über Naturphänomene getroffen werden können, die alleine aufgrund statistischer mathematischer Verfahren gewonnen worden sind. Die Gesetze der statistischen Mechanik beruhen allein auf der Wechselwirkung und der Bewegung einer extrem hohen Zahl von Teilchen, worauf auch ihre Allgemeingültigkeit und Anwendbarkeit auf die unterschiedlichsten Teilgebiete der Physik beruht.

Beispiel: Statistische Mechanik – Die Bernoulligleichung und die Gasgesetze

Drei wichtige Gasgesetze¹⁰⁶ wurden durch Experimente bereits im 17. und 18. Jahrhundert entdeckt. Es war aber mit dem damaligen Wissenstand nicht möglich, diese phänomenologischen Gasgesetze zu begründen, dies wurde erst ermöglicht mit der Entdeckung und Akzeptanz von Atomen und Molekülen, aus denen sich diese Gase zusammensetzen und der Anwendung statistischer Methoden.

¹⁰⁶ Siehe dazu [25] Seite 673 ff.

Das erste Gesetz ist das **Boyle-Mariotte'sche Gesetz**¹⁰⁷, das besagt, dass das Produkt aus Druck und Volumen eines Gases bei konstanter Temperatur konstant ist.

Das zweite Gesetz ist das **Gay-Lussac'sche Gesetz**¹⁰⁸, das besagt, dass der Quotient aus Volumen und absoluter Temperatur konstant ist, wenn der Druck konstant gehalten wird.

Das dritte Gesetz ist das **Amontons'sche Gesetz**¹⁰⁹, das besagt, dass der Quotient aus Druck und absoluter Temperatur konstant ist, wenn das Volumen konstant gehalten wird.

Diese drei Gasgesetze können zusammengefasst werden im **allgemeinen Gasgesetz**¹¹⁰, das einen Zusammenhang zwischen den drei makroskopischen Zustandsgrößen Druck, absoluter Temperatur und Gasdruck herstellt. Es soll in diesem Beispiel gezeigt werden, wie mittels einfacher statistischer Annahmen diese Gesetze aus den mikroskopischen Zustandsgrößen hergeleitet werden können. Damit schafft die Statistik den Übergang von einer mikroskopischen Beschreibung zur makroskopischen Beschreibung der Materie. Die Grundzüge der Herleitung wurden dem Schulbuch Faszination Physik 1+2 entnommen und für diese Arbeit entsprechend ausformuliert und erweitert.¹¹¹

Wir gehen bei der folgenden Herleitung von einem idealen Gas aus. Das bedeutet, dass das Eigenvolumen der Gasteilchen vernachlässigt wird und insbesondere angenommen wird, dass die Gasteilchen untereinander und mit der Gefäßwand nur durch elastische Stöße¹¹² wechselwirken.

Herleitung der Bernoulli-Gleichung für den Gasdruck

Wir nehmen an, dass eine große Anzahl von N Gasteilchen in einem quaderförmigen Behälter mit dem Volumen V eingeschlossen ist. Die Gasteilchen führen ungeordnete Bewegungen aus und üben durch die dadurch erfolgenden Stöße auf die Behälterwände

¹⁰⁷ Dieses Gesetz wurde von Robert Boyle (1627-1691) im Jahre 1622 entdeckt und von Edme Mariotte (1620-1684) im Jahre 1676 experimentell nachvollzogen.

¹⁰⁸ Das Gesetz wurde von Jaques Charles (1746-1823) und Joseph Louis Gay-Lussac (1778-1850) gefunden.

¹⁰⁹ Das Amontons'sche Gesetz wird oft auch als 2. Gay-Lussac'sches Gesetz bezeichnet; Guillaume Amontons (1663-1705)

¹¹⁰ Die erste Formulierung stammt von Émile Clapeyron im Jahr 1834

¹¹¹ Siehe [26] Seite 132 ff.

¹¹² Bei elastischen Stößen ist die Summe der kinetischen Energien der Teilchen vor und nach dem Stoßprozess gleich.

einen Gasdruck auf die Wände aus. Stößt ein Gasteilchen auf die Behälterwand, so wird es reflektiert, es ändert dabei seinen Impuls und überträgt die Impulsdifferenz auf die Wand (Abbildung 49).

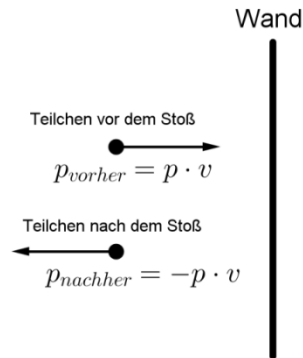


Abbildung 49: Impulsübertrag auf die Behälterwände

Die Impulsänderung¹¹³ und damit der auf die Wand übertragene Impuls beträgt

$$\Delta p = p_{\text{vorher}} - p_{\text{nachher}} = p \cdot v - (-p \cdot v) = 2 \cdot p \cdot v$$

Um den Gesamtimpuls zu erhalten, der in der Zeitdauer Δt auf die Behälterwände übertragen wird, muss man die Gesamtanzahl der Teilchen, die die Wände in der Zeit Δt erreichen, mit der Impulsänderung $2 \cdot p \cdot v$ multiplizieren.

Für die weitere Rechnung sind folgende drei Punkte zu berücksichtigen:

1. Die Gasteilchen haben unterschiedliche Geschwindigkeiten, deshalb muss statt der Geschwindigkeit v die mittlere Geschwindigkeit $\bar{v}$ verwendet werden, so dass wir für den Impulsübertrag den mittleren Impuls $\bar{p}$ verwenden müssen, $\bar{p} = 2 \cdot p \cdot \bar{v}$.
2. Es können nur diejenigen Gasteilchen die Wand erreichen, die während der Zeitdauer Δt höchstens die Strecke $\bar{v} \cdot \Delta t$ von der Behälterwand entfernt sind.
3. Nur $\frac{1}{6}$ der Gasteilchen fliegt in Richtung der Wand.¹¹⁴

¹¹³ Es ist nur nötig, die Impulsänderung der Impulskomponente normal zur Wand zu berücksichtigen, die Impulskomponente parallel zur Wand ist Null, sodass durch sie kein Impuls auf die Wand übertragen wird.

¹¹⁴ Ein Gasteilchen könnte zwar die richtige Geschwindigkeit $\bar{v}$ haben, aber der Geschwindigkeitsvektor könnte in die Richtung einer der anderen 5 Wände gerichtet sein, sodass nur ein Anteil von einem Sechstel der Gasteilchen mit der richtigen Geschwindigkeit tatsächlich die Wand erreicht.

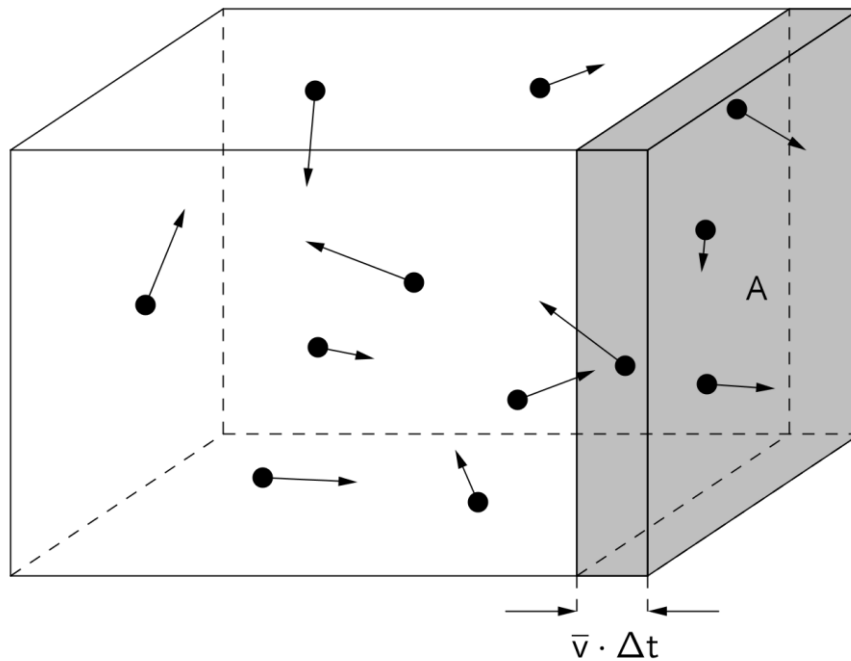


Abbildung 50: Zur Herleitung der Bernoulli-Gleichung

Somit ergibt sich für die Anzahl Z der Gasteilchen, die die Wände erreicht (Abbildung 50)

$$Z \approx \frac{1}{6} \cdot \frac{N}{V} \cdot A \cdot \bar{v} \cdot \Delta t$$

$\frac{N}{V}$... Teilchendichte im Behälter

$A \cdot \bar{v} \cdot \Delta t$... Volumen, in dem sich die Teilchen befinden, die die Behälterwand mit der Fläche A erreichen können

Damit ergibt sich der Impulsübertrag auf die Wand der Fläche A zu

$$\Delta p = \frac{1}{6} \cdot \frac{N}{V} \cdot A \cdot \bar{v} \cdot \Delta t \cdot 2 \cdot m \cdot \bar{v} = \frac{1}{3} \cdot \frac{N}{V} \cdot A \cdot m \cdot \bar{v}^2 \cdot \Delta t$$

Der Kraftstoß auf die Wand¹¹⁵ ergibt sich als Impulsübertrag pro Zeiteinheit

$$F = \frac{\Delta p}{\Delta t} = \frac{1}{3} \cdot \frac{N}{V} \cdot A \cdot m \cdot \bar{v}^2$$

und damit erhält man für den Druck p , der auf die Wand ausgeübt wird

¹¹⁵ Das sieht man folgendermaßen: $F = m \cdot a = m \cdot \frac{\Delta v}{\Delta t} = \frac{m \cdot \Delta v}{\Delta t} = \frac{\Delta p}{\Delta t}$

$$p = \frac{F}{A} = \frac{1}{3} \cdot \frac{N}{V} \cdot m \cdot \overline{v^2} \quad 116$$

Diese Gleichung soll noch etwas umgeformt werden, indem die mittlere kinetische Energie eines Teilchens betrachtet wird.

$$\overline{E_{\text{kin}}} = \frac{1}{2} \cdot m \cdot \overline{v^2} \approx \frac{1}{2} \cdot m \cdot \overline{v}^2 \quad 117$$

Dies in die Gleichung für den Druck eingesetzt ergibt schließlich folgende Gleichung für den Gasdruck, die sogenannte **Bernoulli-Gleichung**:

$$p = \frac{2}{3} \cdot \frac{N}{V} \cdot \overline{E_{\text{kin}}}$$

Die Bernoulli-Gleichung gibt den Zusammenhang zwischen der makroskopischen Größe Druck und der mikroskopischen Größe der mittleren kinetischen Energie an. Um den Zusammenhang des Druckes mit den makroskopischen Größen Temperatur und Volumen herzustellen, müssen noch folgende Überlegungen angestellt werden.

Die innere Energie U eines Systems von N Gasteilchen ist durch die Summe der kinetischen Energien der einzelnen Teilchen gegeben

$$U = \sum_{i=1}^N E_{\text{kin}}^{(i)} = N \cdot \underbrace{\frac{1}{N} \cdot \sum_{i=1}^N E_{\text{kin}}^{(i)}}_{\overline{E_{\text{kin}}}} = N \cdot \overline{E_{\text{kin}}} \quad 118$$

Die absolute Temperatur T eines Systems wird als ein Maß für die innere Energie des Systems definiert, es ist also die innere Energie U proportional zur absoluten Temperatur T . Somit kann folgende Proportionalitätskette angeschrieben werden

$$T \sim U \sim \overline{E_{\text{kin}}}$$

¹¹⁶ Es werden hier dasselbe Größenzeichen p für den Impuls und den Druck verwendet, aus dem Zusammenhang ist aber klar, welche Größe jeweils gemeint ist.

¹¹⁷ Es kann gezeigt werden, dass in sehr guter Näherung $\overline{v^2} = \overline{v}^2$ verwendet werden kann, dass also das mittlere Geschwindigkeitsquadrat etwa dem Quadrat der mittleren Geschwindigkeit der Gasteilchen entspricht.

¹¹⁸ Die innere Energie eines Gases rührt alleine von der Wärmebewegung der Gasteilchen her. Werden Moleküle betrachtet, spielen außer der kinetischen Energie der translatorischen Bewegung auch die Rotationsenergien der Gasteilchen eine Rolle, die hier aber außer Betracht bleiben sollen.

Damit können nun als Anwendung dieser statistischen Beschreibung des idealen Gases die Gasgesetze aus der Bernoulli-Gleichung hergeleitet werden. Es sollen dazu abgeschlossene Systeme betrachtet werden, die Teilchenanzahl N sich also nicht ändert.

Das Boyle-Mariotte'sche Gesetz

Aus der Bernoulli-Gleichung folgt

$$p \cdot V = \frac{2}{3} \cdot N \cdot a \cdot T, \text{ da aus } T \sim \overline{E_{\text{kin}}} \text{ folgt } \overline{E_{\text{kin}}} = a \cdot T \text{ mit einer Konstanten } a.$$

Falls die Temperatur T konstant gehalten wird (*isotherme* Zustandsänderung), ist die rechte Seite der Gleichung konstant und es ergibt sich damit das Boyle-Mariotte'sche Gesetz

$p \cdot V = \text{konstant, falls } T \text{ konstant}$
--

Das Gay-Lussac'sche Gesetz

Aus der Bernoulli-Gleichung folgt

$$p = \frac{2}{3} \cdot \frac{N}{V} \cdot a \cdot T \text{ und daraus}$$

$$\frac{V}{T} = \frac{2}{3} \cdot \frac{N \cdot a}{p}$$

Falls der Druck p konstant gehalten wird (*isobare* Zustandsänderung), ist die rechte Seite der Gleichung konstant und es ergibt sich damit das Gay-Lussac'sche Gesetz

$\frac{V}{T} = \text{konstant, falls } p \text{ konstant}$
--

Das Amontons'sche Gesetz

Aus der Bernoulli-Gleichung folgt

$$p = \frac{2}{3} \cdot \frac{N}{V} \cdot a \cdot T \quad \text{und daraus} \quad \frac{p}{T} = \frac{2}{3} \cdot \frac{N \cdot a}{V}$$

Falls das Volumen V konstant gehalten wird (*isochore* Zustandsänderung), ist die rechte Seite der Gleichung konstant und es ergibt sich damit das Amontons'sche Gesetz

$$\frac{p}{T} = \text{konstant, falls } V \text{ konstant}$$

Das allgemeine Gasgesetz

Aus den obigen drei Gasgesetzen (bei denen jeweils eine makroskopische Zustandsgröße konstant gehalten wird), kann nun das allgemeine Gasgesetz abgeleitet werden, in dem schließlich alle drei Zustandsgrößen Druck, Volumen und Temperatur sowie die Teilchenzahl, gleichzeitig verändert werden können.

Aus dem Amontons'schen Gesetz folgt

$$\frac{p}{T} = \text{konstant} \Rightarrow \frac{p}{T} = \frac{p_0}{T_0} \Rightarrow p = T \cdot \frac{p_0}{T_0}$$

p, T sind dabei ein beliebiger Druck und Temperatur, p_0, T_0 Druck unter Temperatur bei sogenannten Normalbedingungen¹¹⁹.

Aus der Bernoulli-Gleichung folgt dann weiters

$$T \cdot \frac{p_0}{T_0} = p = \frac{2}{3} \cdot \frac{N}{V} \cdot \overline{E_{\text{kin}}} \quad \Rightarrow \quad \overline{E_{\text{kin}}} = \underbrace{\frac{3}{2} \cdot \frac{p_0}{T_0} \cdot \frac{V}{N}}_{k_B} \cdot T = \frac{3}{2} \cdot k_B \cdot T^{120}$$

Die Größe $k_B = \frac{p_0}{T_0} \cdot \frac{V}{N}$ wird als Boltzmann-Konstante¹²¹ bezeichnet. In die Bernoulli-Gleichung eingesetzt ergibt sich

¹¹⁹ Normalbedingungen nach DIN 1343 sind: $p_0 = 101325 \text{ Pa}$ und $T_0 = 273,15 \text{ K}$

¹²⁰ Da $\overline{E_{\text{kin}}} \sim T$ ist der Faktor $\frac{p_0}{T_0} \cdot \frac{V}{N}$ eine Konstante.

$$p = \frac{2}{3} \cdot \frac{N}{V} \cdot \overline{E_{\text{kin}}} = \frac{2}{3} \cdot \frac{N}{V} \cdot \frac{3}{2} \cdot k_B \cdot T = \frac{N}{V} \cdot k_B \cdot T$$

und damit

$$p \cdot V = N \cdot k_B \cdot T.$$

Die übliche Form der allgemeinen Gasgleichung ergibt sich, indem man eine neue Konstante R einführt, die sogenannte universelle Gaskonstante¹²².

Es ist $N = n \cdot N_A$, wobei N_A die Avogadro-Konstante und n die Stoffmenge in mol ist.

Damit lässt sich die Gasgleichung schreiben als $p \cdot V = n \cdot N_A \cdot k_B \cdot T$. Man setzt $R = N_A \cdot k_B$ und erhält das Gasgesetz in einer zweiten Formulierung als

$$p \cdot V = n \cdot R \cdot T$$

Wenn man sich die Herleitung der Bernoulli-Gleichung und davon ausgehend der makroskopischen Zustandsgleichungen ansieht, ist das Entscheidende, dass man mittels statistischer Methoden aufgrund einer großen Anzahl von Teilchen die Schlussfolgerungen auf messbare makroskopische Größen ziehen kann. Es sind die Geschwindigkeiten beziehungsweise Impulse der einzelnen Gasteilchen nicht bekannt und können auch nicht gemessen werden. Trotzdem können aber daraus Gleichungen hergeleitet werden, die in sehr großer Genauigkeit durch Messung makroskopischer Größen verifiziert werden können. Diese Verbindung zwischen Mikro- und Makrophysik muss als einer der ganz großen Durchbrüche in der Physik angesehen werden.

Es soll nur noch erwähnt werden, dass die Vereinfachungen, die am Anfang getroffen worden sind, dazu geführt haben, dass eine Reihe komplizierterer Zustandsgleichungen aufgestellt worden sind. Eine der bekannteren ist die Van-der-Waals'sche

¹²¹ Die Boltzmann-Konstante ergibt sich, indem man den Normaldruck $p_0 = 101325 \text{ Pa}$, die Normaltemperatur $T_0 = 273,15 \text{ K}$ einsetzt und benutzt, dass ein Gas unter Normalbedingungen 1 mol die Anzahl $N = 6,022 \cdot 10^{23}$ von Teilchen enthält und das Volumen $V = 0,024 \text{ m}^3$ einnimmt. Daraus ergibt sich die Boltzmann-Konstante $k_B = 1,38 \cdot 10^{-23} \text{ J} \cdot \text{K}^{-1}$.

¹²² Die universelle Gaskonstante ergibt sich aus dem Produkt aus Avogadro-Konstante und Boltzmann-Konstante als $R = 8,314462 \text{ J} \cdot \text{K}^{-1} \cdot \text{mol}^{-1}$

Zustandsgleichung, die das Eigenvolumen und die gegenseitige Anziehung der Teilchen mitberücksichtigt. Auch gelten die Gasgesetze bei sehr tiefen Temperaturen nahe dem absoluten Nullpunkt nicht mehr.¹²³

Zur Illustration, wie sich diese Größe, Eigenvolumen und gegenseitige Anziehung, auf allgemeine Gasgleichung auswirken, sei die Van-der-Waals-Gleichung angeführt:

$$\left(p + \frac{a \cdot n^2}{V^2}\right) \cdot (V - b \cdot n) = n \cdot R \cdot T$$

Man erkennt, dass zwei zusätzliche Koeffizienten a und b vorkommen, die folgende Bedeutung haben.

1. Der Koeffizient b kommt durch das nichtverschwindende Eigenvolumen der Gasteilchen zustande. Durch das endliche Volumen der Gasteilchen ist das Volumen, das die Teilchen zur Bewegung zur Verfügung haben, kleiner als in der allgemeinen Gasgleichung für das ideale Gas. Der Ausdruck $b \cdot n$ entspricht gerade dem Volumen aller Gasteilchen, um welches sich das Volumen verringert, somit ändert sich das Volumen V in der allgemeinen Gasgleichung in den Ausdruck $V - b \cdot n$. Da n angibt, wie viele Mole Gasteilchen das Gas enthalten, gibt der Betrag von b das Volumen der Teilchen eines Mols des Gases an.
2. Der zweite Term, der den Druck in der Gasgleichung eines idealen Gases abändert, also $\frac{a \cdot n^2}{V^2}$ berücksichtigt, dass sich die Gasteilchen gegenseitig anziehen. Wenn ein Teilchen sich einer Wand annähert, wird es durch die Anziehung der übrigen Teilchen mit einer Kraft zurückgezogen, die proportional zur Teilchendichte, also $\frac{n}{V}$, ist. Die Zahl der Gasteilchen, die in einem bestimmten Zeitintervall die Wand treffen, ist auch proportional zur Teilchendichte, also wieder zu $\frac{n}{V}$. Deshalb ist die

¹²³ Ich möchte anmerken, dass ich diese Herleitungen in zwei 7. Klassen in den Jahren 2013/14 gebracht habe und es von den SchülerInnen zwar als etwas kompliziert empfunden worden ist, aber doch das Entscheidende verstanden worden ist. Es ist mir nicht darauf angekommen, dass die SchülerInnen bei einer schriftlichen Überprüfung die Herleitung genauso niederschreiben konnten, sondern ob sie das Wesentliche daraus verstanden haben, was bei einem Teil der SchülerInnen sicherlich der Fall war. Ich möchte aber explizit noch einmal darauf hinweisen, dass sich die LehrerIn klar sein muss, ob es dem Niveau der Klasse entspricht oder ob sie lieber auf die mathematische Begründung verzichten möchte.

gesamte Druckabnahme proportional zum Quadrat der Teilchendichte, was durch eine multiplikative Konstante a ausgedrückt wird.

Es seien die Koeffizienten für einige Gase angeführt, siehe dazu Tabelle 7.¹²⁴

Gasteilchen	a [in $\text{Pa} \cdot \text{m}^6 \cdot \text{mol}^{-2}$]	b [in $\text{m}^3 \cdot \text{mol}^{-1}$]
Helium	0,0035	$2,38 \cdot 10^{-5}$
Argon	0,136	$3,22 \cdot 10^{-5}$
Wasserstoff (H_2)	0,0247	$2,66 \cdot 10^{-5}$
Stickstoff (N_2)	0,1388	$3,87 \cdot 10^{-5}$
Sauerstoff (O_2)	0,14	$3,186 \cdot 10^{-5}$
Wasserdampf (H_2O)	0,53	$3,05 \cdot 10^{-5}$
Kohlendioxid (CO_2)	0,364	$4,27 \cdot 10^{-5}$

Tabelle 7: Einige Koeffizienten a und b der Van-der-Waals-Gleichung

¹²⁴ Zur Van-der-Waals-Gleichung siehe etwa [25] ab Seite 777; die Koeffizienten wurden auch aus [25] entnommen, Seite 778.

3.7 7. Klasse

3.7.1 Differentialrechnung

Die Differentialrechnung bildet eines der entscheidenden Fundamente der modernen Analysis in der Mathematik genauso wie sie eine der wichtigsten Methoden in der Physik ist. Der Grund dafür liegt darin begründet, dass in der Physik oft Änderungen von Größen betrachtet werden, wofür die Differentialrechnung das geeignete Mittel zur Beschreibung darstellt.

Betrachtet man die Geschichte der Differentialrechnung, so hat die Differentialrechnung ihre Wurzeln sowohl in der Mathematik wie in der Physik gleichermaßen. Ging der eine Begründer der Differentialrechnung, Leibniz¹²⁵, vom geometrischen Problem der Tangente an eine Kurve aus, also von Seiten der Mathematik, so tat dies der zweite Begründer, Newton¹²⁶, von der physikalischen Seite des Problems der Momentangeschwindigkeit eines bewegten Körpers. Die Differentialrechnung wurde sehr schnell in der Physik benutzt und zeitigte große Erfolge, obwohl sie von vielen Mathematikern der damaligen Zeit als unlogisch angesehen worden ist, da der Begriff der von Leibniz und Newton verwendeten infinitesimalen Größen nicht mathematisch exakt fassbar war. Euler¹²⁷ entwickelte bald darauf die auch heute so verwendeten Ableitungsregeln.

Es gelang aber erst Cauchy¹²⁸ am Anfang des 18. Jahrhunderts, die Differentialrechnung mathematisch exakt und widerspruchsfrei zu formulieren. Entscheidend dafür war, dass Cauchy nicht mehr die infinitesimalen Größen verwendete, sondern den Grenzwertbegriff und den Begriff des Differenzenquotienten. Die heutige Form der Differentialrechnung wurde erst Ende des 19. Jahrhunderts schließlich durch Weierstraß¹²⁹ vollendet, der die auch heute gültige Definition des Grenzwertes in die Mathematik eingeführt hat. Das Entscheidende ist aber, dass die Differentialrechnung sehr früh mit der Entwicklung der mathematischen Physik, die mit Newton angesetzt werden kann, als notwendiges Werkzeug ihren Eingang in die Physik gefunden hat.

¹²⁵ Gottfried Wilhelm Leibniz; 1646 – 1716; deutscher Mathematiker und Philosoph

¹²⁶ Isaac Newton; 1642 – 1726; englischer Naturwissenschaftler und Philosoph

¹²⁷ Leonhard Euler; 1707 – 1783; schweizer Mathematiker und Physiker

¹²⁸ Augustin-Louis Cauchy; 1789 – 1857; französischer Mathematiker

¹²⁹ Karl Theodor Wilhelm Weierstraß; 1815 – 1897; deutscher Mathematiker

Beispiel 1: Fermat'sches Prinzip und Reflexionsgesetz

Bekannterweise lautet das Reflexionsgesetz eines Lichtstrahls, der an einer Fläche gespiegelt wird, dass der einfallende und der reflektierte Strahl dieselben Winkel zum Lot einschließen, also $\alpha = \beta$ (Abbildung 51).

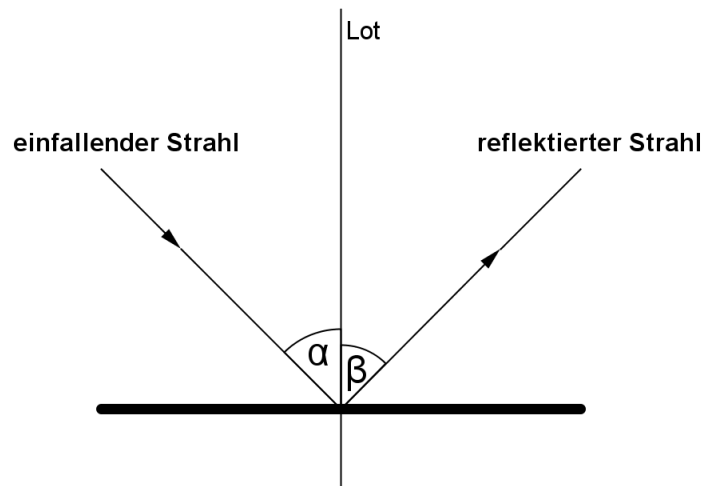


Abbildung 51: Reflexion eines Lichtstrahls

Um das Reflexionsgesetz herzuleiten, kann als einfachste Möglichkeit das Fermat'sche Prinzip verwendet werden. Das Fermat'sche Prinzip besagt, dass ein Lichtstrahl, der zwischen zwei Punkten verläuft, dies auf jenem Weg tut, für den er am wenigsten Zeit benötigt.

Dieses wichtige Prinzip der Physik ist ein Minimalprinzip, wie es auch andere gibt, wie etwa das Prinzip der minimalen potenziellen Energie, die ein Körper zu erreichen versucht. Diese Prinzipien haben den ungemein nützlichen Vorteil, dass man bei ihrer Verwendung oft interne und komplizierte physikalische Vorgänge nicht berücksichtigen, ja nicht einmal genau kennen muss, und doch zu einer Aussage über den Endzustand eines Systems gelangen kann.

Die physikalische Erklärung, weshalb das Fermat'sche Prinzip gilt, liegt darin begründet, dass man von der Vorstellung ausgeht, dass ein Lichtstrahl nicht nur einen Weg, eben den die benötigte Zeit minimalisierenden, nimmt, sondern alle Wege gleichzeitig. Wenn man nun alle diese Wege superponiert, das heißt überlagert, so löschen sich alle Wege durch destruktive Interferenz aus bis auf den einen, den dann der

Lichtstrahl tatsächlich nimmt. In der Quantenmechanik ist diese Vorstellung als die Pfadintegralmethode bekannt, die von Feynman¹³⁰ entwickelt worden war.

Es soll nun im Folgenden mit Hilfe des Fermat'schen Prinzips das Reflexionsgesetz hergeleitet werden. Es liege folgende Situation vor, in der ein Lichtstrahl von einem Punkt A zu einem Punkt B gelangt, indem der Lichtstrahl an einem Spiegel reflektiert wird. Den Punkt B kann man sich als das Auge eines Betrachters vorstellen, der etwa ein gespiegeltes Bild betrachtet, von dem ein einzelner gespiegelter Punkt beobachtet wird. Die beiden (Normal-)Abstände d_1 des Punktes A und d_2 des Punktes B vom Spiegel könnten auch ohne Beschränkung der Allgemeinheit des Ergebnisses als gleich angenommen werden, das es ja keine Rolle spielt, wo genau auf dem gespiegelten Lichtstrahl der Punkt B liegt. Es soll aber trotzdem der allgemeinere Fall angenommen werden, dass die beiden Punkte A und B verschiedene Normalabstände vom Spiegel haben. Die geometrischen Verhältnisse sind aus der Abbildung 52 ersichtlich:

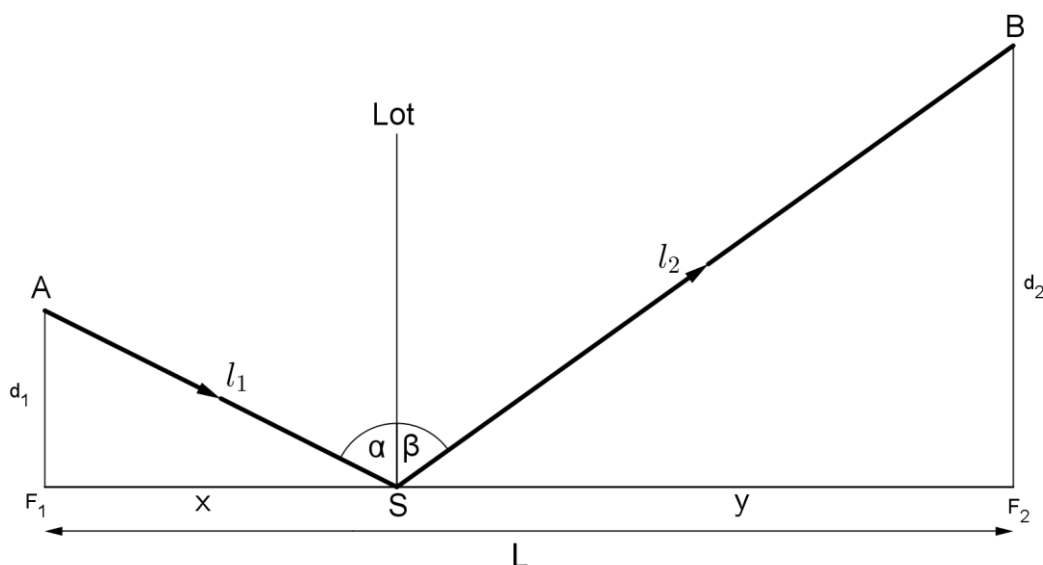


Abbildung 52: Zur Herleitung des Reflexionsgesetzes

Der Lichtstrahl geht vom Punkt A geradlinig zum Punkt S auf dem Spiegel und wird dort reflektiert. Sodann gelangt er vom Punkt S weiter geradlinig zum Punkt B. Der Strahl legt von A nach S die Strecke l_1 und von S nach B die Strecke l_2 zurück. Die Strecke vom Fußpunktes F_1 von A zum Punkt S sei x und die Strecke von S zum

¹³⁰ Richard Feynman; 1918 – 1988; amerikanischer Physiker

Fußpunkt F_2 von B sei y. Die Summe von x und y sei L und konstant. Es wird weiters angenommen, dass das Medium (normalerweise Luft), in dem sich der Lichtstrahl ausbreitet, homogen ist, sodass die Lichtgeschwindigkeit im Medium als konstant angenommen werden kann.

Aus der Abbildung 52 ergibt sich die gesamte zurückgelegte Strecke $l = l_1 + l_2$ des Lichtstrahls als

$$l = \sqrt{d_1^2 + x^2} + \sqrt{d_2^2 + (L - x)^2}.$$

Um diese Strecke zurückzulegen, benötigt der Lichtstrahl die Zeit $t = \frac{l}{c}$, wobei c die Lichtgeschwindigkeit im Medium ist. Also ergibt sich für die Zeit der folgende Ausdruck

$$t(x) = \frac{1}{c} \cdot \left(\sqrt{d_1^2 + x^2} + \sqrt{d_2^2 + (L - x)^2} \right).$$

Dieser Ausdruck ist nun laut dem Fermat'schen Prinzip zu minimalisieren, das heißt es ist mittels Differentialrechnung das Minimum zu bestimmen. Dazu wird die erste Ableitung berechnet und diese dann Null gesetzt.

$$t'(x) = \frac{1}{c} \cdot \left(\frac{2 \cdot x}{2 \cdot \sqrt{d_1^2 + x^2}} + \frac{2 \cdot (L - x) \cdot (-1)}{2 \cdot \sqrt{d_2^2 + (L - x)^2}} \right) = \frac{1}{c} \cdot \left(\frac{x}{\sqrt{d_1^2 + x^2}} - \frac{(L - x)}{\sqrt{d_2^2 + (L - x)^2}} \right) \stackrel{!}{=} 0$$

$$\frac{x}{\sqrt{d_1^2 + x^2}} - \frac{(L - x)}{\sqrt{d_2^2 + (L - x)^2}} = 0$$

$$x \cdot \sqrt{d_2^2 + (L - x)^2} = (L - x) \cdot \sqrt{d_1^2 + x^2}$$

$$x^2 \cdot (d_2^2 + (L - x)^2) = (L - x)^2 \cdot (d_1^2 + x^2)$$

$$x^2 \cdot d_2^2 + \cancel{x^2 \cdot (L - x)^2} = (L - x)^2 \cdot d_1^2 + \cancel{(L - x)^2 \cdot x^2}$$

$$x^2 \cdot (d_2^2 - d_1^2) + 2 \cdot L \cdot d_1^2 \cdot x - L^2 \cdot d_1^2 = 0$$

Hieraus ergibt sich folgende quadratische Gleichung

$$x^2 + \frac{2 \cdot L \cdot d_1^2}{d_2^2 - d_1^2} \cdot x - \frac{L^2 \cdot d_1^2}{d_2^2 - d_1^2} = 0$$

Eingesetzt in die „kleine Lösungsformel“ ergibt sich:

$$\begin{aligned}
x_{1,2} &= -\frac{L \cdot d_1^2}{d_2^2 - d_1^2} \pm \sqrt{\left(\frac{L \cdot d_1^2}{d_2^2 - d_1^2}\right)^2 + \frac{L^2 \cdot d_1^2}{d_2^2 - d_1^2}} \\
&= -\frac{L \cdot d_1^2}{d_2^2 - d_1^2} \pm \sqrt{\frac{\cancel{L^2 \cdot d_1^4} + L^2 \cdot d_1^2 \cdot d_2^2 - \cancel{L^2 \cdot d_1^4}}{(d_2^2 - d_1^2)^2}} \\
&= -\frac{L \cdot d_1^2}{d_2^2 - d_1^2} \pm \frac{L \cdot d_1 \cdot d_2}{d_2^2 - d_1^2} \\
&= \frac{L \cdot d_1 \cdot (\pm d_2 - d_1)}{(d_2 - d_1) \cdot (d_2 + d_1)}
\end{aligned}$$

Daraus ergeben sich die beiden Lösungen:

$$\begin{aligned}
x_1 &= \frac{L \cdot d_1 \cdot \cancel{(d_2 - d_1)}}{\cancel{(d_2 - d_1)} \cdot (d_2 + d_1)} = \frac{L \cdot d_1}{d_2 + d_1} \\
x_2 &= \frac{-L \cdot d_1 \cdot \cancel{(d_2 + d_1)}}{(d_2 - d_1) \cdot \cancel{(d_2 + d_1)}} = -\frac{L \cdot d_1}{(d_2 - d_1)}
\end{aligned}$$

Für die y-Strecken ergeben sich:

$$\begin{aligned}
y_1 &= L - x_1 = L - \frac{L \cdot d_1}{d_2 + d_1} = \frac{L \cdot d_2 + L \cdot d_1 - L \cdot d_1}{d_2 + d_1} = \frac{L \cdot d_2}{d_2 + d_1} \\
y_2 &= L + \frac{L \cdot d_1}{(d_2 - d_1)}
\end{aligned}$$

Man sieht, dass die Lösungen x_2 und y_2 ausgeschlossen werden müssen. Denn ist $d_2 > d_1$, so ist $x_2 < 0$, ist im andern Fall $d_2 < d_1$, so ist $y_2 < 0$.

Daraus ergeben sich für die Winkel α und β folgende Beziehungen:

$$\begin{aligned}
\tan \alpha &= \frac{x_1}{d_1} = \frac{L}{d_2 + d_1} \\
\tan \beta &= \frac{y_1}{d_2} = \frac{L}{d_2 + d_1}
\end{aligned}$$

Daraus folgt schließlich, dass $\tan \alpha = \tan \beta$ ist und daraus das Reflexionsgesetz

$\alpha = \beta$

Es bleibt nur noch zu zeigen, dass das erhaltene Extremum x_1 tatsächlich eine Minimumstelle ist. Das Standardrezept dazu ist, die 2. Ableitung zu bilden und den erhaltenen Wert einzusetzen.

Ist $t''\left(\frac{L \cdot d_1}{d_2 + d_1}\right) > 0$, so liegt eine Minimumstelle vor, ist $t''\left(\frac{L \cdot d_1}{d_2 + d_1}\right) < 0$ eine Maximumstelle, im Falle gleich Null liegt eine Sattelstelle vor.

Instruktiver als diese Vorgangsweise kann es sein, sich die Steigung links und rechts der Extremstelle der Funktion anzusehen, insbesondere dann, wenn wie in diesem Beispiel die 2. Ableitung recht kompliziert ist.

Da x_1 im Intervall $[0, L]$ die einzige Extremstelle ist, bietet es sich an, die 1. Ableitung an den Intervallgrenzen zu betrachten:

$$t'(0) = \frac{1}{c} \cdot \left(\frac{x}{\sqrt{d_1^2 + x^2}} - \frac{(L-x)}{\sqrt{d_2^2 + (L-x)^2}} \right) \Big|_{x=0} = -\frac{1}{c} \cdot \frac{L}{\sqrt{d_2^2 + L^2}} < 0$$

$$t'(L) = \frac{1}{c} \cdot \left(\frac{x}{\sqrt{d_1^2 + x^2}} - \frac{(L-x)}{\sqrt{d_2^2 + (L-x)^2}} \right) \Big|_{x=L} = \frac{1}{c} \cdot \frac{L}{\sqrt{d_2^2 + L^2}} > 0$$

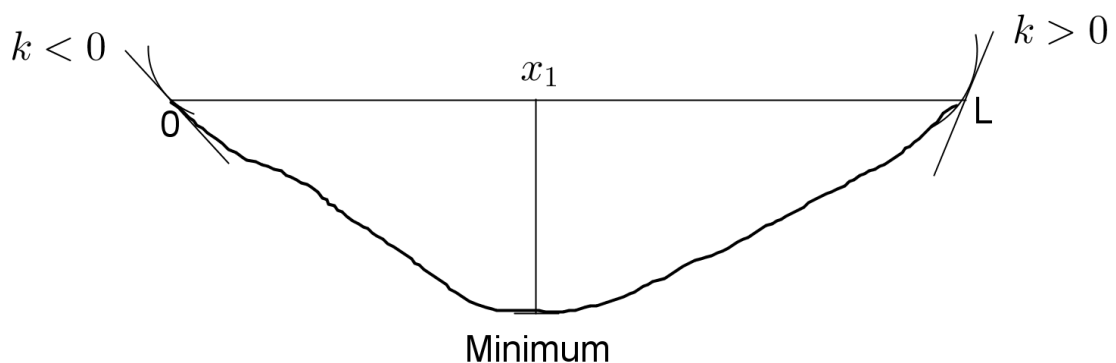


Abbildung 53: Bestimmung der Extremstelle anhand des Funktionsverlaufs

Aus der Abbildung 53 erkennt man den einzig möglichen Funktionsverlauf, insbesondere, dass die Extremstelle x_1 tatsächlich nur eine Minimumstelle sein kann. Es wurde in der Abbildung bewusst darauf verzichtet, den genauen Funktionsverlauf zu

zeichnen, da gezeigt werden soll, dass bereits aus den beiden Steigungen an den Stellen $x_1 = 0$ und $x_1 = L$ geschlossen werden kann, dass die Stelle x_1 nur eine Minimumstelle sein kann.

Man hätte statt dieser doch etwas langwierigen Herleitung, die für den Mathematikunterricht sicherlich instruktiv wäre, dies für den Physikunterricht auch kürzer herleiten können. Geht man von der ersten Ableitung aus und setzt diese null, so erhält man wie oben gezeigt folgende Bedingung für ein Extremum

$$\frac{x}{\sqrt{d_1^2 + x^2}} - \frac{(L-x)}{\sqrt{d_2^2 + (L-x)^2}} = 0$$

Man sieht nun aus der Abbildung 52: Zur Herleitung des Reflexionsgesetzes, dass folgende Beziehungen gelten:

$$\sin \alpha = \frac{x}{\sqrt{d_1^2 + x^2}} \quad \text{und} \quad \sin \beta = \frac{(L-x)}{\sqrt{d_2^2 + (L-x)^2}}$$

Und damit ergibt sich $\sin \alpha - \sin \beta = 0$ beziehungsweise $\sin \alpha = \sin \beta$, woraus sofort das Reflexionsgesetz folgt, $\alpha = \beta$. Allerdings ist mit dieser sehr kurzen Herleitung nicht gezeigt worden, dass es sich tatsächlich um ein Minimum handelt, worauf aber im Physikunterricht sicherlich verzichtet werden darf.

Beispiel 2: Fermat'sches Prinzip und Brechung

Es werde nun der Fall behandelt, dass ein Lichtstrahl auf die Grenzfläche zweier verschiedener Medien fällt und vom Medium 1 ins Medium 2 eintritt. Physikalisch unterscheiden sich zwei Medien darin, dass die Lichtgeschwindigkeit in jedem Material kleiner als die Vakuumlichtgeschwindigkeit ist¹³¹. Die Folge hiervon ist, dass ein

¹³¹ Die Lichtgeschwindigkeit in Vakuum beträgt exakt 299792,458 km/s. Die Lichtgeschwindigkeit in Luft beträgt dagegen nur 299706 km/s, ist also um 0,03% geringer als im Vakuum. In Wasser beträgt die Lichtgeschwindigkeit nur mehr 225408 km/s, das sind nur mehr etwa 75,2% der Vakuumlichtgeschwindigkeit. In Diamant sinkt die Lichtgeschwindigkeit sogar auf 123881 km/s, ist also nur mehr 41,3 % der Lichtgeschwindigkeit in Vakuum (deshalb auch das schöne Farbenspiel, das „Feuer“ bei geschliffenen Diamanten, da eine sehr starke Brechung des Lichtes auftritt).

Lichtstrahl, der in einem Winkel α zum Lot auf die Grenzfläche fällt, nach dem Übertritt in anderen Medium einen anderen Winkel β zum Lot hat. Dieses Phänomen wird Brechung genannt und ist die Voraussetzung vieler optischer Geräte wie etwa Linsen und Prismen.

Das Brechungsgesetz, das die quantitativen Zusammenhänge zwischen einfallendem und gebrochenem Strahl beschreibt, lautet

$$\frac{\sin \alpha}{\sin \beta} = \frac{v_1}{v_2}$$

Wobei v_1 und v_2 die Lichtgeschwindigkeit im Medium 1 und 2 ist. Diese Beziehung soll im Folgenden mit Hilfe des Fermat'schen Prinzips hergeleitet werden.

Die Abbildung 54 zeigt die geometrischen Verhältnisse bei einem Lichtstrahl, der von einem Punkt A aus einen Punkt B erreicht. Dabei wird die Grenzfläche, an der die beiden Medien aneinanderstoßen, durchlaufen und die Lichtgeschwindigkeit wechselt von der Geschwindigkeit v_1 im Medium 1 zu einer Geschwindigkeit v_2 im Medium 2. Gesucht ist der Weg, den der Lichtstrahl vom Punkt A zum Punkt B nimmt, wenn man annimmt, dass das Fermat'sche Prinzip gilt.

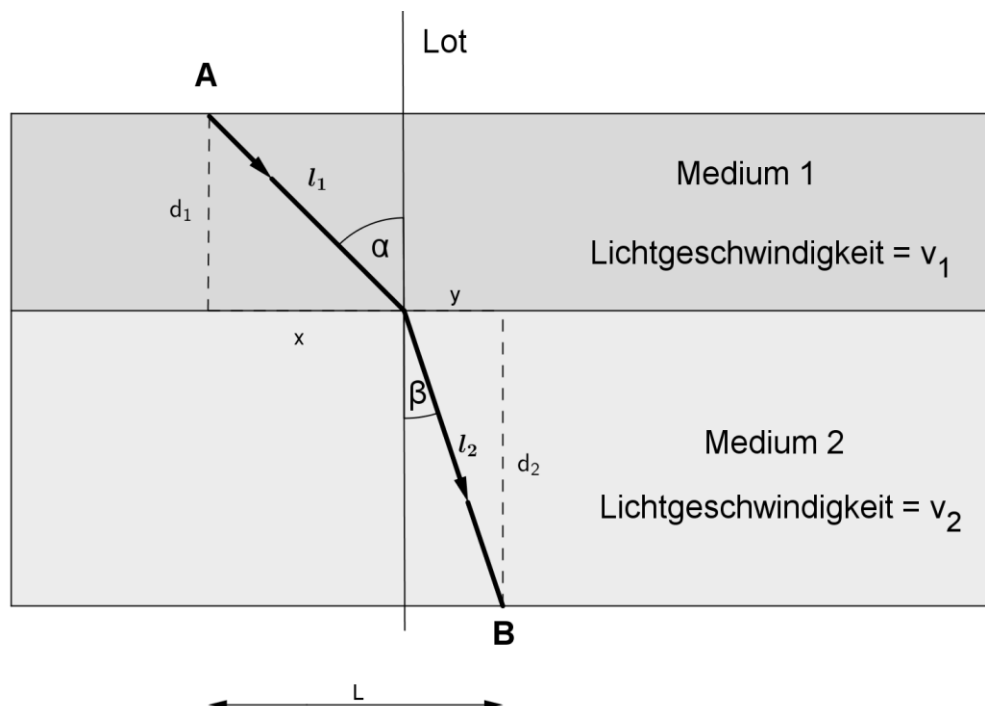


Abbildung 54: Zur Herleitung des Brechungsgesetzes

Die Zeit t , die der Lichtstrahl von A nach B benötigt, beträgt

$$t(x) = \frac{l_1}{v_1} + \frac{l_2}{v_2} = \frac{1}{v_1} \cdot \sqrt{d^2 + x^2} + \frac{1}{v_2} \cdot \sqrt{d^2 + (L-x)^2}.$$

Die 1. Ableitung nach x ist dann

$$t'(x) = \frac{x}{v_1 \cdot \sqrt{d^2 + x^2}} - \frac{L-x}{v_2 \cdot \sqrt{d^2 + (L-x)^2}}$$

Durch Nullsetzen der 1. Ableitung erhält man wiederum das Minimum.

$$\frac{x}{v_1 \cdot \sqrt{d^2 + x^2}} - \frac{L-x}{v_2 \cdot \sqrt{d^2 + (L-x)^2}} = 0$$

Es ist nun nicht notwendig, diese Gleichung nach x aufzulösen (was auch ein ziemlich schwieriges Unterfangen wäre), sondern man sieht aus der Abbildung, dass folgende geometrischen Beziehungen gelten:

$$\sin \alpha = \frac{x}{\sqrt{d^2 + x^2}} \quad \text{und} \quad \sin \beta = \frac{L-x}{\sqrt{d^2 + (L-x)^2}}$$

Also lässt sich obige Gleichung schreiben als

$$\frac{\sin \alpha}{v_1} - \frac{\sin \beta}{v_2} = 0$$

beziehungsweise als

$$\frac{\sin \alpha}{\sin \beta} = \frac{v_1}{v_2},$$

was gerade das Brechungsgesetz ist.

Üblicherweise wird das Brechungsgesetz nicht mit den Lichtgeschwindigkeiten in den Medien, sondern durch die Brechungsindizes angegeben. Der Brechungsindex wird einfach als das Verhältnis der Lichtgeschwindigkeit im Vakuum zur Lichtgeschwindigkeit im entsprechenden Medium definiert:

$$n_1 := \frac{c}{v_1} \quad \text{und} \quad n_2 := \frac{c}{v_2}$$

Somit ergibt sich dann das sogenannte Brechungsgesetz nach Snellius als

$\frac{\sin \alpha}{\sin \beta} = \frac{n_2}{n_1}$
--

Betrachten wir noch die Frage, wann der gebrochene Strahl zum Lot hin oder vom Lot weg gebrochen wird. Aus dem Brechungsgesetz lässt sich für den Sinus des gebrochenen Strahls schreiben

$$\sin \beta = \frac{n_1}{n_2} \cdot \sin \alpha.$$

Wenn der Strahl von einem dünneren Medium in ein dichteres Medium eindringt, ist der Quotient $\frac{n_1}{n_2} < 1$. Also muss $\sin \beta < \sin \alpha$ sein und damit $\beta < \alpha$, der Strahl wird also zum Lot hin gebrochen. Dass aus der Ungleichung $\sin \beta < \sin \alpha$ folgt, dass auch $\beta < \alpha$ gilt, folgt daraus, dass die Sinusfunktion für das Intervall $\left[0; \frac{\pi}{2}\right]$, in dem die Winkel α und β liegen, streng monoton wachsend ist.

Dringt der Strahl umgekehrt von einem dichteren Medium in ein dünneres Medium ein, so ist $\frac{n_1}{n_2} > 1$, damit $\sin \beta > \sin \alpha$ und damit wiederum $\beta > \alpha$. Der Strahl wird somit vom Lot weg gebrochen.

Beispiel 3: Definition der Geschwindigkeit und Beschleunigung

In der Unterstufe werden die Geschwindigkeit und die Beschleunigung noch als einfache Quotienten definiert. Damit kann natürlich nur die geradlinig gleichförmige Bewegung rechnerisch erfasst werden, was aber für die Unterstufe auch völlig ausreichend ist, es ist also

$$v = \frac{s}{t} \quad \text{und} \quad a = \frac{v}{t} \quad ^{132}$$

v ... Geschwindigkeit des Körpers, $[v] = \text{m} \cdot \text{s}^{-1}$

a ... Beschleunigung des Körpers, $[a] = \text{m} \cdot \text{s}^{-2}$

s ... zurückgelegter Weg, $[s] = \text{m}$

t ... benötigte Zeit, $[t] = \text{s}$

Im Schulbuch Prisma 2 Physik 2 wird die Beschleunigung als eine Beschleunigung aus dem Stand heraus definiert¹³³. In der Definitionsformel $a = \frac{v}{t}$ wird die Geschwindigkeit v als Momentangeschwindigkeit („die Geschwindigkeit, die der Tachometer im Auto anzeigt“) und die Zeit t als Beschleunigungszeit erklärt. Das ist eine Erklärung, die so für diese Altersstufe genügt, da hier sogar zwischen der Durchschnitts- und der Momentangeschwindigkeit unterschieden wird.

In der Oberstufe werden die Begriffe der Geschwindigkeit und Beschleunigung erneut aufgegriffen und exaktifiziert. Während manche Schulbücher wie der Big Bang von Apolin den Übergang zur Momentangeschwindigkeit und Momentanbeschleunigung mathematisch nicht weiter ausführen, macht dies aber etwa das Schulbuch Physik compact 5 in einem Erweiterungskapitel¹³⁴. Es wird im Kapitel 3.4 die mittlere Geschwindigkeit definiert als

$$\text{Geschwindigkeit} = \frac{\text{zurückgelegter Weg}}{\text{dazu benötigte Zeit}}$$

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

¹³² Für die Geschwindigkeit wird in den Schulbüchern der Unterstufe oft $v = \frac{s}{t}$ geschrieben, für die

Beschleunigung $a = \frac{v}{t}$. Bei der Geschwindigkeitsdefinition gibt es keine Probleme, da s immer eine Wegstrecke, also einen in der Zeit t zurückgelegten Weg angibt. Für die Beschleunigung sollte man aber diese Schreibweise nicht verwenden, da sie zu falschen Anwendungen führen kann. Sei etwa folgende Angabe gegeben: Ein Körper bewegt sich in der Zeit t = 5 s gleichförmig um eine Wegstrecke 100 m weiter. Für die Geschwindigkeit ergibt sich dann korrekterweise $v = 20 \text{ m/s}$. Wird jetzt aber diese Geschwindigkeit in die Gleichung für die Beschleunigung eingesetzt, ergibt sich $a = \frac{20}{5} = 4 \frac{\text{m}}{\text{s}^2}$, was natürlich Unsinn ist, da die

Beschleunigung bei einer geradlinig gleichförmigen Bewegung null ist. Der Fehler liegt natürlich darin begründet, dass hier keine Geschwindigkeitsdifferenz in die Gleichung für die Beschleunigung eingesetzt wurde, die natürlich null wäre, woraus sich dann die korrekte Beschleunigung $a = 0 \text{ m/s}^2$ ergibt.

¹³³ Siehe [29] Seite 56

¹³⁴ Siehe [22] Kapitel 3.14 Seite 57

$\vec{v}$... mittlere Geschwindigkeit, $[\vec{v}] = \text{m} \cdot \text{s}^{-1}$

$\Delta \vec{s}$... Wegdifferenz, $[\Delta \vec{s}] = \text{m}$

Δt ... Zeitintervall, $[\Delta t] = \text{s}$

Es wird nur verbal erklärt,

*Die mittlere Geschwindigkeit wird umso besser der tatsächlichen Geschwindigkeit entsprechen, je kleiner das Zeitintervall Δt gewählt wird.*¹³⁵

Im selben Kapitel wird sodann die mittlere Beschleunigung definiert als

$$\text{Beschleunigung} = \frac{\text{Geschwindigkeitsänderung}}{\text{dazu benötigte Zeit}}$$

$$\vec{a} = \frac{\Delta \vec{v}}{\Delta t}$$

$\vec{a}$... mittlere Beschleunigung, $[\vec{a}] = \text{m} \cdot \text{s}^{-2}$

$\Delta \vec{v}$... Geschwindigkeitsdifferenz, $[\Delta \vec{v}] = \text{m} \cdot \text{s}^{-1}$

Δt ... Zeitintervall, $[\Delta t] = \text{s}$

Und wieder wird verbal erklärt,

*Wieder wird diese mittlere Beschleunigung umso genauer der tatsächlichen Beschleunigung entsprechen, je kleiner das Zeitintervall Δt ist.*¹³⁶

Um jetzt auch mathematisch exakt zu einer brauchbaren Fassung dieser beiden verbalen Annäherungen an die tatsächliche Geschwindigkeit und Beschleunigung, das heißt der Momentangeschwindigkeit und Momentanbeschleunigung, zu gelangen, muss der Grenzübergang $\Delta t \rightarrow 0$ vollzogen werden.

$$\text{Momentangeschwindigkeit: } \vec{v} = \lim_{\Delta t \rightarrow 0} \frac{\Delta \vec{s}}{\Delta t} = \frac{d\vec{s}(t)}{dt} = \vec{s}'(t)$$

$$\text{Momentanbeschleunigung: } \vec{a} = \lim_{\Delta t \rightarrow 0} \frac{\Delta \vec{v}}{\Delta t} = \frac{d\vec{v}(t)}{dt} = \vec{v}'(t)$$

¹³⁵ Siehe [22] Seite 41

¹³⁶ Siehe [22] Seite 41

Es kann jetzt auch im Weiteren ein direkter Zusammenhang zwischen der Momentanbeschleunigung und dem Weg hergestellt werden.

$$\bar{a} = \frac{d\bar{v}(t)}{dt} = \frac{d}{dt} \left(\frac{d\bar{s}(t)}{dt} \right) = \frac{d^2\bar{s}(t)}{dt^2} = \bar{s}''(t)$$

In der Physik ist es üblich, die Zeitableitungen mit einem Punkt über der physikalischen Größe zu schreiben, also $\dot{\bar{s}}(t)$ statt $\bar{s}'(t)$, $\dot{\bar{v}}(t)$ statt $\bar{v}'(t)$ und $\ddot{\bar{s}}(t)$ statt $\bar{s}''(t)$.

Zusammengefasst gelten also folgende grundlegenden Beziehungen zwischen den Größen Weg, Momentangeschwindigkeit und Momentanbeschleunigung in „physikalischer Schreibweise“:

$\begin{aligned}\bar{v}(t) &= \dot{\bar{s}}(t) \\ \bar{a}(t) &= \dot{\bar{v}}(t) = \ddot{\bar{s}}(t)\end{aligned}$
--

Damit kann nun für jede Bewegung, falls der zurückgelegte Weg als Funktion der Zeit gegeben ist, die entsprechende Momentangeschwindigkeit und Momentanbeschleunigung für jeden Zeitpunkt berechnet werden.

Für die Lösung der umgekehrten Aufgabe, das heißt aus einer vorgegebenen Beschleunigungsfunktion beziehungsweise Geschwindigkeitsfunktion die Funktion des Weges herzuleiten, ist die Integralrechnung notwendig. Damit lassen sich dann alle kinematischen¹³⁷ Aufgaben lösen, das heißt die Bestimmung der Größen Weg, Geschwindigkeit und Beschleunigung aus den anderen.

Beispiel 4: Einige weitere Größen in der Physik, die durch eine Ableitung definiert sind

Viele Größen in der Physik werden durch Ableitungen definiert, oft sind es Zeitableitungen, was darin begründet liegt, dass die zeitliche Änderung einer Größe oftmals eine wichtige physikalische Bedeutung hat. Es sollen in diesem Beispiel nur

¹³⁷ Die Kinematik ist ein Teilgebiet der Physik, das sich mit Bewegungsaufgaben beschäftigt, ohne die dabei auftretenden oder zugrundeliegenden Kräfte zu berücksichtigen. Werden auch die Kräfte mit einbezogen, spricht man in der Physik von der Kinetik.

einige Größen angegeben werde, um die Wichtigkeit der Differentialrechnung in der Physik vor Augen zu stellen.

Die Leistung P:

Die Leistung ist der Quotient aus der verrichteten Arbeit und der Zeit, in der diese Arbeit verrichtet wird. Im Falle einer konstanten Leistung kann dies als einfacher Quotient geschrieben werden.

$$P = \frac{W}{t}$$

W ... verrichtete Arbeit, [W] = J (Joule)

P ... Leistung, [P] = J · s⁻¹ = Watt (W)

Ist die Leistung zeitlich nicht konstant, also veränderlich, kann die Momentanleistung als Differentialquotient geschrieben werden als

$$P(t) = \frac{dW}{dt} = \dot{W}$$

Die Winkelgeschwindigkeit ω

Die Winkelgeschwindigkeit eines rotierenden Körpers gibt an, wie schnell sich der Drehwinkel mit der Zeit ändert. Für eine gleichförmige Drehbewegung kann die Winkelgeschwindigkeit wieder als einfacher Quotient geschrieben werden.

$$\omega = \frac{\varphi}{t}$$

φ ... Drehwinkel, ein Winkel ist immer dimensionslos

ω ... Winkelgeschwindigkeit, [ω] = s⁻¹

Ist die Änderung des Drehwinkels zeitlich nicht konstant, kann die momentane Winkelgeschwindigkeit geschrieben werden als

$$\omega(t) = \frac{d\varphi(t)}{dt} = \dot{\varphi}(t)$$

Die Winkelgeschwindigkeit ist eigentlich eine vektorielle Größe, wobei die Richtung des Winkelgeschwindigkeitsvektors parallel zur Drehachse gerichtet ist. Somit schreibt sich die Definition der Winkelgeschwindigkeit als

$$\vec{\omega}(t) = \frac{d\vec{\varphi}(t)}{dt} = \dot{\vec{\varphi}}(t)$$

wobei dem Drehwinkel $\vec{\varphi}$ ein Vektor zugeordnet wird, der normal auf die Rotationsebene steht.

Die Elektrische Stromstärke I:

Die elektrische Stromstärke ist als die zeitliche Änderung der Ladung definiert.

$$I(t) = \frac{dQ(t)}{dt} = \dot{Q}(t)$$

Q ... elektrische Ladung, $[Q] = A \cdot s = \text{Coulomb (C)}$

I ... elektrische Stromstärke, $[I] = \text{Ampere (A)}$

Neben diesen Beispielen gibt es noch zahlreiche weitere Größen, die durch einen Differentialquotienten definiert werden, aber diese drei Beispiele sollten genügen, die Wichtigkeit der Differentialrechnung vor Augen zu führen.

Beispiel 5: Geschwindigkeit und Beschleunigung bei einer Drehbewegung

Drehbewegungen spielen in der Natur eine sehr wichtige Rolle, man denke zum Beispiel an die Bewegung eines Planeten um die Sonne oder die Eigenrotation eines Planeten. Es soll in diesem Beispiel angenommen werden, dass sich ein Körper auf

einer Kreisbahn mit dem Abstand r zur Rotationsachse bewegt. Üblicherweise würde man in der Physik für ein solches Problem Polarkoordinaten verwenden, es soll aber in kartesischen Koordinaten behandelt werden, da diese den SchülerInnen meist geläufiger sind. Die Bewegung des Körpers soll mit der Zeit t parametrisiert werden und hieraus sollen dann die Geschwindigkeits- und Beschleunigungsvektoren, insbesondere auch deren Richtung, bestimmt werden.

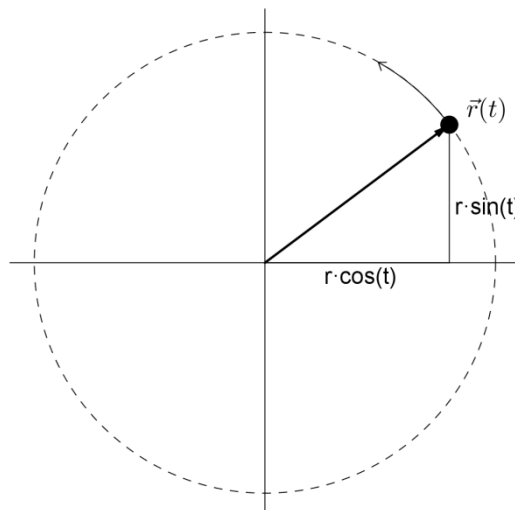


Abbildung 55: Drehbewegung einer Punktmasse

In der Abbildung 55 ist die Situation der Drehbewegung veranschaulicht. Der Körper sei als Massenpunkt durch den rotierenden Radiusvektor $\vec{r}(t)$ beschrieben. $\vec{r}(t)$ kann entweder durch seine kartesischen Koordinaten $(x|y)$ oder aber durch seine Polarkoordinaten $(r|\varphi)$ angegeben werden. Betrachten wir seine kartesischen Koordinaten, so ergibt sich aus obiger Abbildung direkt

$$\vec{r}(t) = \begin{pmatrix} r \cdot \cos(t) \\ r \cdot \sin(t) \end{pmatrix}$$

beziehungsweise in Koordinatenform geschrieben

$$x = r \cdot \cos(t)$$

$$y = r \cdot \sin(t)$$

Dabei soll r konstant sein und der Parameter t eine reelle Zahl größer null. Somit startet die Kreisbewegung für $t = 0$ beim Punkt $(r|0)$.

Es soll nun die Geschwindigkeit dieser Kreisbewegung berechnet werden. Für die Geschwindigkeit gilt

$$\vec{v}(t) = \dot{\vec{r}}(t)$$

beziehungsweise in Koordinatenform geschrieben

$$v_x(t) = \dot{x}(t) \quad \text{und} \quad v_y(t) = \dot{y}(t).$$

Damit ergibt sich für die Geschwindigkeitskomponenten

$$v_x(t) = \frac{d}{dt}(r \cdot \cos(t)) = -r \cdot \sin(t)$$

$$v_y(t) = \frac{d}{dt}(r \cdot \sin(t)) = r \cdot \cos(t)$$

Der Geschwindigkeitsvektor kann somit folgendermaßen geschrieben werden

$$\vec{v}(t) = \begin{pmatrix} -r \cdot \sin(t) \\ r \cdot \cos(t) \end{pmatrix}$$

Der Geschwindigkeitsvektor steht immer normal auf dem Radiusvektor, was am einfachsten dadurch gezeigt werden kann, dass das Skalarprodukt dieser beiden Vektoren Null ergibt. Damit ist der Geschwindigkeitsvektor natürlich auch immer tangential zur Kreisbahn des Körpers.

$$\vec{r}(t) \cdot \vec{v}(t) = \begin{pmatrix} r \cdot \cos(t) \\ r \cdot \sin(t) \end{pmatrix} \cdot \begin{pmatrix} -r \cdot \sin(t) \\ r \cdot \cos(t) \end{pmatrix} = -r^2 \cdot \cos(t) \cdot \sin(t) + r^2 \cdot \sin(t) \cdot \cos(t) = 0$$

Es soll nun noch die Beschleunigung der Kreisbewegung berechnet werden. Für die Beschleunigung gilt

$$\vec{a}(t) = \dot{\vec{v}}(t) = \ddot{\vec{r}}(t)$$

beziehungsweise in Koordinatenform geschrieben

$$a_x(t) = \dot{v}_x(t) = \ddot{x}(t) \quad \text{und} \quad a_y(t) = \dot{v}_y(t) = \ddot{y}(t).$$

Damit ergibt sich für die Beschleunigungskomponenten

$$a_x(t) = \frac{d}{dt}(-r \cdot \sin(t)) = -r \cdot \cos(t)$$

$$a_y(t) = \frac{d}{dt}(r \cdot \cos(t)) = -r \cdot \sin(t)$$

Der Beschleunigungsvektor kann somit folgendermaßen geschrieben werden

$$\vec{a}(t) = \begin{pmatrix} -r \cdot \cos(t) \\ -r \cdot \sin(t) \end{pmatrix}$$

Der Beschleunigungsvektor zeigt damit immer in die entgegengesetzte Richtung wie der Radiusvektor, er ist also immer zur Rotationsachse hin gerichtet. Mit anderen Worten heißt das, dass sich der Körper, wenn er sich auf einer Kreisbahn bewegt, immer im „freien Fall“ auf die Rotationsachse befindet. Für die Bewegung des Mondes um die Erde heißt das etwa, dass der Mond ununterbrochen im freien Fall auf die Erde fällt.

Es seien die Geschwindigkeit und die Beschleunigung noch in der Abbildung 56 veranschaulicht.

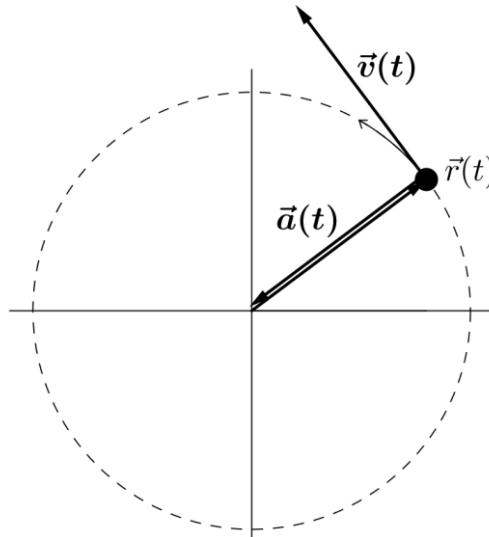


Abbildung 56: Geschwindigkeit und Beschleunigung einer Drehbewegung

3.7.2 Kegelschnitte

Kegelschnitte kommen in der Physik hauptsächlich in astronomischen Fragestellungen und in der Optik vor. Bahnbewegungen von Himmelskörpern können Ellipsenbahnen, Hyperbelbahnen oder Parabelbahnen sein. Satelliten bewegen sich auf Kreisbahnen um die Erde. Parabelformen kommen etwa bei Autoscheinwerfern vor oder auch bei Radioteleskopen oder als Parabolantennen zum Empfang von Mikrowellensignalen von Telekommunikationssatelliten. Des Weiteren kommen parabolische Bahnbewegungen in der Mechanik etwa beim schiefen Wurf vor.

Beispiel 1: Halley'scher Komet

Es findet sich im Mathematikbuch der siebten Klasse von Götz, Reichel folgende instruktive Aufgabe, die die Größenverhältnisse und Bahnen der Himmelskörper im Sonnensystem sehr schön veranschaulichen hilft.¹³⁸

744	Der HALLEY'sche Komet bewegt sich auf einer ellipsenförmigen Bahn; die Sonne ist einer ihrer Brennpunkte. Am 9. Februar 1986 hatte er wieder einmal seinen sonnennächsten Punkt durchschritten. Sein Abstand von der Sonne betrug 0,587 AE (= Astronomische Einheit = Mittlere Entfernung Erde – Sonne). Seine größte Entfernung von der Sonne beträgt 35 AE. Berechne die Halbachsen der Bahnellipse des Kometen, beschreibe die Bahnkurve durch eine Gleichung und fertige eine Zeichnung an (1 AE $\triangleq$ 2 mm)! Zeichne ferner die nahezu kreisförmige Bahn der Erde und die des Jupiters (Entfernung von der Sonne 5,20 AE) ein!	
-----	--	--

Abbildung 57: Eine Aufgabe zum Halley'schen Kometen

Lösung:

Zur Veranschaulichung seien in der Abbildung 58 die wichtigsten Kenngrößen einer Ellipse in Erinnerung gerufen.

¹³⁸ [23] Seite 192, Aufgabe 744

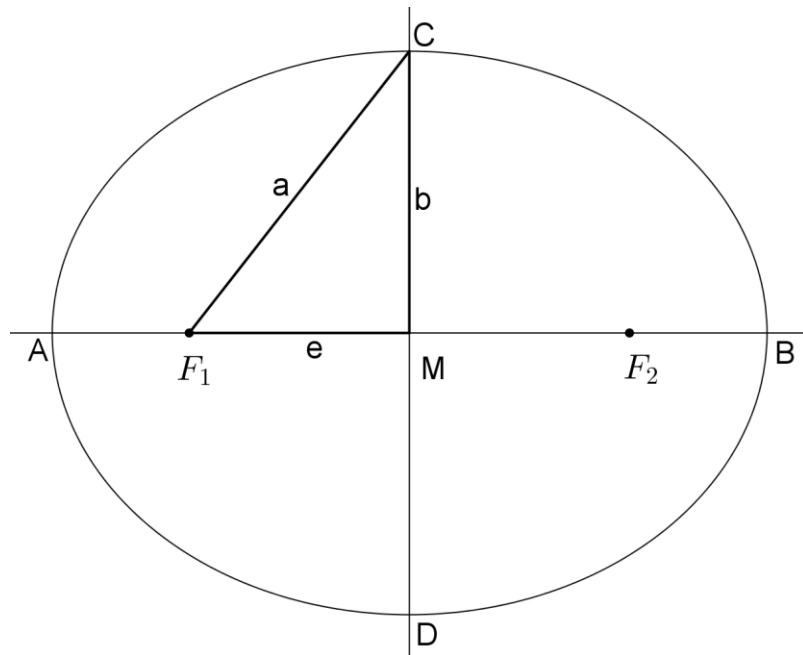


Abbildung 58: Die Ellipse

Die große Halbachse ergibt sich zu

$$2 \cdot a = 0,587 \text{ AE} + 35 \text{ AE} = 35,587 \text{ AE} \Rightarrow a \approx 17,79 \text{ AE}$$

und daraus die lineare Exzentrizität e als

$$e = a - \overline{AF_1} = 17,7935 \text{ AE} - 0,587 \text{ AE} = 17,21 \text{ AE}$$

und hieraus dann die kleine Bahnachse

$$b = \sqrt{a^2 - e^2} = \sqrt{17,7935^2 - 17,2065^2} \approx 4,53 \text{ AE}$$

Damit lässt sich die Gleichung der elliptischen Bahn des Halley'schen Kometen schreiben als

$$\frac{x^2}{316,48} + \frac{y^2}{20,52} = 1$$

Die Bahnen der Erde, des Jupiter und des Halley'schen Kometen sind in der Abbildung 59 dargestellt.

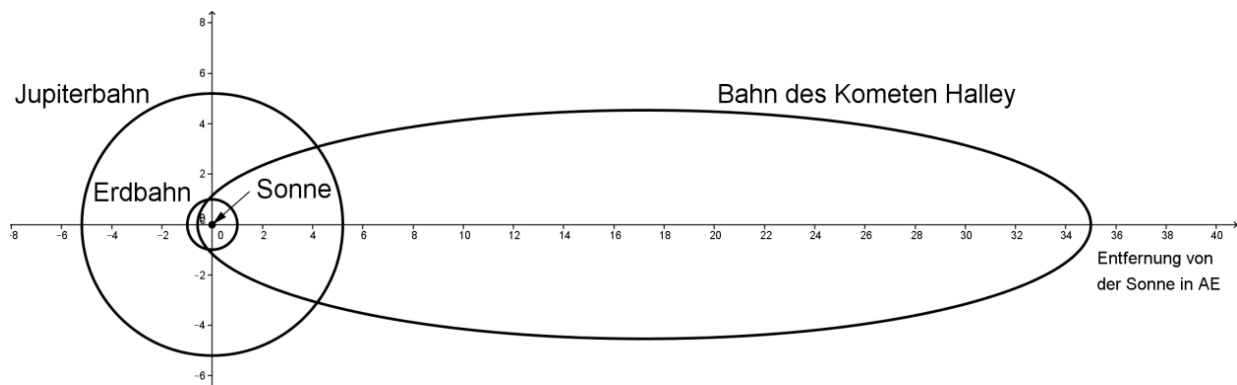


Abbildung 59: Die Bahnkurve des Halley'schen Kometen

Beispiel 2: Elliptische, hyperbolische und parabolische Spiegel

Dieses Beispiel ist der Optik entnommen und betrifft drei Arten speziell geformter Spiegel, den elliptischen Spiegel, den hyperbolischen Spiegel und den Parabolspiegel. Grundlage der besonderen Eigenschaften dieser drei Spiegel sind spezielle Eigenschaften der Kegelschnitte Ellipse, Hyperbel und Parabel.

Der Mathematikunterricht beschränkt sich beim Kapitel Kegelschnitte zumeist auf die algebraischen Methoden, dabei geht es dann darum, Gleichungen von Kegelschnitten aufzustellen, Geraden mit den Kegelschnitten zu schneiden und Tangenten zu bestimmen. Was dabei leider immer zu kurz kommt, sind die geometrischen Besonderheiten, die die Kegelschnitte aufweisen. Aber gerade diese speziellen Eigenschaften werden in vielen interessanten und alltäglichen Anwendungen in der Physik und Technik verwendet. Viele dieser Anwendungen sind direkte Folgerungen der geometrischen Eigenschaften der Kegelschnitte, weshalb es auch instruktiv sein kann, die Eigenschaften auf geometrischem statt auf algebraischem Weg zu zeigen.

Da auch in der Mathematik geometrische Beweise bedauerlicherweise immer mehr an Bedeutung verlieren, sollen die entsprechenden geometrischen Beweise gezeigt werden. Dies kann auch durchaus lehrreich für den Physikunterricht sein, da es für die Physik in einigen Bereichen sehr produktiv sein kann, wenn die SchülerInnen die Fähigkeit erlangen, sich bildliche Vorstellungen von physikalischen Gegebenheiten zu machen.

Elliptische Spiegel

Die zugrunde liegende geometrische Eigenschaft für die Konstruktion von elliptischen Spiegeln besagt folgendes:

Die Normale zur Tangente halbiert den Winkel zwischen den Leitstrahlen.

Das kann man mit folgendem geometrischen Beweis zeigen, der in den Grundzügen und in aller Kürze im Schulbuch Mathematik verstehen zu finden ist.¹³⁹

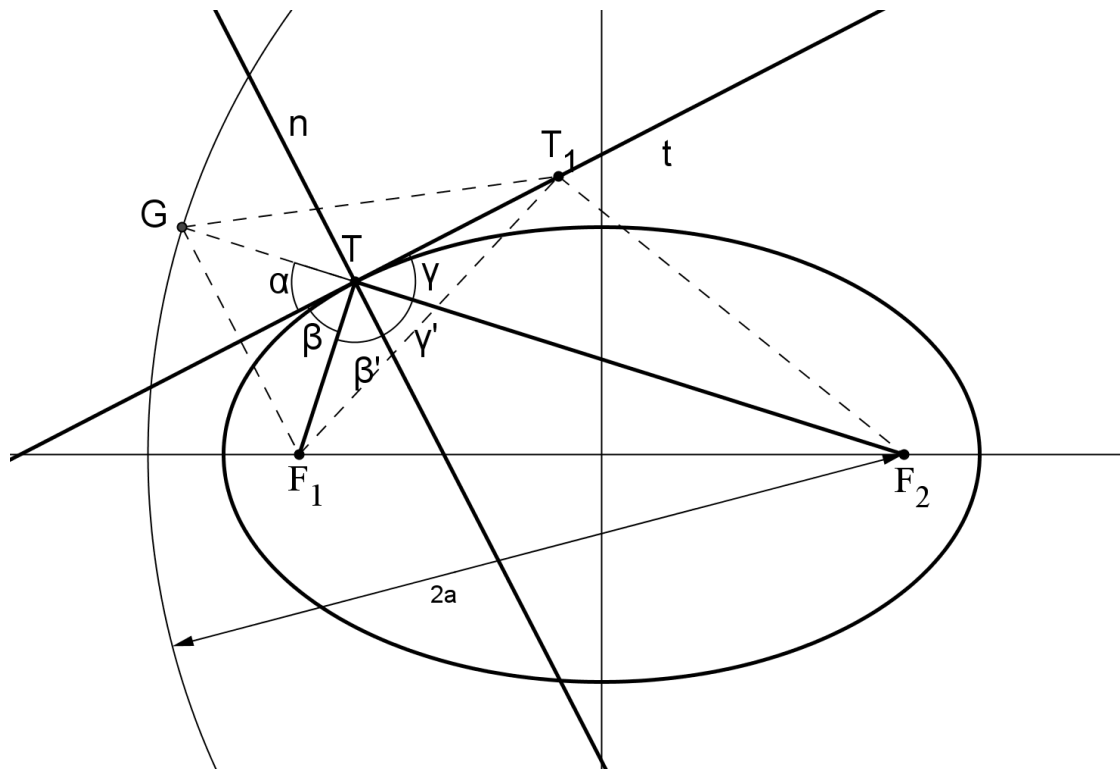


Abbildung 60: Ein geometrischer Beweis zur Ellipse

Man zeichne einen beliebigen Punkt T auf der Ellipse und zeichne einen Kreis mit dem Mittelpunkt F_2 und dem Radius $2a = F_1T + F_2T$. Man verlängere die Leitlinie TF_2 und ermittle den Schnittpunkt mit dem Kreis, dieser Punkt sei mit G bezeichnet.¹⁴⁰

Da $F_2T + TG = F_2T + TF_1 = 2a$ ist, muss $TG = TF_1$ gelten.

¹³⁹ Siehe [13], Seite 168; der Beweis wurde von mir als Beweisidee verwendet und abgeändert.

¹⁴⁰ Den Kreis bezeichnet man als 1. Leitkreis und den Punkt G als 1. Gegenpunkt von T .

Zeichnet man nun die Streckensymmetrale der Strecke F_1G , so muss diese durch den Ellipsenpunkt T gehen, da ja T von G und F_1 gleich weit entfernt ist.

Es bleibt noch zu zeigen, dass die Gerade t wirklich eine Tangente der Ellipse im Punkt T ist. Dazu nimmt man einen beliebigen Punkt T_1 auf der Geraden an. Man erkennt unmittelbar, dass $F_1T_1 + T_1F_2 = GT_1 + T_1F_2 > 2a$ sein muss. Das bedeutet nun aber, dass T_1 kein Ellipsenpunkt sein kann und infolgedessen T der einzige Punkt der Geraden t ist, sodass t eine Tangente der Ellipse ist.

Weiters folgt $\alpha = \beta$ (da t Streckensymmetrale) sowie $\alpha = \gamma$ (α und γ sind Scheitelwinkel) und hieraus $\beta = \gamma$. Da β' und γ' komplementäre Winkel zu β und γ sind, müssen auch diese gleich sein und es folgt daraus die behauptete Eigenschaft, dass die Normale zur Tangente den Winkel zwischen den Leitlinien halbiert, mit anderen Worten gleich der Winkelsymmetralen der Leitlinien ist.

Aus dieser mathematischen Eigenschaft der Ellipse kann nun eine für die Optik interessante Folgerung gemacht werden. Reflexion an einer gekrümmten Fläche bedeutet eine Reflexion an der Tangente dieser Fläche, da man sich die Ellipse in einem Punkt durch ein kleines Tangentenstück ersetzt denken kann. Die Normale auf die Tangente entspricht dann in der Optik dem Lot auf die Spiegelfläche. Wenn nun ein Lichtstrahl durch einen Brennpunkt auf einen Punkt der Ellipse fällt, wird dieser so reflektiert, dass der reflektierte Strahl durch den anderen Brennpunkt der Ellipse gehen muss, dies ist in der Abbildung 61 veranschaulicht.

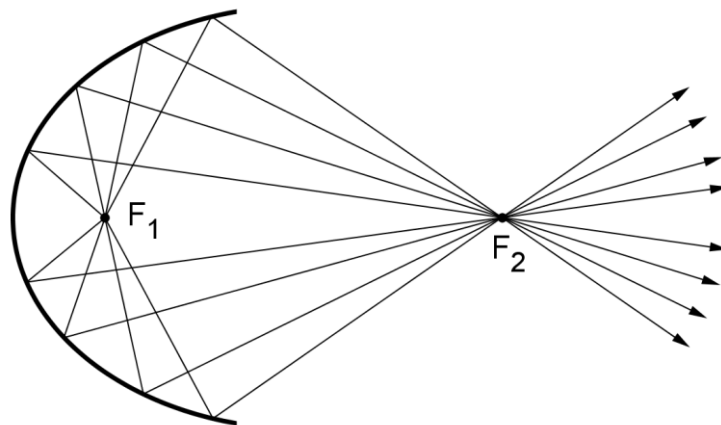


Abbildung 61: Ein Ellipsenspiegel

Anwendung findet solch ein Ellipsenspiegel etwa in Zahnarztpraxen. Die Lampe befindet sich im einen Brennpunkt des Spiegels und alle Lichtstrahlen werden im zweiten Brennpunkt vereinigt, der sich im Mund des Patienten befindet.

Hyperbelspiegel

Die zugrunde liegende geometrische Eigenschaft für die Konstruktion von hyperbolischen Spiegeln besagt folgendes:

Die Tangente halbiert den Winkel zwischen den Leitstrahlen.

Das kann man mit folgendem geometrischen Beweis zeigen, der in den Grundzügen ebenfalls im Schulbuch Mathematik verstehen 7 zu finden ist.¹⁴¹

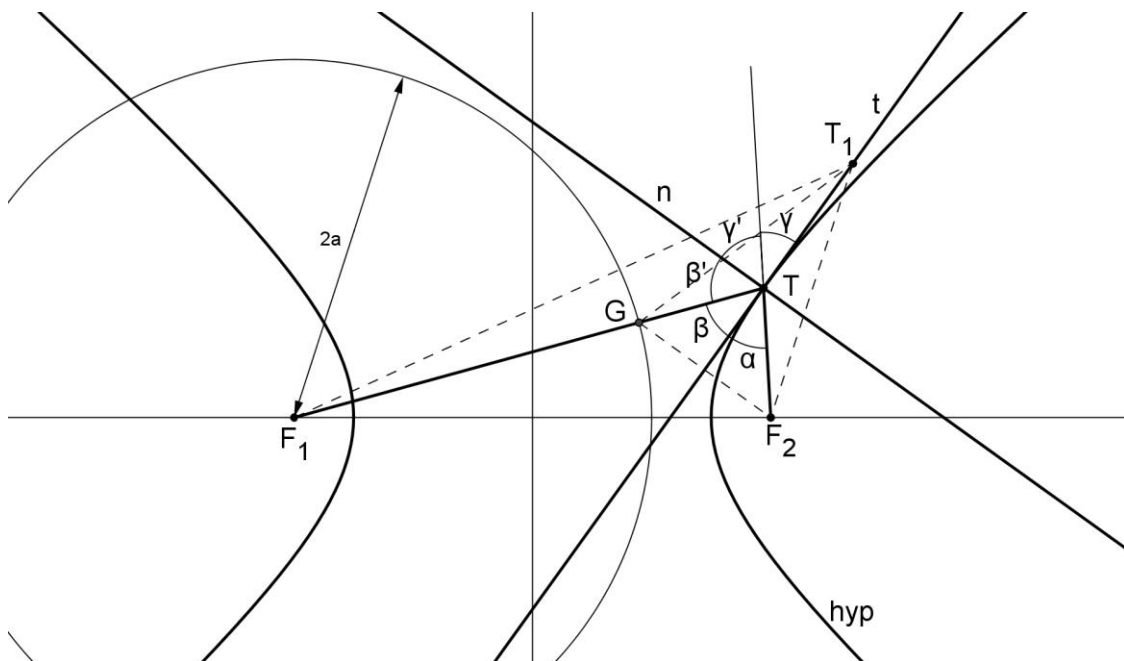


Abbildung 62: Ein geometrischer Beweis zur Hyperbel

¹⁴¹ Siehe [13], Seite 169; der Beweis wurde von mir als Beweisidee verwendet und abgeändert.

Man zeichne auf der Hyperbel einen beliebigen Punkt T und zeichne einen Kreis mit dem Mittelpunkt F_1 und dem Radius $2a = F_1T - F_2T$. Man schneide die Leitlinie TF_1 mit dem Kreis und ermittle den Schnittpunkt G .¹⁴²

Da $F_1T - TG = F_1T - TF_2 = 2a$ ist, muss $TG = TF_2$ gelten.

Zeichnet man nun die Streckensymmetrale der Strecke F_2G , so muss diese durch den Hyperbelpunkt T gehen, da ja T von G und F_2 gleich weit entfernt ist.

Es bleibt wiederum noch zu zeigen, dass die Gerade t wirklich eine Tangente im Punkt T ist. Dazu nimmt man einen beliebigen Punkt T_1 auf der Geraden an. Man erkennt unmittelbar, dass $F_1T_1 - T_1F_2 = F_1T_1 - T_1G < 2a$ sein muss. Das bedeutet nun aber, dass T_1 kein Hyperbelpunkt sein kann und infolgedessen T der einzige Punkt der Geraden t ist, sodass t eine Tangente der Hyperbel ist.

Weiters folgt $\alpha = \beta$, da t Streckensymmetrale von F_2G ist. Damit ist die Eigenschaft gezeigt, dass die Tangente den Winkel zwischen den Leitlinien halbiert, mit anderen Worten gleich der Winkelsymmetralen der Leitlinien ist.

Genauso wie beim Ellipsenspiegel kann man aus dieser mathematischen Eigenschaft der Hyperbel eine wichtige Folgerung für die Optik gewinnen. Aus der Abbildung 62 folgt weiter, dass $\alpha = \gamma$ (α und γ sind Scheitelwinkel) und daraus $\beta = \gamma$. Da β' und γ' komplementäre Winkel zu β und γ sind, müssen auch diese gleich sein. Die Normale auf die Tangente entspricht wieder dem Lot, da der Einfallswinkel und der reflektierte Winkel gleich sind und damit das Reflexionsgesetz erfüllt ist. Weiters ist aus der Abbildung 62 ersichtlich dass die Verlängerung des reflektierten Strahls durch den Brennpunkt F_2 geht.

So gesehen kann die Eigenschaft einer Hyperbel auch folgendermaßen formuliert werden:

¹⁴² Den Kreis bezeichnet man als 1. *Leitkreis* der Hyperbel und den Punkt G als 1. *Gegenpunkt* von T .

Ein von einem Brennpunkt einer Hyperbel ausgehender Lichtstrahl wird in einem Punkt der Hyperbel so reflektiert, als ob er vom anderen Brennpunkt kommen würde.¹⁴³

Dies ist in der Abbildung 63 veranschaulicht.

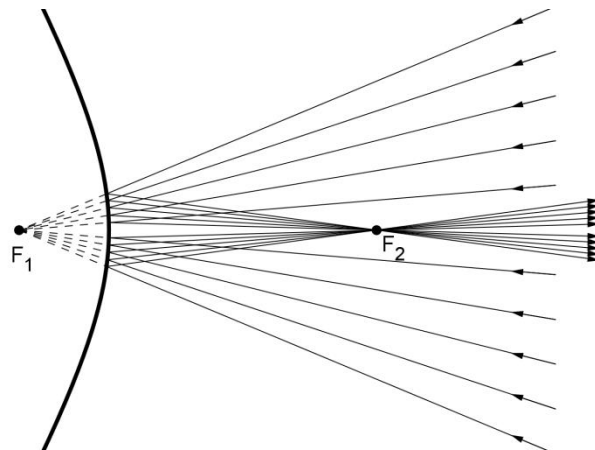


Abbildung 63: Ein Hyperbelspiegel

Parabolspiegel

Die zugrunde liegende geometrische Eigenschaft für die Konstruktion von Parabolspiegeln besagt folgendes:

Die Tangente halbiert den Winkel zwischen der Leitstrecke und der Strecke, die vom Punkt normal auf die Leitlinie steht.

Das kann man mit folgendem geometrischen Beweis zeigen, der in den Grundzügen im Schulbuch Mathematik verstehen 7 zu finden ist.¹⁴⁴

¹⁴³ Siehe [13] Seite 169

¹⁴⁴ Siehe [13], Seite 169; der Beweis wurde von mir als Beweisidee verwendet und abgeändert.

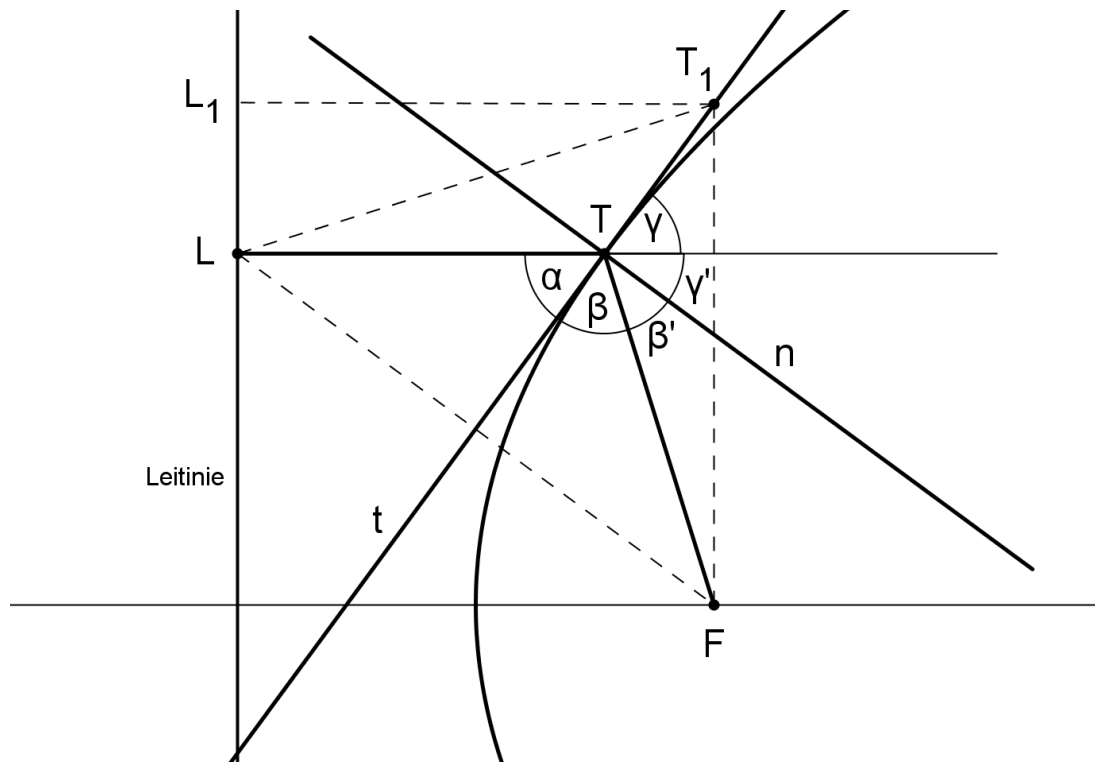


Abbildung 64: Ein geometrischer Beweis zur Parabel

Man zeichne auf der Parabel einen beliebigen Punkt T . Laut Definition einer Parabel ist die Leitstrecke TF und der (Normal)Abstand des Punktes T von der Leitlinie gleich lang. Also liegt T auf der Streckensymmetralen t von LF .

Es bleibt wiederum noch zu zeigen, dass die Gerade t wirklich eine Tangente im Punkt T ist. Dazu nimmt man einen beliebigen Punkt T_1 auf der Geraden an. Man erkennt unmittelbar, dass $L_1T_1 \neq LT_1 = T_1F$ ist. Das bedeutet aber, dass T_1 kein Parabelpunkt sein kann und infolgedessen T der einzige Punkt der Geraden t ist, sodass t eine Tangente der Parabel ist.

Weiters folgt $\alpha = \beta$, da t eine Streckensymmetrale von LF ist. Damit ist die Eigenschaft gezeigt, dass die Tangente den Winkel zwischen der Leitstrecke und der Strecke, die vom Punkt normal auf die Leitlinie steht halbiert.

Genauso wie beim Ellipsen- und Hyperbelspiegel kann man aus dieser mathematischen Eigenschaft der Parabel eine wichtige Folgerung für die Optik gewinnen. Aus der Abbildung 64 folgt weiter, dass $\alpha = \gamma$ (α und γ sind Scheitelwinkel) und daraus $\beta = \gamma$.

Da β' und γ' komplementäre Winkel zu β und γ sind, müssen auch diese gleich sein. Die Normale auf die Tangente entspricht wieder dem Lot, da der Einfallswinkel und der reflektierte Winkel gleich sind und damit das Reflexionsgesetz erfüllt ist. Weiters ist aus der Abbildung 64 ersichtlich, dass die Verlängerung der Strecke LT parallel zur Parabelachse verläuft.

So gesehen kann die Eigenschaft einer Parabel auch folgendermaßen formuliert werden:

Ein vom Brennpunkt einer Parabel ausgehender Lichtstrahl wird an der Parabel parallel zur Parabelachse reflektiert und umgekehrt werden zur Parabelachse parallele Lichtstrahlen durch den Brennpunkt reflektiert.¹⁴⁵

Dies ist in der Abbildung 65 veranschaulicht.

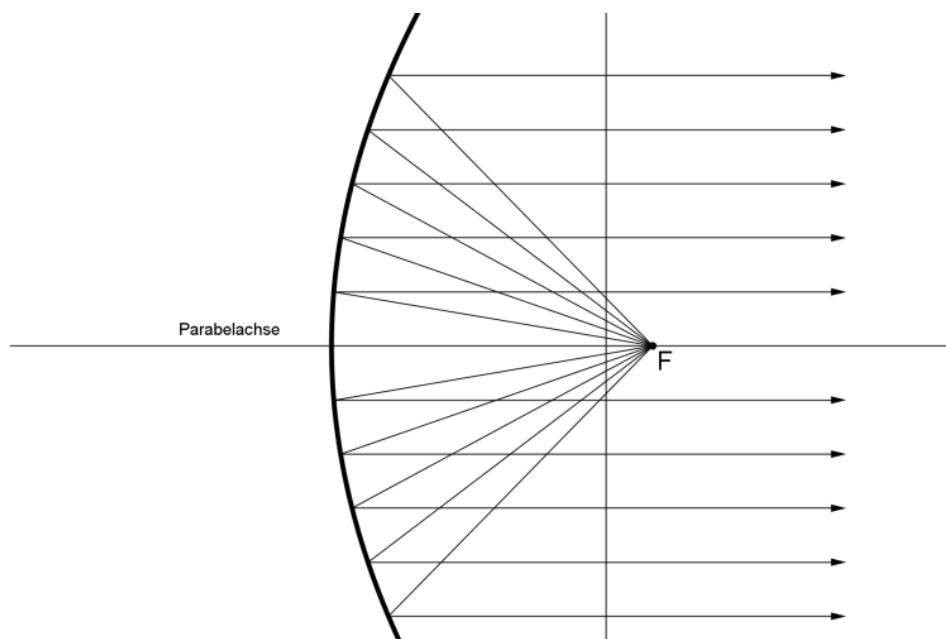


Abbildung 65: Ein Parabolspiegel

Sphärische Spiegel, also Spiegel, deren Spiegelfläche ein Teil einer Kugelschale ist, sind zwar einfacher und kostengünstiger herzustellen als Spiegel, deren Spiegelflächen keine sphärische Form haben, allerdings weisen sie die sogenannte sphärische Aberration auf. Achsenparallele Strahlen, die an einem sphärischen Spiegel reflektiert

¹⁴⁵ Siehe [13] Seite 170

werden, schneiden sich nicht im Brennpunkt des Spiegels, was natürlich die Bildqualität beungünstigt. Parabolische Spiegel weisen diese Aberration nicht auf, weshalb parallele Lichtstrahlen alle durch den Brennpunkt gehen beziehungsweise umgekehrt Lichtstrahlen, die vom Brennpunkt ausgehen, als achsenparallele Strahlen reflektiert werden.

Anwendungen der Parabolspiegel sind etwa Autoscheinwerfer, bei denen sich eine Lichtquelle im Brennpunkt des Parabolspiegels befindet und alle Strahlen beim Fernlicht als achsenparallele Strahlen reflektiert werden.

Es soll noch als Abschluss dieser Anwendungen der Kegelschnitte kurz das Cassegrain-Teleskop beschrieben werden. Der Aufbau ist in Abbildung 66 ersichtlich.

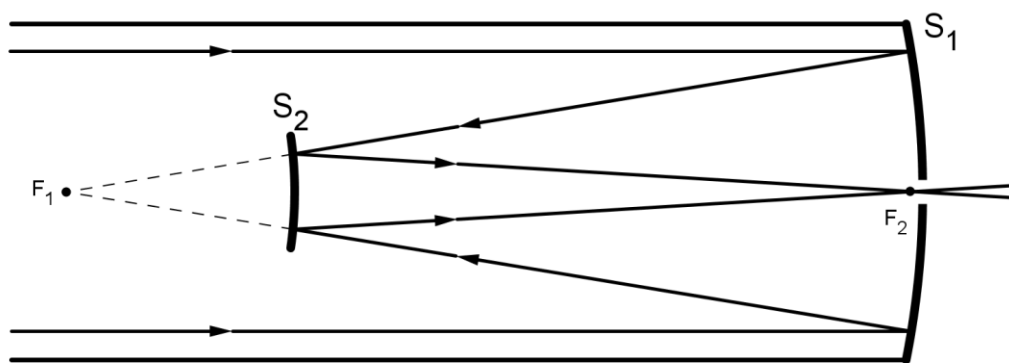


Abbildung 66: Das Cassegrain-Teleskop

Die von einem weit entfernten Objekt (Stern, Galaxien) eintreffenden Lichtstrahlen können als parallel angesehen werden und treffen auf den Hauptspiegel S₁. Dieser ist ein konkav-parabolischer Spiegel, der die Lichtstrahlen in den Brennpunkt F₁ reflektiert. Bevor die Lichtstrahlen den Brennpunkt F₁ erreichen, treffen sie auf den Sekundärspiegel S₂. Dieser ist ein konvex-hyperbolischer Spiegel, der so angeordnet ist, dass sein Brennpunkt mit dem Brennpunkt des Hauptspiegels zusammenfällt. Deswegen werden die Lichtstrahlen vom Sekundärspiegel in seinen zweiten Brennpunkt reflektiert, der beim Hauptspiegel liegt. Durch eine kleine Öffnung im Hauptspiegel verlassen die Lichtstrahlen das Teleskop, werden durch das Okular parallel ausgerichtet und können dann mit dem Auge oder einer Kamera betrachtet werden.

3.7.3 Komplexe Zahlen

Komplexe Zahlen werden im Mathematikunterricht in der siebten Klasse behandelt, allerdings fehlt es meist völlig an konkreten Anwendungen, die außerhalb der Mathematik liegen. Gerade bei LehrerInnen, die nicht Physik beziehungsweise Chemie unterrichten, fehlt das Hintergrundwissen und auch die Erfahrung der konkreten Anwendungsmöglichkeiten der komplexen Zahlen.

Das liegt sicherlich auch daran, dass die Anwendungsbeispiele auf dem Niveau der Schulmathematik recht rar sind, einziges auch praktisch anwendbares Beispiel hierzu sind Zeigerdiagramme aus der Wechselstromlehre. In der Hochschulphysik dagegen sind komplexe Zahlen beziehungsweise im Weiteren die komplexe Analysis als Erweiterung der reellen Analysis als wichtige Methoden, mathematische Probleme sehr effektiv lösen zu können, nicht wegzudenken und sozusagen alltägliches Handwerkszeug des Physikers.¹⁴⁶ Es sei auch darauf hingewiesen, dass komplexe Zahlen besonders in der Quantenmechanik eine überaus wichtige Rolle spielen, was aber natürlich über das Niveau der Schulmathematik weit hinausgeht.¹⁴⁷

Da wie schon erwähnt, die einzig wirklich wichtige Anwendung der komplexen Zahlen in der Wechselstromlehre zu finden ist, soll die Anwendung der komplexen Zahlen auch anhand von Zeigerdiagrammen bei Wechselstromkreisen ausführlicher erklärt werden. In neueren Lehrbüchern wird dieses Thema leider sehr stiefmütterlich behandelt und mit dieser Darstellung hoffe ich die Vorzüge herausstreichen zu können, da, sobald die komplexen Widerstände erklärt sind, mit diesen sehr weitreichende und interessante Anwendungen behandelt werden können.

¹⁴⁶ Als Beispiel sei etwa genannt die Verwendung komplexer Integrationsmethoden in der Elektrodynamik, wenn es bei beschleunigten Ladungen um die Berechnung retardierter Potentiale geht. Im Zuge der Berechnung sind Integrale zu berechnen, deren Integrationsweg über singuläre Punkte geht, was durch die Benutzung komplexer Integrale leicht zu bewerkstelligen ist. Man kann in der komplexen Zahlenebene einen Integrationsweg wählen, der um die problematischen singulären Punkte herumführt und das Integral mit Hilfe des Residuensatzes sehr einfach auswerten.

¹⁴⁷ So ist die Schrödingergleichung eine komplexe Differentialgleichung, wobei auch die Lösungen im Allgemeinen komplexe Funktionen sind. Auch die Pauli-Matrizen sind komplexwertige Matrizen, die zur Beschreibung des Spins etwa von Elektronen verwendet werden.

Beispiel 1: Zeigerdiagramme

In der Physik kommen in vielen Bereichen periodische Vorgänge vor, die natürlich mit Hilfe der trigonometrischen Funktionen dargestellt werden können. Eine alternative Methode ist die Darstellung dieser Funktionen durch einen rotierenden Zeiger als Vektor und im Weiteren dann als eine komplexe Zahl in der Gauß'schen Zahlenebene. Zur Anwendung kommt das Zeigermodell vor allem in der Schwingungslehre, der Wellenoptik, der Quantenmechanik und, was auch im Folgenden vorgestellt werden soll, in der Wechselstromlehre.

Eine Sinusschwingung, und nur solche kommen für des Zeigermodell in Frage, kann allgemein durch die Funktionsgleichung

$$y(t) = \hat{y} \cdot \sin(\omega t + \varphi_0)$$

angeschrieben werden.

Die Grundidee des Zeigermodells ist eine sehr einfache. Ein Zeiger (Vektor) drehe sich mit der konstanten Winkelgeschwindigkeit ω um den Ursprung, wobei seine Länge der Amplitude $\hat{y}$ des entsprechenden periodischen Vorgangs entspricht.

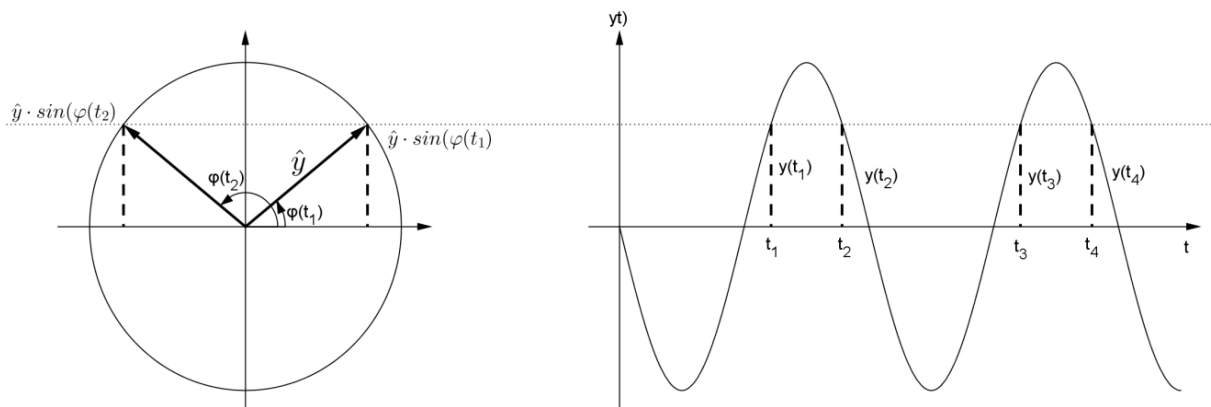


Abbildung 67: Sinusschwingung und Zeigerdiagramm

Aus der Abbildung 67 erkennt man unmittelbar, dass der y-Komponente des Zeigers die Auslenkung der periodischen Funktion entspricht, also lässt sich schreiben

$$\hat{y} \cdot \sin(\varphi(t)) = \hat{y} \cdot \sin(\omega t + \varphi_0).$$

Entsprechend kann natürlich jede periodische Bewegung auch alternativ durch die Kosinusfunktion beschrieben werden, also

$$y(t) = \hat{y} \cdot \cos(\omega t + \varphi_0) .$$

In diesem Fall entspricht die x-Komponente des Zeigers der periodischen Funktion, wenn gilt

$$\hat{y} \cdot \cos(\varphi(t)) = \hat{y} \cdot \cos(\omega t + \varphi_0) .$$

Beide Darstellungen sind natürlich äquivalent, da sich die Sinus- und Kosinusfunktion nur um einen Phasenwinkel von $\frac{\pi}{2}$ unterscheiden, es gilt also immer

$$\cos(x) = \sin\left(x + \frac{\pi}{2}\right) .$$

Es steht grundsätzlich frei, welche Darstellung, Sinus oder Kosinus, man wählt, entscheidend ist nur, dass man die gewählte Darstellung in der Folge beibehält.

Es gelten folgende Zuordnungen zwischen dem Zeigermodell und der Schwingungsdarstellung:

Zeigerlänge	↔	Schwingungsamplitude
y-Komponente des Zeigers	↔	Auslenkung der Schwingung
Momentanwinkel des Zeigers	↔	Phasenwinkel
Winkelgeschwindigkeit	↔	Kreisfrequenz
Drehzahl	↔	Frequenz
Umlaufdauer	↔	Periodendauer

In einem weiteren Schritt soll der Zeiger als komplexe Zahl aufgefasst werden, wie in der Abbildung 68 dargestellt ist.

Der Zeiger wird nun als komplexe Zahl $z = a + i \cdot b$ geschrieben

$$z(t) = \hat{y} \cdot (\cos(\varphi(t)) + i \cdot \sin(\varphi(t)))$$

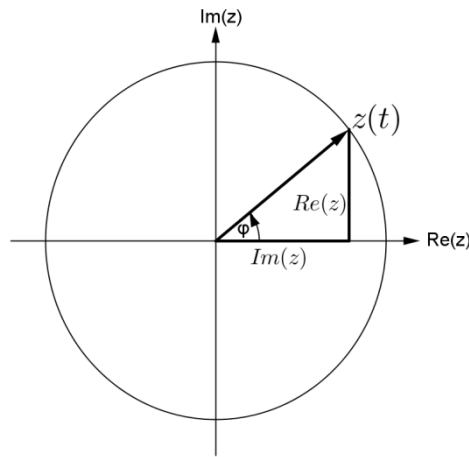


Abbildung 68: Komplexe Darstellung des Zeigermodells

Die Schwingung entspricht dabei dem Realteil oder dem Imaginärteil der komplexen Zahl

$$y(t) = \operatorname{Re}(z(t)) = \hat{y} \cdot \cos(\varphi(t)) = \hat{y} \cdot \cos(\omega t + \varphi_0)$$

oder falls als Darstellung der Sinus gewählt wird

$$y(t) = \operatorname{Im}(z(t)) = \hat{y} \cdot \sin(\varphi(t)) = \hat{y} \cdot \sin(\omega t + \varphi_0).$$

Aus dem Mathematikunterricht der siebten Klasse sollte die Euler'sche Formel für die Exponentialdarstellung einer komplexen Zahl bekannt sein, sie lautet

$$e^{i \cdot x} = \cos x + i \cdot \sin x$$

Für eine komplexe Zahl $z = a + i \cdot b$ folgt dann

$$\underline{\hat{y}} \cdot e^{i(\omega t + \varphi_0)} = \underline{\hat{y}} \cdot e^{i \cdot \varphi_0} \cdot e^{i \cdot \omega \cdot t} = \underline{\hat{y}} \cdot e^{i \cdot \omega \cdot t}.$$

$\underline{\hat{y}} = \hat{y} \cdot e^{i \cdot \varphi_0}$ ist dabei die komplexe zeitunabhängige Amplitude, die den Drehwinkel zur Zeit $t = 0$ angibt und die Länge $\hat{y}$ hat. $e^{i \cdot \omega \cdot t}$ stellt einen Zeiger der Länge 1 dar¹⁴⁸, der die zeitliche Abhängigkeit des periodischen Vorganges beinhaltet.

¹⁴⁸ Dass eine komplexe Zahl der Form $e^{i \cdot \omega \cdot t}$ immer die Länge 1 hat, erkennt man leicht:

$$|e^{i \cdot \omega \cdot t}| = |\cos \omega t + i \cdot \sin \omega t| = \sqrt{\cos^2 \omega t + \sin^2 \omega t} = 1$$

Beispiel 2: Zeigerdiagramme bei Widerständen, Spulen und Kondensatoren

In diesem Beispiel soll gezeigt werden, wie man die Beziehungen zwischen der Spannung und der Stromstärke in einem Wechselstromkreis durch die Verwendung von Zeigern beziehungsweise in komplexer Darstellung einfach geometrisch darstellen kann. Es sollen die Spannungsabfälle an den drei Bauteilen Widerstand, Spule und Kondensator in einem Wechselstromkreis ermittelt werden, wobei die Bauteile in diesem Beispiel noch einzeln betrachtet werden.

Es soll angenommen werden, dass in allen drei Fällen der Spannungsabfall an den Bauteilen so geartet ist, dass die Spannung $U(t)$ einen Strom $I(t) = \hat{I} \cdot \sin(\omega \cdot t + \varphi_i)$ beziehungsweise $I(t) = \hat{I} \cdot e^{i(\omega t + \varphi_i)}$ zur Folge hat, wobei $\hat{I}$ die (reelle) Amplitude und φ_i der Nullphasenwinkel des Stromes ist. Es ist also jeweils die Spannung zu bestimmen, die diesen Strom verursacht, wobei die Kreisfrequenz unverändert bleibt, allerdings der Phasenwinkel sich in bestimmter Weise ändern wird.

Es soll gezeigt werden, wie im Folgenden die Berechnungen in komplexer Schreibweise durchgeführt werden, wobei komplexe Größen wie in der Wechselstromtechnik üblich mit einem Unterstrich als solche gekennzeichnet werden.

Seien der Strom und die Spannung im Allgemeinen wie folgt angeschrieben:

$$\underline{I}(t) = \hat{I} \cdot e^{i(\omega t + \varphi_i)} \quad \text{und} \quad \underline{U}(t) = \hat{U} \cdot e^{i(\omega t + \varphi_u)}$$

Analog zum Ohm'schen Gesetz in der Gleichstromtechnik, also $U = R \cdot I$ beziehungsweise $\frac{U}{I} = R$, wird in der Wechselstromtechnik der Quotient von (komplexer) Spannung und (komplexer) Stromstärke als (komplexer) Widerstand oder *Impedanz* $\underline{Z}$ bezeichnet. Die Impedanz ist im Allgemeinen komplexwertig und kann wie jede komplexe Zahl in Real- und Imaginärteil aufgespalten werden:

$$\underline{Z} = \operatorname{Re} \underline{Z} + i \cdot \operatorname{Im} \underline{Z}$$

Der Realteil der Impedanz wird als *Wirkwiderstand*¹⁴⁹ (*Resistanz*) R , der Imaginärteil als *Blindwiderstand*¹⁵⁰ (*Reaktanz*) X bezeichnet.

¹⁴⁹ Im Wirkwiderstand wird elektrische Energie unumkehrbar in thermische, mechanische oder chemische Energie umgewandelt.

Der *Scheinwiderstand* Z , ergibt sich dann nach dem pythagoreischen Lehrsatz als

$$Z = \sqrt{R^2 + X^2}.$$

A) Widerstand in einem Wechselstromkreis

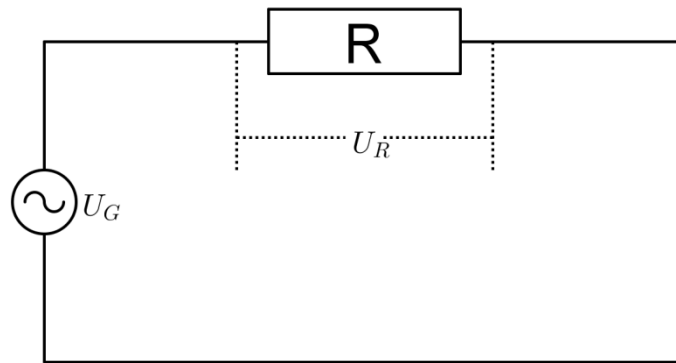


Abbildung 69: Ein Widerstand in einem Wechselstromkreis

Es soll eine Situation wie in Abbildung 69 vorliegen, dass eine ideale Wechselspannungsquelle¹⁵¹ an einen Widerstand R angeschlossen ist. Nach der 1. Kirchhoff'schen Regel¹⁵² gilt

$$\underline{U}_G - \underline{U}_R = 0$$

Laut dem Ohm'schen Gesetz gilt für den Spannungsabfall am Widerstand

$$\underline{U}_R = \underline{I} \cdot R$$

sodass insgesamt folgt

$$\boxed{\underline{U}_R(t) = R \cdot \underline{I}(t)}$$

Also muss gelten

¹⁵⁰ Im Blindwiderstand wird keine elektrische Energie in thermische, mechanische oder chemische Energie umgewandelt, sie pendelt vielmehr zwischen Erzeuger und Verbraucher hin und her, das können zum Beispiel auch ein Kondensator und eine Spule in einem Wechselstromkreis sein.

¹⁵¹ Bei einer idealen Wechselspannungsquelle können der innerer Widerstand, die Selbstinduktivität und die Kapazität der Spannungsquelle vernachlässigt werden.

¹⁵² Auch Maschensatz genannt, wonach sich in einem geschlossenen Stromkreis die Summe der elektrischen Spannungen zu Null addieren muss. Die Kirchhoff'schen Regeln gelten in einem Wechselstromkreis genauso wie in einem Gleichstromkreis, weshalb sie als Berechnungsgrundlage herangezogen werden können.

$$\hat{U} \cdot e^{i(\omega t + \varphi_u)} = R \cdot \hat{I} \cdot e^{i(\omega t + \varphi_i)}$$

$$R = \frac{\hat{U}}{\hat{I}} \cdot e^{i(\varphi_u - \varphi_i)} = \frac{\hat{U}}{\hat{I}} \cdot (\cos(\varphi_u - \varphi_i) + i \cdot \sin(\varphi_u - \varphi_i))$$

Da R eine reelle Größe ist, muss der Imaginärteil des vorigen Ausdrucks verschwinden, also gelten

$$\sin(\varphi_u - \varphi_i) = 0 \Leftrightarrow \varphi_u = \varphi_i.$$

Das heißt, dass bei einem ohmschen Widerstand in einem Wechselstromkreis die Nullphasenwinkel immer gleich sind, was bedeutet, dass es zu keiner Phasenverschiebung zwischen dem Strom und der Spannung kommt, man spricht hier von Gleichphasigkeit.

Für die Impedanz folgt dann

$$\underline{Z} = \frac{\underline{U}}{\underline{I}} = \frac{\hat{U}}{\hat{I}} = R$$

Die Verhältnisse zwischen Strom und Spannung sind in Abbildung 70 veranschaulicht, wobei sowohl die Schwingungsdarstellung als auch die Darstellung in der komplexen Ebene gezeigt ist.

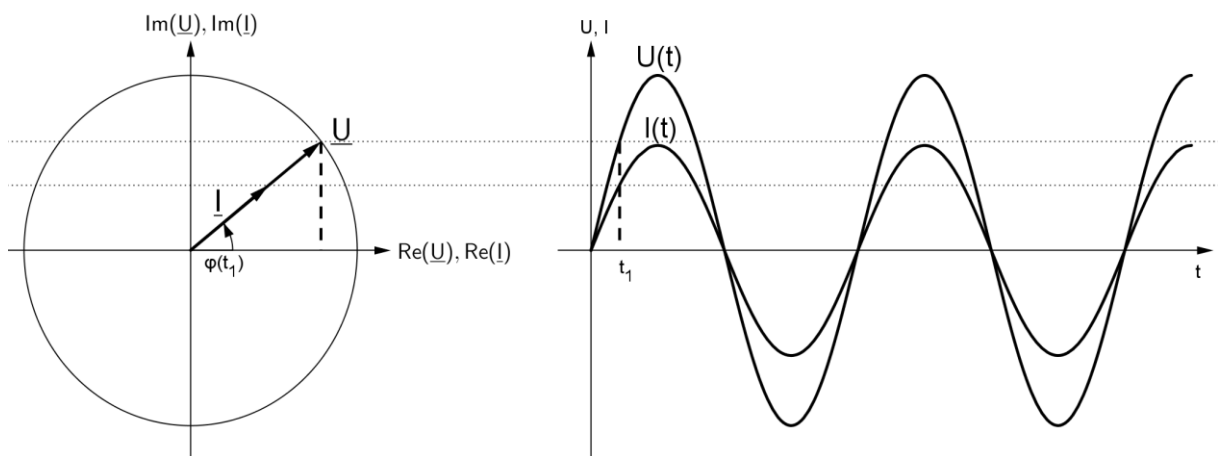


Abbildung 70: Phasenbeziehung beim Widerstand in einem Wechselstromkreis

B) Spule in einem Wechselstromkreis

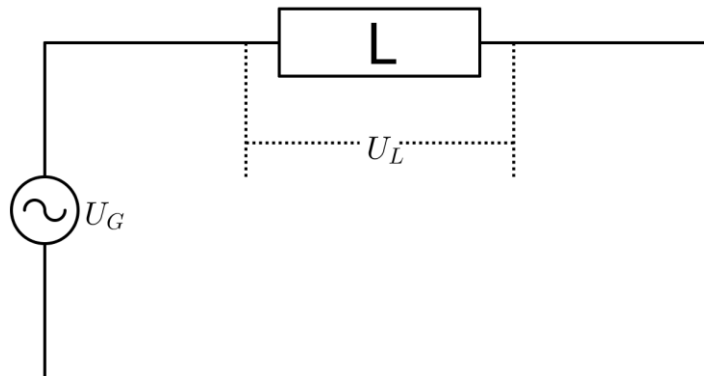


Abbildung 71: Eine Spule in einem Wechselstromkreis

Es soll jetzt eine Situation wie in Abbildung 71 vorliegen, dass eine Wechselspannungsquelle an eine Spule L angeschlossen ist.

Nach der 1. Kirchhoff'schen Regel gilt

$$\underline{U}_G - \underline{U}_L = 0$$

Für den Spannungsabfall an einer Spule gilt

$$\underline{U}_L = L \cdot \frac{d\underline{I}}{dt},$$

wobei L die Induktivität der Spule gemessen in Henry angibt. Sodann ergibt sich

$$\underline{U}_G = L \cdot \frac{d\underline{I}}{dt}$$

Für die Ableitung des Stromes nach der Zeit ergibt sich

$$\frac{d\underline{I}}{dt} = \frac{d}{dt} \left(\hat{I} \cdot e^{i(\omega t + \varphi_i)} \right) = i \cdot \omega \cdot \hat{I} \cdot e^{i(\omega t + \varphi_i)} = i \cdot \omega \cdot \underline{I}$$

Man erhält als Beziehung zwischen der Spannung an der Induktivität und dem Strom

$$\boxed{\underline{U}_L(t) = i \cdot \omega L \cdot \underline{I}(t)}$$

Die Multiplikation einer komplexen Größe mit der imaginären Einheit bedeutet, dass die Größe um den Winkel $\frac{\pi}{2}$ gegen den Uhrzeigersinn gedreht wird. Das erkennt man sehr leicht aus der Euler'schen Formel, da

$$e^{i \cdot \frac{\pi}{2}} = \cos \frac{\pi}{2} + i \cdot \sin \frac{\pi}{2} = i \quad \text{gilt.}$$

Mit anderen Worten bedeutet das, dass der Strom um $\frac{\pi}{2}$ hinter der Spannung zurückbleibt beziehungsweise die Spannung dem Strom um $\frac{\pi}{2}$ voreilt, das entspricht einer Viertelschwingung.

Für die Impedanz folgt dann

$$\underline{Z}_L = i \cdot \omega L$$

Der induktive Blindwiderstand ergibt sich als $X_L = \omega L$ und man erkennt dabei, dass dieser mit steigender Kreisfrequenz zunimmt. Der komplexe Widerstand wird auf der positiven imaginären Achse aufgetragen und steht damit normal auf den Wirkwiderstand. Er ist ein konstanter komplexer Zeiger, der allerdings nicht rotiert, da er keine Zeitabhängigkeit aufweist.

Die Verhältnisse hierzu sind in Abbildung 72 veranschaulicht.

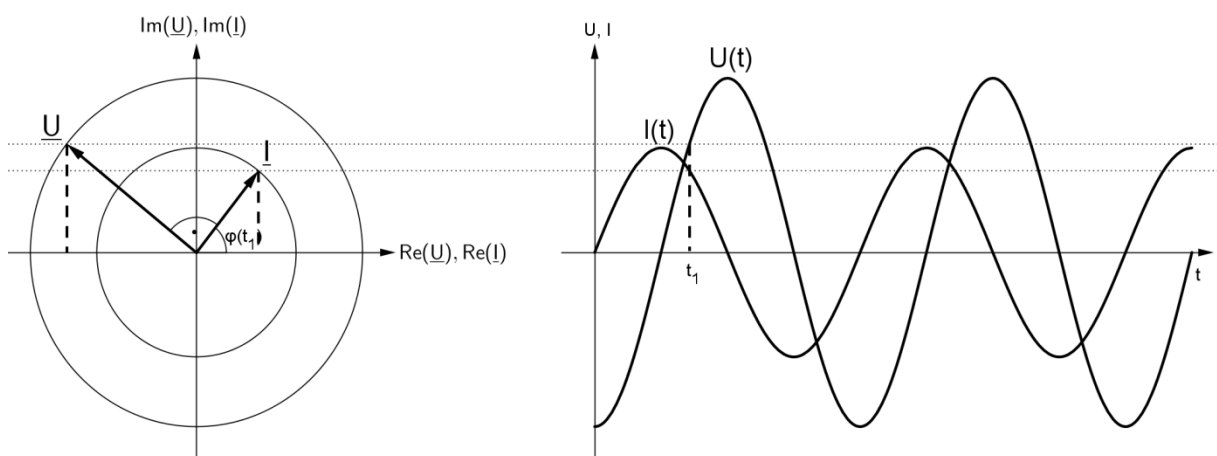


Abbildung 72: Phasenbeziehung an einer Spule in einem Wechselstromkreis

C) Kondensator in einem Wechselstromkreis

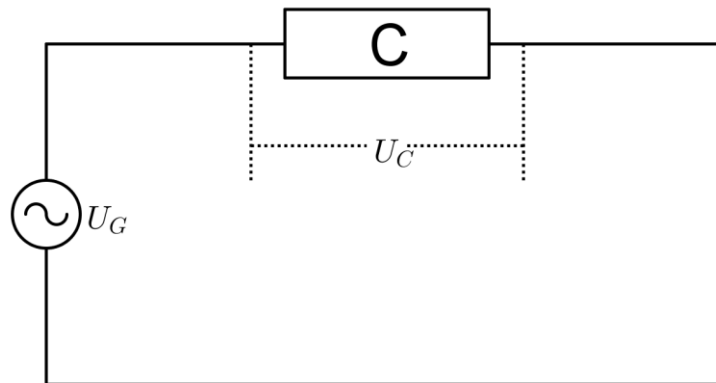


Abbildung 73: Ein Kondensator in einem Wechselstromkreis

Es soll jetzt die Situation wie in Abbildung 73 vorliegen, dass eine Wechselspannungsquelle an einen Kondensator C angeschlossen ist.

Nach der 1. Kirchhoff'schen Regel gilt

$$\underline{U}_G - \underline{U}_C = 0$$

Für den Spannungsabfall an einem Kondensator gilt

$$\underline{I}(t) = C \cdot \frac{d\underline{U}_C}{dt} ,$$

wobei C die Kapazität des Kondensators gemessen in Farad angibt.

Daraus folgt

$$\underline{I}(t) = C \cdot \frac{d}{dt} \left(\hat{U} \cdot e^{i(\omega t + \varphi_u)} \right) = i \cdot \omega C \cdot \hat{U} \cdot e^{i(\omega t + \varphi_u)} = i \cdot \omega C \cdot \underline{U}_C(t)$$

$$\underline{U}_C(t) = \frac{1}{i \cdot \omega C} \cdot \underline{I}(t) = -\frac{i}{\omega C} \cdot \underline{I}(t)$$

Man erhält als Beziehung zwischen der Spannung am Kondensator und dem Strom

$$\underline{U}_C(t) = -i \cdot \frac{1}{\omega C} \cdot \underline{I}(t)$$

Die Multiplikation einer komplexen Größe mit der negativen imaginären Einheit bedeutet, dass die Größe um den Winkel $\frac{\pi}{2}$ im Uhrzeigersinn gedreht wird. Das erkennt man sehr leicht aus der Euler'schen Formel, da

$$e^{i\left(-\frac{\pi}{2}\right)} = \cos\left(-\frac{\pi}{2}\right) + i \cdot \sin\left(-\frac{\pi}{2}\right) = -i \quad \text{gilt.}$$

Mit anderen Worten bedeutet das, dass der Strom um $\frac{\pi}{2}$ der Spannung voreilt beziehungsweise die Spannung um $\frac{\pi}{2}$ hinter dem Strom zurückbleibt. Für die Impedanz folgt dann

$$\underline{Z}_C = -\frac{i}{\omega C}$$

Der kapazitive Blindwiderstand ergibt sich als $X_C = -\frac{1}{\omega C}$ und man erkennt dabei,

dass dieser mit steigender Kreisfrequenz abnimmt. Der komplexe Widerstand wird auf der negativen imaginären Achse aufgetragen und steht damit wieder normal auf den Wirkwiderstand. Er ist ein konstanter komplexer Zeiger, der allerdings wiederum nicht rotiert, da er keine Zeitabhängigkeit aufweist. Die Verhältnisse hierzu sind in Abbildung 74 veranschaulicht.

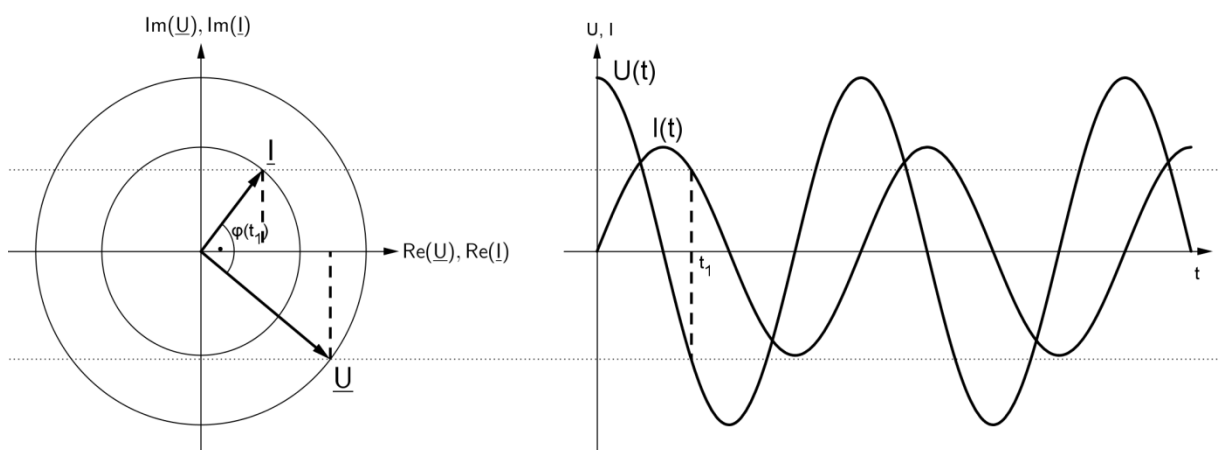


Abbildung 74: Phasenbeziehung an einem Kondensator in einem Wechselstromkreis

Es seien abschließend in Abbildung 75 die komplexen Widerstände (Impedanzen) eines ohmschen Widerstandes, einer Spule und eines Kondensators im Wechselstromkreis in der komplexen Ebene zusammenfassend dargestellt.

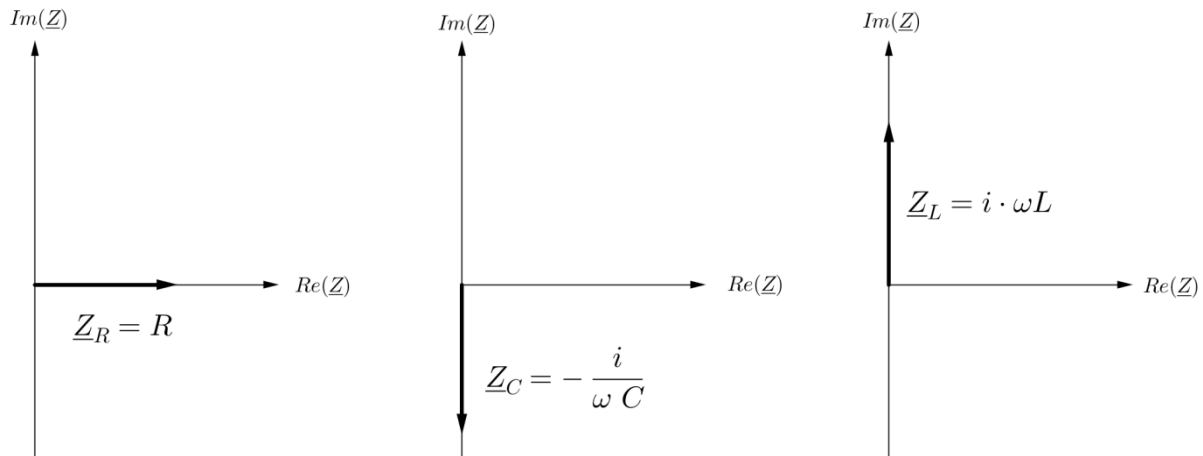


Abbildung 75: Komplexe Widerstände am ohmschen Widerstand, einer Spule und einer Kapazität

Beispiel 3: Schaltkreise aus mehreren Bauteilen

Im Beispiel 1 und Beispiel 2 wurde die Darstellung der Spannung und des Stromes sowie des komplexen Widerstandes (Impedanz) im komplexen Zeigerdiagramm vorgestellt, aber darin ist der eigentliche Nutzen dieser Darstellung noch nicht zum Ausdruck gekommen. Das soll nun mit diesem Beispiel geschehen. Sind einmal die komplexen Widerstände im Unterricht erklärt, ist die weitere Vorgangsweise sehr einfach und es können sowohl Serien- wie Parallelschwingkreise, Hoch- und Tiefpässe, sowie weiterführend etwa die Verhältnisse beim Drehstrom sehr leicht und ohne viele mathematische Ableitungen fortgeführt werden.

Es seien jetzt in einem Wechselstromkreis je zwei Bauteile in Reihe geschaltet, ein Widerstand und ein Kondensator (Abbildung 76a), ein Widerstand und eine Spule (Abbildung 76b) und eine Spule und ein Kondensator (Abbildung 76c).

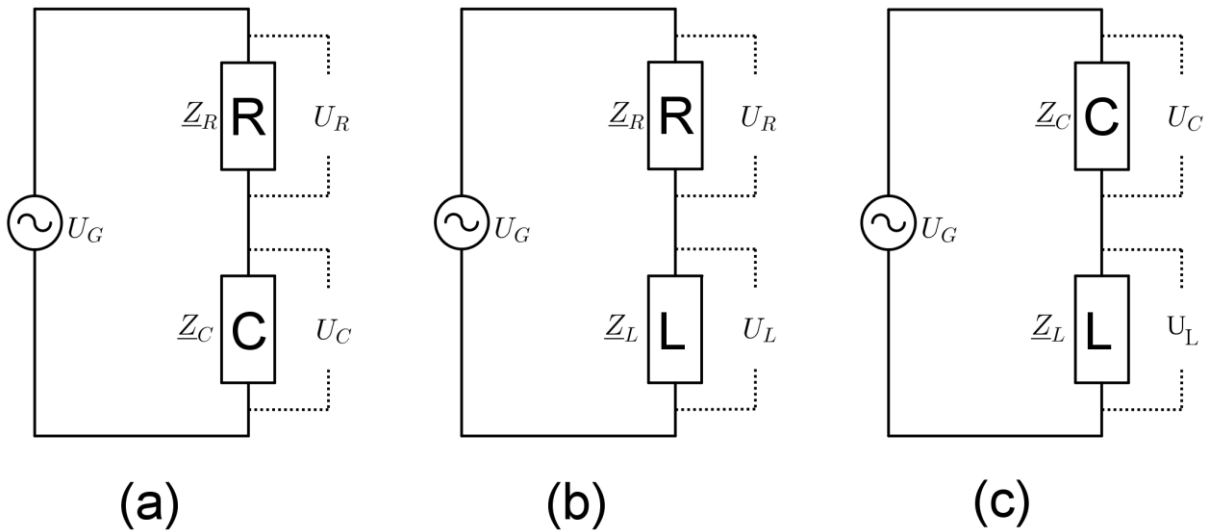


Abbildung 76: Wechselstromkreise mit zwei Bauteilen in Serie

In einem Wechselstromkreis gelten genauso die Regeln für Serien- oder Parallelschaltungen sowie die Kirchhoff'schen Regeln wie in einem Gleichstromkreis. Sind die Bauteile in Serie geschaltet, gibt man am einfachsten den Strom vor und bestimmt von diesem ausgehend die an jedem Bauteil anfallenden Spannungen, wie das auch im vorigen Beispiel 1 geschehen ist. Die (komplexen) Gesamtwiderstände ergeben sich dann sehr einfach durch Addition der (komplexen) Einzelwiderstände, es gilt also

$$\begin{aligned}\underline{Z}_{\text{RC}} &= \underline{Z}_{\text{R}} + \underline{Z}_{\text{C}} = R - \frac{i}{\omega C} \\ \underline{Z}_{\text{RL}} &= \underline{Z}_{\text{R}} + \underline{Z}_{\text{L}} = R + i \cdot \omega L \\ \underline{Z}_{\text{LC}} &= \underline{Z}_{\text{L}} + \underline{Z}_{\text{C}} = i \cdot \left(\omega L - \frac{1}{\omega C} \right)\end{aligned}$$

Die Beträge der Impedanzen ergeben sich durch übliche Betragsbildung einer komplexen Zahl als $|z| = \sqrt{z \cdot \bar{z}}$, wobei $\bar{z}$ die konjugiert komplexe Zahl zu z ist.

$$\begin{aligned}|\underline{Z}_{\text{RC}}| &= \sqrt{\left(R - \frac{i}{\omega C}\right) \cdot \left(R + \frac{i}{\omega C}\right)} = \sqrt{R^2 + \frac{1}{(\omega C)^2}} \\ |\underline{Z}_{\text{RL}}| &= \sqrt{(R + i \cdot \omega L) \cdot (R - i \cdot \omega L)} = \sqrt{R^2 + (\omega L)^2} \\ |\underline{Z}_{\text{LC}}| &= \left| \omega L - \frac{1}{\omega C} \right|\end{aligned}$$

Die Zusammenhänge für die komplexen Widerstände (Impedanzen) sind in Abbildung 77 veranschaulicht.

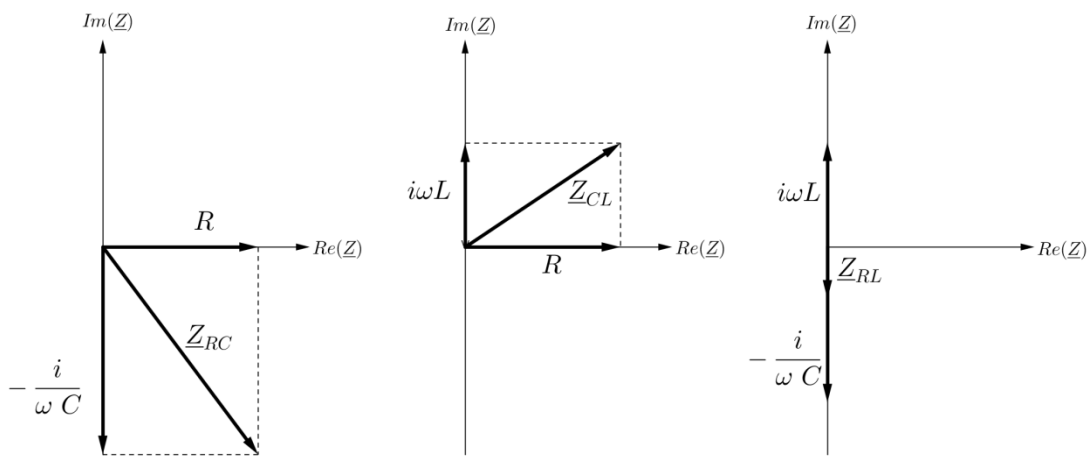


Abbildung 77: Impedanzen im Serienwechselstromkreis mit zwei Bauteilen

Interessant hierbei und wichtig in Anwendungen ist, dass sich die induktiven und kapazitiven Blindwiderstände gegenseitig zu Null addieren können, da sie ungleiches Vorzeichen haben und beide auf der imaginären Achse liegen. Der ohmsche Widerstand dagegen kann nicht zum Verschwinden gebracht werden, es ist also immer ein gewisser Energieverlust vorhanden.

Beispiel 4: Komplexe Widerstände beim Serienschwingkreis

Sind ein ohmscher Widerstand, eine Kapazität und eine Spule in Reihe geschaltet, so spricht man von einem Serienschwingkreis (Abbildung 78).

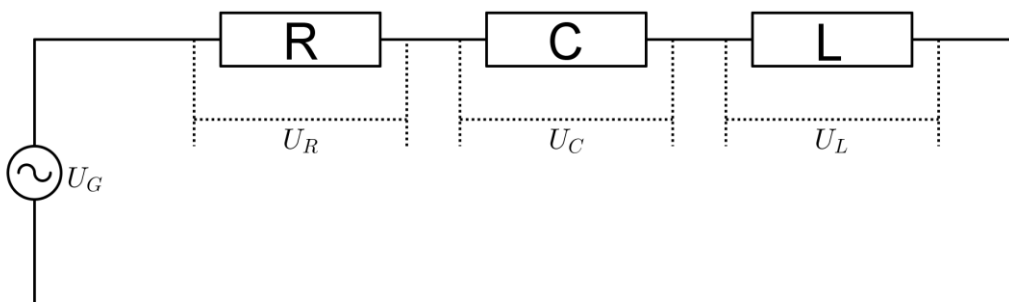


Abbildung 78: Ein Serienschwingkreis

Der komplexe Gesamtwiderstand (Impedanz) ergibt sich sofort aus der Summe der komplexen Einzelwiderstände zu

$$\underline{Z}_{\text{ges}} = \underline{Z}_R + \underline{Z}_C + \underline{Z}_L = R + i \cdot \left(\omega L - \frac{1}{\omega C} \right)$$

mit dem Betrag

$$|\underline{Z}_{\text{ges}}| = \sqrt{\left(R + i \cdot \left(\omega L - \frac{1}{\omega C} \right) \right) \cdot \left(R - i \cdot \left(\omega L - \frac{1}{\omega C} \right) \right)} = \sqrt{R^2 + \left(\omega L - \frac{1}{\omega C} \right)^2}$$

Man erkennt, dass der Betrag der Gesamtimpedanz beim Serienschwingkreis ein Minimum hat, wenn der zweite Summand verschwindet, wenn also gilt

$$\omega L - \frac{1}{\omega C} = 0 \Leftrightarrow \omega^2 LC = 0 \Leftrightarrow \omega = \frac{1}{\sqrt{LC}}$$

Das bedeutet, dass der komplexe Widerstand eines Serienschwingkreises mit der Kapazität C und der Induktivität L bei der Schwingungsfrequenz $\omega = (\sqrt{LC})^{-1}$ ein Minimum hat, man spricht von der *Resonanzfrequenz* des Serienschwingkreises.

Es seien noch die Zusammenhänge für die komplexen Widerstände (Impedanzen) des Serienschwingkreises in Abbildung 79 veranschaulicht.

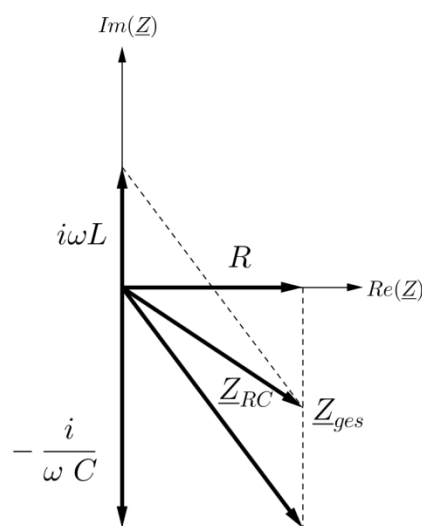


Abbildung 79: Die komplexen Widerstände beim Serienschwingkreis

3.8 8. Klasse

3.8.1 Integralrechnung

Die Integration ist die Umkehroperation der Differentiation. Das bedeutet, wird eine Funktion $f(x)$ integriert, so sucht man eine Funktion $F(x)$, deren Ableitung nach der Variablen x die Ausgangsfunktion $f(x)$ ist. Die Funktion $F(x)$ ist nicht eindeutig festgelegt, man darf immer eine konstante Funktion addieren, da die Ableitung der konstanten Funktion null ergibt.

In der Physik bedeutet dies, dass man immer Anfangs- oder Randbedingungen mit angeben muss, wenn man ein eindeutiges Ergebnis erhalten möchte. Man erhält durch die Integration als Ergebnis eine Funktionenschar und durch die Anfangs- oder Randbedingungen legt man fest, welche Funktion aus dieser Funktionenschar eine Lösung der gegebenen Aufgabe darstellt. Möchte man etwa die Bewegungsgleichung durch Integration lösen, muss die Anfangsgeschwindigkeit und die Anfangsposition zu einer bestimmten Zeit (meist zur Zeit $t = 0$) angegeben werden.

Die Dimension eines Integrals einer Funktion erhält man durch die Multiplikation der Dimension des Integranden mit der Dimension der Integrationsvariablen. Wird zum Beispiel das Integral $\int_0^1 v(t) dt$ gebildet, wobei $v(t)$ eine Geschwindigkeitsfunktion darstellen soll, so ist die Dimension des Integranden $[v] = m \cdot s^{-1}$ und die Dimension der Integrationsvariablen $[t] = s$. Für die Dimension des Integrals ergibt sich dann $(m \cdot s^{-1}) \cdot s = m$, also die Dimension einer Länge.

Im Folgenden sollen mehrere Beispiele vorgestellt werden, wie in der Physik die Integration zur Berechnung diverser Größen verwendet werden kann. Es soll nicht unerwähnt bleiben, dass die Integration in der Physik sicherlich eine der häufigsten und grundlegenden Rechenoperationen darstellt und in jedem Teilgebiet der Physik zur Anwendung kommt.

Beispiel 1: Arbeit in einem Kraftfeld

Das Beispiel ist dem Schulbuch Mathematik 8 von Götz, Reichel¹⁵³ entnommen.

- 412 Berechne die Arbeit (in Joule), die notwendig ist, um einen Körper mit der Masse 1 kg von der Erdoberfläche in 300 km Höhe zu bringen! ($M_{\text{Erde}} = 6 \cdot 10^{24} \text{ kg}$, $r_{\text{Erde}} = 6370 \text{ km}$).
- 413 Berechne die Arbeit (in Joule), die notwendig ist, um einen Körper mit der Masse 1 kg von der Erdoberfläche aus dem Gravitationsfeld der Erde zu entfernen!

Dieses erste Beispiel soll etwas genauer behandelt werden, um die Vorgangsweise der Integration, wie sie der Physiker in der Praxis durchführt, besser nachvollziehen zu können.

Auf einen Körper wirkt die Gravitationskraft $F(r) = G \cdot \frac{M_{\text{Erde}} \cdot m}{r^2}$ ¹⁵⁴, wobei $M_{\text{Erde}} = 6 \cdot 10^{24} \text{ kg}$ die Masse der Erde und m die Masse des im Gravitationsfeld der Erde bewegten Körpers ist. $G = 6,673 \cdot 10^{-11} \text{ m}^3 \text{ kg}^{-1} \text{ s}^{-2}$ ist die Gravitationskonstante. Da die Gravitationskraft eine Funktion der Entfernung vom Erdmittelpunkt ist, wird sie für die kleine Wegstrecke Δs als konstant angenommen. Für diese kleine Wegstrecke gilt dann angenähert

$$\Delta W = F(s) \cdot \Delta s$$

Wird die gesamte Wegstrecke in kleine Wegdifferenzen Δs_i aufgeteilt und alle Produkte $F(s_i) \cdot \Delta s_i$ aufsummiert, ergibt sich ein näherungsweiser Ausdruck für die Arbeit, die verrichtet wurde.

¹⁵³ Siehe [27] Aufgabe 412 Seite 110

¹⁵⁴ Die Kraft kann hier als skalare Größe angeschrieben werden, da nur Bewegungen parallel zur Kraft eine Rolle spielen. Wird ein Körper eine bestimmte Höhe gehoben, muss nur die senkrechte Bewegung berücksichtigt werden, die Bewegung normal dazu liefert keinen Beitrag zur Arbeit.

$$W \approx \sum_{i=1}^n F(s_i) \cdot \Delta s_i$$

Dieser Näherung ist umso genauer, je kleiner die Wegstrecken Δs_i gewählt werden. Lässt man die Wegstrecken gegen null gehen beziehungsweise die Anzahl der Wegstrecken gegen unendlich, so ergibt sich der exakte Wert der Arbeit als Grenzwert dieser Summenbildung und man schreibt dafür

$$W = \lim_{n \rightarrow \infty} \sum_{i=1}^n F(s_i) \cdot \Delta s_i = \int_{r_0}^{r_1} F(s) \cdot dr$$

als das Integral der Funktion in den Grenzen r_0 bis r_1 .

Wird die Gravitationskraft für die Kraft $\vec{F}$ eingesetzt, erhält man

$$W = \int_{r_0}^{r_1} F(s) \cdot ds = \int_{r_0}^{r_1} G \cdot \frac{M_{\text{Erde}} \cdot m}{r^2} \cdot dr = G \cdot M_{\text{Erde}} \cdot m \cdot \int_{r_0}^{r_1} \frac{dr}{r^2}.$$

Die Auswertung ergibt für dieses Integral

$$W = G \cdot M_{\text{Erde}} \cdot m \cdot \int_{r_0}^{r_1} \frac{dr}{r^2} = G \cdot M_{\text{Erde}} \cdot m \cdot \left(-\frac{1}{r} \right) \Big|_{r_0}^{r_1} = G \cdot M_{\text{Erde}} \cdot m \cdot \left(\frac{1}{r_0} - \frac{1}{r_1} \right)$$

Also gilt allgemein: Wenn ein Körper in einem Gravitationsfeld aus einer Höhe r_0 auf eine Höhe r_1 gehoben wird, muss die Arbeit

$$W = G \cdot M_{\text{Erde}} \cdot m \cdot \left(\frac{1}{r_0} - \frac{1}{r_1} \right)$$

verrichtet werden.

Damit können nun die beiden gestellten Aufgaben einfach durch Einsetzen gelöst werden.

Aufgabe 412

$$W = 6,673 \cdot 10^{-11} \cdot 6 \cdot 10^{24} \cdot 1 \cdot \left(\frac{1}{6,37 \cdot 10^6} - \frac{1}{6,67 \cdot 10^6} \right) \approx 2,83 \cdot 10^6 \text{ J}$$

Antwort: Um eine Masse von 1 kg von der Erdoberfläche aus auf eine Höhe von 300 km zu bringen, ist eine Arbeit von etwa $2,83 \cdot 10^6 \text{ J}$ notwendig.¹⁵⁵

Aufgabe 413

Es muss jetzt das folgende Integral berechnet werden

$$W = G \cdot M_{\text{Erde}} \cdot m \cdot \int_{r_0}^{\infty} \frac{dr}{r^2} = G \cdot M_{\text{Erde}} \cdot m \cdot \left(-\frac{1}{r} \right) \Big|_{r_0}^{\infty} = \frac{G \cdot M_{\text{Erde}} \cdot m}{r_0}$$

Eingesetzt erhält man

$$W = \frac{6,673 \cdot 10^{-11} \cdot 6 \cdot 10^{24} \cdot 1}{6,37 \cdot 10^6} \approx 6,29 \cdot 10^7 \text{ J}$$

Beispiel 2: Dehnungsarbeit einer Feder

Die Dehnungsarbeit beim Spannen einer Feder wurde bereits im Kapitel 3.2.1, Beispiel 2 mittels des Flächeninhalts berechnet. Dieses Ergebnis soll nun mit Hilfe der Integralrechnung noch einmal hergeleitet werden.

Das Beispiel ist dem Schulbuch Mathematik 8 von Götz, Reichel¹⁵⁶ entnommen.

¹⁵⁵ Hätte man statt des Gravitationsgesetzes angenommen, dass die Arbeit gegen die (konstante) Gewichtskraft verrichtet wird, ergäbe sich eine Arbeit von $W = m \cdot g \cdot s = 1 \cdot 9,81 \cdot 3 \cdot 10^5 \text{ J} = 2,943 \cdot 10^6 \text{ J}$. Dieses Ergebnis ist etwas kleiner als das Ergebnis von $2,83 \cdot 10^6 \text{ J}$, da die Gravitationskraft mit der Höhe abnimmt.

¹⁵⁶ Siehe [27] Aufgabe 404 Seite 109

404 Eine Schraubenfeder ($k = 150 \text{ N m}^{-1}$) ist gegenüber dem unbelasteten Zustand um 15 cm gedehnt. Welche Arbeit (in Joule) ist nötig, um die Feder weitere 20 cm zu dehnen?

Für die Auslenkung einer Feder gilt das Hooke'sche Gesetz $F(x) = k \cdot x$. Um eine Feder um eine kleine Auslenkung Δx zu dehnen, muss die Arbeit $\Delta W = F(x) \cdot \Delta x = k \cdot x \cdot \Delta x$ verrichtet werden. Für die gesamte verrichtete Arbeit ergibt sich dann

$$W = k \cdot \int_{x_A}^{x_E} x \cdot dx = k \cdot \left. \frac{x^2}{2} \right|_{x_A}^{x_E} = \frac{k}{2} \cdot (x_E^2 - x_A^2)$$

Durch Einsetzen erhält man

$$W = \frac{150}{2} \cdot (0,35^2 - 0,15^2) \text{ J} = 7,5 \text{ J}$$

Es werden 7,5 J benötigt, um die Feder zu dehnen.

Beispiel 3: Der Massenmittelpunkt¹⁵⁷

Der Massenmittelpunkt eines Massensystems vereinfacht die Berechnungen in der Mechanik auf zwei Weisen. Erstens dient er dazu, einen komplexen Körper auf einen Massenpunkt zu reduzieren, dem dieselbe Masse zugewiesen wird wie der Körper besitzt. Mit der Reduktion eines Körpers auf einen Massenpunkt ist es möglich, die Bahnkurve eines Körpers unter der Einwirkung einer Kraft vereinfacht zu berechnen. Als Beispiel diene ein Hammer, der geworfen wird. Der Hammer kann um seinen Massenmittelpunkt rotieren und beschreibt somit eine sehr komplizierte Bahn, der Massenmittelpunkt selbst bewegt sich aber auf der einfachen Flugparabel, die einem schiefen Wurf entspricht. Eine weitere Anwendung des Massenmittelpunktes besteht darin, dass, wenn man die Berechnungen im sogenannten Schwerpunktsystem

¹⁵⁷ Der Massenmittelpunkt und der Schwerpunkt fallen in den meisten Fällen zusammen. Genau genommen müsste man zwischen den beiden Begriffen unterscheiden, da der Schwerpunkt durch die Erdanziehung bestimmt wird, der Massenmittelpunkt aber ausschließlich von der Körpergeometrie und der Massenverteilung bestimmt ist. Wenn diese Anziehung nicht homogen ist, kann sich der Schwerpunkt vom Massenmittelpunkt unterscheiden.

durchführt, diese sich sehr vereinfachen können. Das Schwerpunktsystem hat dabei den Massenmittelpunkt als Koordinatenursprung.

Betrachten wir als erstes den Massenmittelpunkt zweier Punktmassen mit den Massen m_1 und m_2 in folgender Anordnung (Abbildung 80):

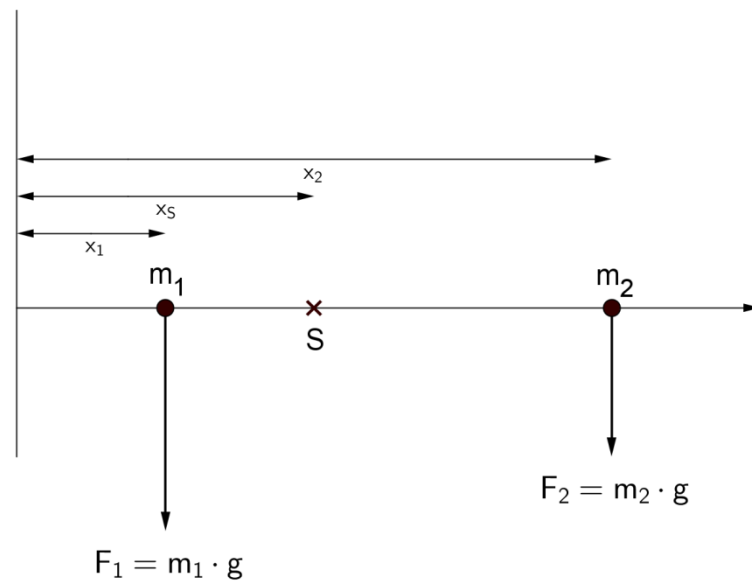


Abbildung 80: Massenmittelpunkt zweier Punktmassen

Es soll die x-Koordinate so bestimmt werden, dass die beiden Massen im Gleichgewicht sein sollen, wenn der Drehpunkt im Massenmittelpunkt liegt. Es wird das Drehmoment beider Massen berechnet:

$$\begin{aligned}(x_S - x_1) \cdot m_1 \cdot g &= (x_2 - x_S) \cdot m_2 \cdot g \\ (m_1 + m_2) \cdot x_S &= m_1 \cdot x_1 + m_2 \cdot x_2 \\ x_S &= \frac{m_1 \cdot x_1 + m_2 \cdot x_2}{M}\end{aligned}$$

Der Massenmittelpunkt zweier Punktmassen ergibt sich somit zu

$$x_S = \frac{m_1 \cdot x_1 + m_2 \cdot x_2}{M},$$

wobei M die Gesamtmasse des Systems ist.

Das kann sehr einfach auf n Massenpunkte verallgemeinert werden mit

$$x_s = \frac{m_1 \cdot x_1 + m_2 \cdot x_2 + \dots + m_n \cdot x_n}{M} = \frac{1}{M} \cdot \sum_{i=1}^n m_i \cdot x_i$$

Im nächsten Schritt stellen wir uns einen in x-Richtung ausgedehnten Körper mit einer kontinuierlichen Massenverteilung vor, diese sei durch $m(x)$ gegeben. Der Körper sei in kleine Massenstücke $\Delta m(x_1), \Delta m(x_2), \dots, \Delta m(x_n)$ aufgeteilt. Damit lässt sich die Massenmittelpunktskoordinate x_s folgendermaßen schreiben:

$$x_s \approx \frac{1}{M} \cdot \sum_{i=1}^n x_i \cdot \Delta m_i$$

Wenn man die Massenstücke immer kleiner werden lässt und schließlich den Grenzwert mit $\Delta m \rightarrow 0$ bildet, erhält man das Integral für den Massenmittelpunkt¹⁵⁸

$$x_s = \frac{1}{M} \int x \cdot dm$$

Ein Beispiel soll die Berechnung des Massenmittelpunktes verdeutlichen.¹⁵⁹

Berechne den Massenmittelpunkt eines rechtwinkligen Dreiecks mit den Kathetenlängen L und der Gesamtmasse M , wobei die Masse M homogen verteilt sei (siehe Abbildung 81)!

Es sei ρ die Massendichte pro Flächeneinheit¹⁶⁰. Daraus ergibt sich für die Gesamtmasse des Dreiecks

$$M = \rho \cdot \frac{L^2}{2}$$

¹⁵⁸ Ist der Körper räumlich ausgedehnt, erhält man die y- und z-Koordinate des Massenmittelpunktes analog:

$$y_s = \frac{1}{M} \cdot \int y \cdot dm \quad \text{und} \quad z_s = \frac{1}{M} \cdot \int z \cdot dm$$

¹⁵⁹ Siehe dazu [28] Seite 298

¹⁶⁰ Man stelle sich eine unendlich dünne Platte vor, die nur eine Längen- und Breitenausdehnung besitzt. Die Flächenmassendichte ist dann die Masse pro Flächeneinheit und hat die Einheit $\text{kg} \cdot \text{m}^{-2}$.

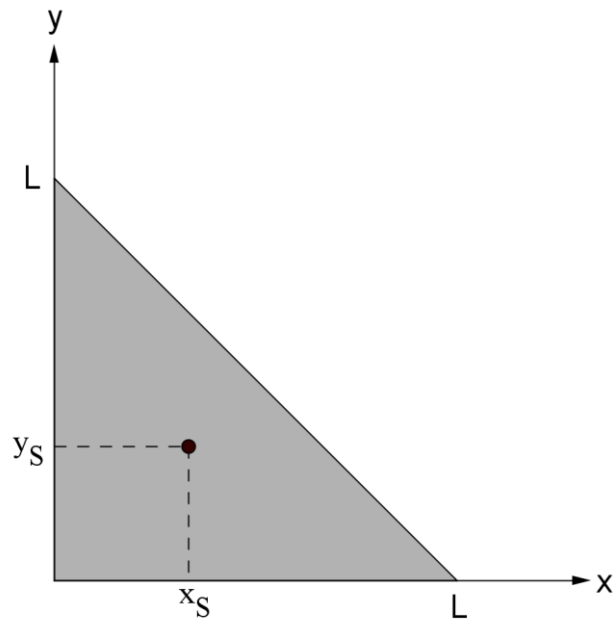


Abbildung 81: Massenmittelpunkt eines Dreiecks

Die lineare Funktion, die das Dreieck begrenzt ist gegeben durch $y = -x + L$. Damit ergibt sich für dm folgender Ausdruck

$$dm = \rho \cdot y(x) \cdot dx = \rho \cdot (L - x) \cdot dx$$

Das kann in die Formel für den Massenmittelpunkt eingesetzt werden

$$x_S = \frac{1}{M} \cdot \int_{x_A}^{x_E} x \cdot dm = \frac{2}{\rho \cdot L^2} \cdot \rho \cdot \int_0^L x \cdot (L - x) dx = \frac{2}{L^2} \cdot \int_0^L x \cdot L - x^2 dx$$

Die Auswertung des Integrals ergibt

$$x_S = \frac{2}{L^2} \cdot \left(\frac{x^2 \cdot L}{2} - \frac{x^3}{3} \right) \Big|_0^L = \frac{2}{L^2} \cdot \left(\frac{L^3}{2} - \frac{L^3}{3} \right) = \frac{2}{L^2} \cdot \frac{L^3}{6} = \frac{L}{3}$$

Aufgrund der vorliegenden Symmetrie muss die y -Koordinate des Massenmittelpunktes ebenfalls den Wert $\frac{L}{3}$ haben.

Also lässt sich die Massenmittelpunktskoordinate S angeben als

$$S = (x_S \mid y_S) = \left(\frac{L}{3}, \frac{L}{3} \right).$$

Beispiel 4: Das Trägheitsmoment

Das Trägheitsmoment spielt bei den Rotationen dieselbe Rolle wie die Masse bei den Translationen¹⁶¹. Viele Gleichungen der Rotation und Translation können völlig analog geschrieben werden, wenn man statt der Masse m das Trägheitsmoment I einsetzt.

Beispielsweise ist die kinetische Energie für translatorische Bewegungen gegeben durch $E_{\text{kin}} = \frac{1}{2} \cdot m \cdot v^2$. Für Rotationen lautet der entsprechende Ausdruck $E_{\text{rot}} = \frac{1}{2} \cdot I \cdot \omega^2$, wobei ω die Winkelgeschwindigkeit ist.

Das Trägheitsmoment ist für n Punktmassen m_i , die sich im (Normal)Abstand r_i von der Drehachse befinden, durch die Formel

$$I = \sum_{i=1}^n m_i \cdot r_i^2$$

gegeben.

Der Körper soll nun wieder eine kontinuierlich verteilte Massenverteilung aufweisen. Wenn man die Massenstücke Δm_i , aus denen man sich den Körper aufgebaut denkt, immer kleiner werden lässt und schließlich den Grenzwert mit $\Delta m \rightarrow 0$ bildet, erhält man das Integral für das Trägheitsmoment eines Körpers

$$I = \int r^2 dm$$

Zwei Beispiele sollen die Berechnung des Trägheitsmoments verdeutlichen.

1. Trägheitsmoment eines homogenen langen dünnen Stabes

Der Stab habe die Länge L und die Gesamtmasse M , die Drehachse sei normal zum Stab und befinde sich in der Entfernung a von einem Ende des Stabes (Abbildung 82).

¹⁶¹ Siehe auch Kapitel 3.6.4, Beispiel 1 und Beispiel 2

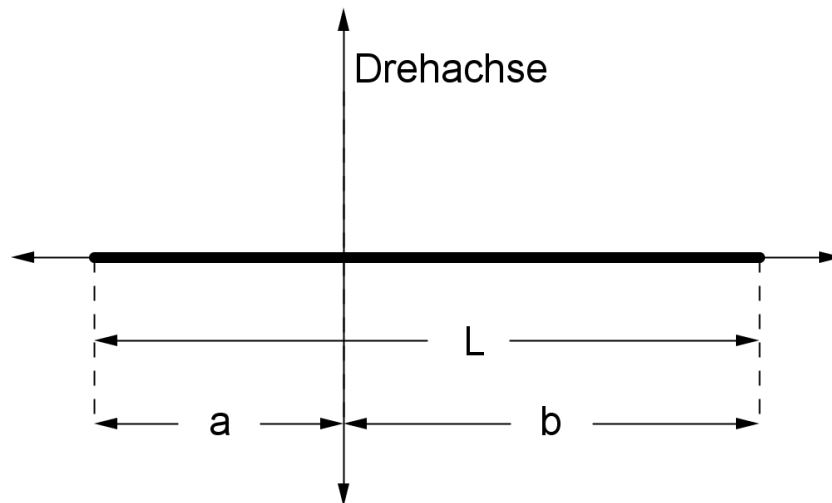


Abbildung 82: Trägheitsmoment eines rotierenden Stabes

Es sei ρ die (Linien)Massendichte¹⁶² des Stabes, dann ist $\rho = \frac{M}{L}$. Damit ergibt sich folgender Ausdruck für eine kleine Teilmasse $dm = \rho \cdot dx$ und somit das Trägheitsmoment

$$I = \int_a^b x^2 \cdot dm = \frac{M}{L} \cdot \int_a^b x^2 \cdot dx = \frac{M}{L} \cdot \frac{x^3}{3} \Big|_a^b = \frac{M}{3 \cdot L} \cdot (b^3 - a^3)$$

Sei speziell die Drehachse im Mittelpunkt des Stabes, dann ist $a = -\frac{1}{2}$ und $b = \frac{1}{2}$. Damit ergibt sich das Trägheitsmoment

$$I = \frac{M}{3 \cdot L} \cdot \left(\left(\frac{L}{2} \right)^3 - \left(-\frac{L}{2} \right)^3 \right) = \frac{M}{3 \cdot L} \cdot \frac{L^3}{4} = \frac{M \cdot L^2}{12}$$

Sei die Drehachse jetzt an einem Ende des Stabes, dann ist $a = 0$ und $b = L$. Damit ergibt sich

$$I = \frac{M}{3 \cdot L} \cdot (L^3 - 0) = \frac{M \cdot L^2}{3}$$

¹⁶² Man stelle sich einen unendlich dünnen Stab vor, der nur eine Längenausdehnung besitzt. Die Linienmassendichte ist dann die Masse pro Längeneinheit und hat die Einheit $\text{kg} \cdot \text{m}^{-1}$.

Damit sieht man, dass das Trägheitsmoment eines Körpers von der Lage der Drehachse abhängig ist. Derselbe Körper kann ganz verschiedene Trägheitsmomente besitzen je nachdem, welche Drehachse man betrachtet.

2. Trägheitsmoment einer Schwungscheibe

Schwungscheiben sind Zylinder mit einer großen Masse, die durch ihre Rotation hohe Energiemengen speichern können. Deswegen sind sie in vielen technischen Anwendungen bis hin zu einfachen Spielzeugautos zu finden. Es sei eine Schwungscheibe mit dem Radius R und der Höhe H gegeben, die die Gesamtmasse M haben soll, die homogen über den Vollzylinder verteilt ist. Die Drehachse sei die Symmetrieachse des Zylinders (vergleiche Abbildung 83).

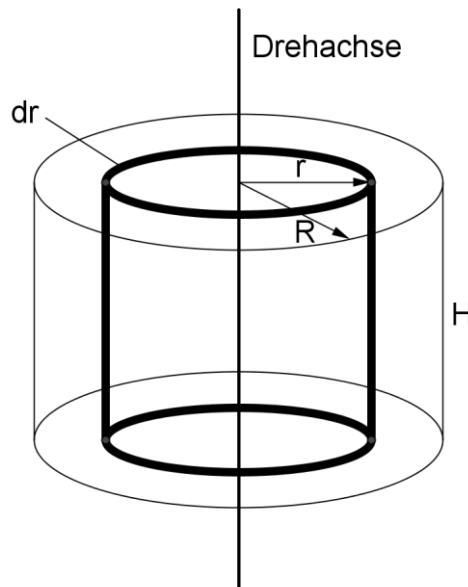


Abbildung 83: Trägheitsmoment einer Schwungscheibe

Alle Massenelemente, die sich im selben Abstand von der Drehachse befinden, tragen denselben Beitrag zum Trägheitsmoment bei. Deswegen betrachten wir einen sehr dünnen Hohlzylinder mit der Dicke dr . Die Masse des Hohlzylinders beträgt näherungsweise

$$dm = \rho \cdot 2\pi r \cdot H \cdot dr$$

wobei $\rho = \frac{M}{V} = \frac{M}{R^2\pi H}$ die Massendichte der Schwungscheibe ist.

Damit ergibt sich für das Trägheitsmoment

$$I = 2\pi H\rho \cdot \int_0^R r^3 dr = \frac{2\pi HM}{R^2\pi H} \cdot \frac{r^4}{4} \Big|_0^R = \frac{2M}{R^2} \cdot \frac{R^4}{4} = \frac{M \cdot R^2}{2}$$

$$I = \frac{M \cdot R^2}{2}$$

Das heißt, die über den Vollzylinder gleichmäßig verteilte Masse hat dieselbe Wirkung wie ein Massenpunkt, der die halbe Masse des Schwungrades hat und im Abstand R um die Drehachse rotiert!

Beispiel 5: Mittlere Leistung von Wechselstrom

Die elektrische Leistung $P(t)$ des elektrischen Stroms ist durch das Produkt der Stromstärke $I(t)$ und der Spannung $U(t)$ gegeben

$$P(t) = I(t) \cdot U(t) .$$

Die elektrische Stromstärke und die elektrische Spannung eines Wechselstromes sind gegeben durch $I(t) = I_0 \cdot \sin(\omega t)$ und $U(t) = U_0 \cdot \sin(\omega t)$, wobei $\omega = \frac{2\pi}{T}$ und T die Dauer einer Schwingung ist. Für die elektrische Momentanleistung ergibt sich dann

$$P(t) = I_0 \cdot U_0 \cdot \sin^2(\omega t) .$$

Über einen Zeitraum t verrichtet der elektrische Strom die elektrische Arbeit W , die als Leistung multipliziert mit der Zeit berechnet wird. Nun ist beim Wechselstrom die Leistung nicht konstant, so dass die elektrische Arbeit durch Integration berechnet werden muss. Da der elektrische Strom ein periodischer Vorgang ist, genügt es, eine Periode T zu betrachten.

Die elektrische Arbeit berechnet sich zu

$$W = I_0 U_0 \cdot \int_0^T \sin^2(\omega t) dt$$

Die mittlere Leistung wird nun als die konstante Leistung definiert, die über eine Periode dieselbe elektrische Arbeit liefern würde, somit soll gelten

$$\bar{P} \cdot T = W = \int_0^T I(t) \cdot U(t) dt .$$

Damit ergibt sich die Definition der mittleren Leistung zu

$$\boxed{\bar{P} = \frac{1}{T} \cdot \int_0^T I(t) \cdot U(t) dt}$$

Konkret ergibt sich dann für die mittlere Wechselstromleistung

$$\bar{P} = \frac{I_0 \cdot U_0}{T} \cdot \int_0^T \sin^2(\omega t) dt$$

Das Integral lässt sich durch partielle Integration lösen.

$$\begin{aligned} \int_0^T \sin^2(\omega t) dt &= \int_0^T \underbrace{\sin\left(\frac{2\pi t}{T}\right)}_{u'} \cdot \underbrace{\sin\left(\frac{2\pi t}{T}\right)}_v dt \\ &= \underbrace{-\frac{T}{2\pi} \cdot \cos\left(\frac{2\pi t}{T}\right) \cdot \sin\left(\frac{2\pi t}{T}\right)}_{=0} \bigg|_0^T + \frac{T}{2\pi} \cdot \frac{2\pi}{T} \cdot \int_0^T \cos^2\left(\frac{2\pi t}{T}\right) dt \\ &= \int_0^T 1 - \sin^2\left(\frac{2\pi t}{T}\right) dt \\ &= T - \int_0^T \sin^2\left(\frac{2\pi t}{T}\right) dt \end{aligned}$$

Damit ist

$$\begin{aligned} \int_0^T \sin^2\left(\frac{2\pi t}{T}\right) dt &= T - \int_0^T \sin^2\left(\frac{2\pi t}{T}\right) dt \\ 2 \cdot \int_0^T \sin^2\left(\frac{2\pi t}{T}\right) dt &= T \\ \int_0^T \sin^2\left(\frac{2\pi t}{T}\right) dt &= \frac{T}{2} \end{aligned}$$

Für die mittlere Leistung ergibt sich daraus

$$\bar{P} = \frac{I_0 \cdot U_0}{T} \cdot \frac{T}{2} = \frac{I_0 \cdot U_0}{2}$$

Das heißt, auf längere Zeit ist bei einem Wechselstrom die konstante Leistung $\frac{I_0 \cdot U_0}{2}$ wirksam. Siehe auch Abbildung 84.

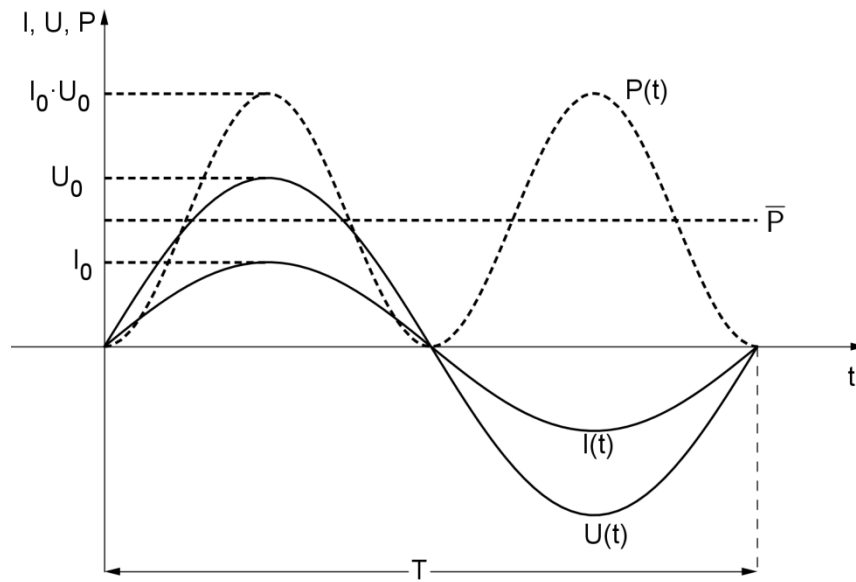


Abbildung 84: Momentan- und mittlere Wechselstromleistung

3.8.2 Differentialgleichungen

Dieses Kapitel über Differentialgleichungen stellt den Übergang her von der Schulphysik zur Physik, die dann weiter an den Universitäten gelehrt wird. Das Kapitel markiert sozusagen die Grenze dessen, bis zu der die Schulphysik zu gehen vermag und wo die Mathematik dann auch nicht mehr ausreichend ist, die zum Schulstoff gehört.

Wird ein Physiker gefragt, in welcher Form der Mathematik die bestimmenden Gleichungen der Physik geschrieben sind, kann als Antwort sicher gegeben werden, dass das in Form von Differentialgleichungen geschieht. Im Schulunterricht ist das Thema Differentialgleichungen sicher ein wenn überhaupt vernachlässigtes. Ein Grund mag darin liegen, dass LehrerInnen, die nur Mathematik unterrichten, nicht die entsprechenden Vorkenntnisse beziehungsweise den Einblick in die Wichtigkeit von Differentialgleichungen in der Physik haben. Es kann im Schulunterricht sicher nur sehr eingeschränkt auf das Thema eingegangen werden, es sollte aber meines Erachtens doch zumindest zur Sprache gebracht werden und eventuell eine einfache Differentialgleichung hergeleitet und gelöst werden, um die Vorgangsweise zu erläutern. Zumindest aber die eine mündliche Matura ist das Thema lohnend und sollte von den SchülerInnen, da man doch in diesem Fall ein größeres Interesse voraussetzen kann, verstanden und behandelt werden.

Einige der wichtigsten Differentialgleichungen in der Physik seien im folgenden ersten Beispiel genannt und erläutert.

Beispiel 1: Drei wichtige Differentialgleichungen der Physik

1. Die Bewegungsgleichung der Mechanik

Das 2. Newton'sche Gesetz lautet in Vektorform $m \cdot \vec{a} = \vec{F}$. Dabei ist die Beschleunigung $\vec{a}(t)$ die zweite Ableitung des Ortes nach der Zeit, somit kann die Bewegungsgleichung geschrieben werden als

$$m \cdot \frac{d^2 \vec{x}(t)}{dt^2} = \vec{F}(\vec{x}, t)$$

Dabei ist $\vec{F}(\vec{x}, t)$ die äußere Kraft, die im allgemeinen Fall vom Ort $\vec{x}$ als auch von der Zeit t abhängen kann. Ist die äußere Kraft nur vom Ort abhängig und nicht von der Zeit, so vereinfacht sich für zeitunabhängige Probleme die Bewegungsgleichung zu

$$m \cdot \frac{d^2 \vec{x}(t)}{dt^2} = \vec{F}(\vec{x})$$

Die äußere Kraft $\vec{F}$ kann natürlich verschiedenste Formen annehmen und macht damit das Lösen der Bewegungsgleichung entsprechend schwierig. Man denke nur an das Problem der Bewegung eines Planeten um die Sonne. Im einfachsten Fall ist die äußere Kraft die Gravitationskraft, mit der der Planet von der Sonne angezogen wird. Als Lösung erhält man die bekannten Ellipsenbahnen der Planeten. Für genauere Berechnungen müssen aber auch die gegenseitigen Anziehungskräfte der anderen Körper im Sonnensystem mitberücksichtigt werden. Dadurch ist es etwa nicht möglich, etwa das Dreikörperproblem, wenn drei Körper sich gegenseitig anziehen, mathematisch in geschlossener Form zu lösen, es müssen dann numerische Methoden herangezogen werden.

Um die Wichtigkeit dieser Gleichung zu begreifen, muss man sich klarmachen, dass mit dieser Differentialgleichung bei bekannten äußeren Kräften alle Bewegungsvorgänge gelöst werden können, das heißt diese Gleichung beschreibt die gesamte Kinetik.

Im nächsten Beispiel 2 soll diese Gleichung dann exemplarisch für den freien Fall gelöst werden.

2. Die 4 Maxwellgleichungen der Elektrodynamik

Eine der großartigsten Errungenschaften der Physik des 19. Jahrhunderts war die Aufstellungen dieser vier gekoppelten Differentialgleichungen durch Maxwell¹⁶³.

¹⁶³ James Clerk Maxwell; 1831 – 1879, schottischer Physiker, entwickelte die vier nach ihm benannten Maxwell-Gleichungen, die die Grundlage der Elektrodynamik darstellen sowie die kinetische Gastheorie, welche die statistische Mechanik begründete.

Die Mathematik dieser Gleichungen geht über den Schulstoff weit hinaus und soll hier auch nicht ansatzweise versucht werden. Das Schwierige, aber auch was diese Gleichungen so interessant machen, ist, dass es sich um sogenannte gekoppelte Differentialgleichungen handelt, das heißt die Gleichungen hängen voneinander ab und müssen als Gesamtheit gelöst werden.

Es soll nur darauf hingewiesen werden, dass mit diesen 4 Gleichungen die gesamte Elektrodynamik hergeleitet werden kann, sie verknüpfen die bis dahin getrennten Begriffe des elektrischen und des magnetischen Feldes zum Begriff des elektromagnetischen Feldes. Auch die gesamte Theorie der Optik, solange es sich dabei um die klassische Optik handelt, ist in diesen Gleichungen enthalten.

Die Gleichungen sehen folgendermaßen aus:

$$\begin{aligned}\operatorname{div} \vec{E} &= \frac{\rho}{\epsilon_0} \\ \operatorname{div} \vec{B} &= 0 \\ \operatorname{rot} \vec{E} &= -\frac{\partial \vec{B}}{\partial t} \\ \operatorname{rot} \vec{B} &= \mu_0 \vec{j} + \mu_0 \epsilon_0 \frac{\partial \vec{E}}{\partial t}\end{aligned}$$

Dabei sind $\vec{E}$ und $\vec{B}$ die elektrische und magnetische Feldstärke, ρ und $\vec{j}$ die elektrische Ladungsdichte Stromdichte und ϵ_0 und μ_0 sind bestimmte Konstanten. Div und rot (Divergenz und Rotation) sind sogenannte Differentialoperatoren, in denen die Ableitungen der Größen $\vec{E}$ und $\vec{B}$ vorkommen.

Das Symbol ∂ , das in allen diesen vier Gleichungen vorkommt ist, steht für eine sogenannte partielle Ableitung. Das sind Ableitungen nach einer Variablen einer Größe, wenn diese Größe von mehreren Variablen abhängen kann. Die beiden Feldstärken $\vec{E}$ und $\vec{B}$ hängen sowohl vom Ort als auch von der Zeit ab. Da der Ort durch die drei Koordinaten x, y, z angegeben wird, können die Feldstärken von vier Variablen abhängen, drei Orts- und einer Zeitkoordinate, man schreibt dann $\vec{E}(x, y, z, t)$ beziehungsweise $\vec{B}(x, y, z, t)$. Wird nun etwa die elektrische Feldstärke

$\vec{E}$ nach der Zeit abgeleitet, so schreibt man nicht $\frac{d\vec{E}(x,y,z,t)}{dt}$, sondern stattdessen $\frac{\partial \vec{E}(x,y,z,t)}{\partial t}$.

Auch ohne die Mathematik hinter den Gleichungen verstehen zu können, sieht man doch in der dritten und vierten Maxwellgleichung, dass die elektrische und magnetische Feldstärke voneinander abhängen. Die vorkommenden ersten Zeitableitungen von $\vec{E}$ und $\vec{B}$ sind in den Gleichungen mit der jeweils anderen Feldstärke verknüpft. Ein elektrisches Feld, das sich mit der Zeit ändert, ruft ein magnetisches Feld hervor und umgekehrt ein zeitlich veränderliches magnetisches Feld ein elektrisches Feld.

Genauso stellt man sich Licht vor. Licht ist nichts Anderes als zeitlich veränderliche elektrische und magnetische Felder, die sich gegenseitig erzeugen, dabei stehen die elektrischen und magnetischen Feldvektoren normal aufeinander und beide normal auf die Ausbreitungsrichtung des Lichtes. Licht ist somit eine transversale Welle.

3. Die Schrödingergleichung der Quantenmechanik

Beschreibt die Bewegungsgleichung der Mechanik die Bewegung von makroskopischen Körpern, wie etwa Planeten im Sonnensystem, so kann diese Bewegungsgleichung für die Bewegung der Elektronen im Atom nicht mehr verwendet werden, sie verliert ihre Gültigkeit im mikroskopischen Bereich. Zur Beschreibung atomarer Vorgänge ist eine neue Beschreibung notwendig, die quantenmechanische, und die zugrundeliegende Differentialgleichung ist die Schrödingergleichung¹⁶⁴.

Die Schrödingergleichung gilt normalerweise in drei Dimensionen, etwa wenn Elektronen in Atomen beschrieben werden sollen, aber für gewisse Aufgabenstellungen genügt es, die Gleichung in nur einer Dimension

¹⁶⁴ Erwin Schrödinger; 1887 – 1961, österreichischer Physiker, er stellte die nach ihm benannte Schrödingergleichung im Jahr 1926 auf.

aufzuschreiben. In einer Dimension lautet die (zeitabhängige) Schrödingergleichung folgendermaßen:

$$-\frac{\hbar^2}{2m} \cdot \frac{\partial^2 \psi(x,t)}{\partial x^2} + E_{\text{Pot}}(x,t) \cdot \psi(x,t) = -i\hbar \cdot \frac{\partial \psi(x,t)}{\partial t}$$

Hierbei ist $\hbar$ das sogenannte Planck'sche Wirkungsquantum, eine der fundamentalen Naturkonstanten der Physik, neben der Lichtgeschwindigkeit und der Gravitationskonstanten. $\hbar$ ist eigentlich das Planck'sche Wirkungsquantum h geteilt durch 2π und hat den Wert $6,626 \cdot 10^{-34} \text{ J} \cdot \text{s}$.

Bemerkenswert ist, dass mit dieser Bewegungsgleichung nicht mehr die Bewegung etwa eines Elektrons in einem Atom an sich beschrieben wird, denn der klassische Begriff der Bahnkurve verliert in der Quantenmechanik seine Bedeutung und deswegen muss eine andere Beschreibungsmöglichkeit gefunden werden. Diese neue Art der Beschreibung liefert der Begriff der Aufenthaltswahrscheinlichkeit eines Teilchens. Was die Schrödingergleichung liefert, sind Aussagen darüber, wo sich das Elektron mit einer bestimmten Wahrscheinlichkeit aufhält. Dies wird durch die sogenannte Wellenfunktion $\psi(x,t)$ erfüllt. Die Wellenfunktion selbst hat in dieser Form keine physikalische Bedeutung, sie kann sogar komplexwertig sein. Erst das Quadrat von $\psi(x,t)$, also $|\psi(x,t)|^2$ ist eine reelle Zahl und hat die Bedeutung einer Aufenthaltswahrscheinlichkeit des Teilchens am Ort x zu einer bestimmten Zeit t .

Dieses Quadrat der Wellenfunktion $|\psi(x,t)|^2$ ist dann auch die Mathematik, die hinter den Orbitalen, den Aufenthaltsräumen der Elektronen im Atom, steckt, die man in der Schule in Physik, aber mehr noch in der Chemie lernt und die essentiell sind, um chemische Verbindungen und Reaktionen angemessen beschreiben und verstehen zu können.

Beispiel 2: Lösen der Bewegungsgleichung für den schiefen Wurf

Betrachten wir einen Körper, der zur Zeit $t = 0$ unter einem Winkel α zur Horizontalen abgeschossen wird, das heißt, er erhält zur Zeit $t = 0$ die Geschwindigkeit v_0 vermittelt und wird sich dann für die Zeiten $t > 0$ selbst überlassen. Der Abschuss soll bei $x = 0$ erfolgen in einer Höhe h_0 . Die Situation zur Zeit $t = 0$ ist in der Abbildung 85 veranschaulicht.

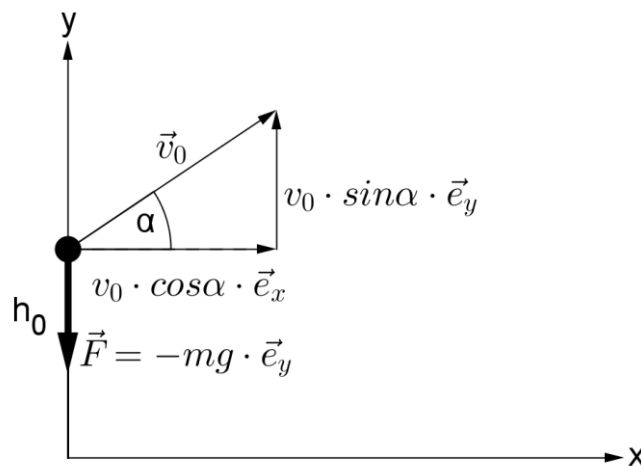


Abbildung 85: Der schiefe Wurf zur Zeit $t = 0$

Ausgangspunkt ist die Differentialgleichung der Bewegung, das heißt das zweite Newton'sche Gesetz $\vec{F} = m \cdot \vec{a}$, wobei für diese Aufgabenstellung die Behandlung in zwei Dimensionen genügt, da die Bewegung in einer Ebene stattfindet. Die Bewegungsgleichung lautet somit

$$m \cdot \begin{pmatrix} a_x \\ a_y \end{pmatrix} = \begin{pmatrix} F_x \\ F_y \end{pmatrix}$$

Als einzige wirkende Kraft kommt das Gewicht des Körpers zum Tragen, das zu jedem Zeitpunkt t konstant ist, das heißt, dass $F_x = 0$ und $F_y = -m \cdot g$, also

$$m \cdot \begin{pmatrix} a_x \\ a_y \end{pmatrix} = \begin{pmatrix} 0 \\ -m \cdot g \end{pmatrix} = m \cdot \begin{pmatrix} 0 \\ -g \end{pmatrix}$$

Nach Division beider Seiten durch m lässt sich die Bewegungsgleichung für den schiefen Wurf endgültig schreiben als

$$\begin{pmatrix} a_x \\ a_y \end{pmatrix} = \begin{pmatrix} 0 \\ -g \end{pmatrix} \quad (1)$$

Da die Kraft auf den Körper immer in die negative y-Richtung zeigt, erhält sie ein negatives Vorzeichen. Diese Vektorgleichung kann nun unter der Berücksichtigung, dass die Beschleunigung die zweite Ableitung des Ortes nach der Zeit ist, in zwei Koordinatengleichungen geschrieben werden.

$$\begin{aligned} a_x(t) = \ddot{x}(t) &= \frac{d^2 x(t)}{dt^2} = 0 \\ a_y(t) = \ddot{y}(t) &= \frac{d^2 y(t)}{dt^2} = -g \end{aligned} \quad (2)$$

Das ist ein System von zwei Differentialgleichungen in zwei Variablen, die nicht gekoppelt sind, das heißt, sie sind nach ihren Koordinaten getrennt. Deshalb vereinfacht sich die Lösung beträchtlich auf die getrennte Lösung der beiden Koordinatengleichungen.

Lösung der 1. Gleichung:

Beide Seiten der ersten Gleichung von (2) werden integriert, wobei sich dann, da die Ableitung der Geschwindigkeit die Beschleunigung ist, ergibt

$$\int \ddot{x}(t) dt = \int 0 dt \quad (3)$$

mit der Lösung

$$\dot{x}(t) = c_1.$$

c_1 ist dabei eine Konstante, die noch aus den Anfangsbedingungen bestimmt werden muss. Jetzt können wieder beide Seiten integriert werden, wobei die Geschwindigkeit die erste Ableitung des Ortes nach der Zeit ist und man erhält

$$\int \dot{x}(t) dt = \int c_1 dt$$

mit der Lösung

$$x(t) = c_1 \cdot t + c_2 \quad (4)$$

Wieder ist c_2 eine Integrationskonstante, die aus den Anfangsbedingungen bestimmt werden muss.

Die Anfangsbedingungen lauten $\dot{x}(t=0) = v_0 \cdot \cos \alpha$ und $x(t=0) = 0$.

Aus (3) ergibt sich

$$\dot{x}(t) = c_1 \Rightarrow \dot{x}(t=0) = c_1 = v_0 \cdot \cos \alpha$$

Aus (4) ergibt sich

$$x(t) = v_0 \cdot \cos \alpha \cdot t + c_2 \Rightarrow x(t=0) = c_2 = 0$$

Die Lösung der Bewegung in x-Richtung ist somit

$$x(t) = v_0 \cdot \cos \alpha \cdot t.$$

Lösung der 2. Gleichung:

Beide Seiten der zweiten Gleichung von (2) können wie bei der Lösung der 1. Gleichung integriert werden und man erhält

$$\int \ddot{y}(t) dt = -\int g dt$$

mit der Lösung

$$\dot{y}(t) = -g \cdot t + c_3 \quad (5)$$

wobei c_3 wieder eine noch zu bestimmende Integrationskonstante ist. Jetzt wird (5) wieder integriert

$$\int \dot{y}(t) dt = \int -g \cdot t + c_3 dt$$

mit der Lösung

$$y(t) = -\frac{g}{2} \cdot t^2 + c_3 \cdot t + c_4 \quad (6)$$

und der Integrationsvariablen c_4 .

Die Anfangsbedingungen lauten $\dot{y}(t=0) = v_0 \cdot \sin \alpha$ und $y(t=0) = h_0$.

Aus (5) ergibt sich

$$\dot{y}(t) = -g \cdot t + c_3 \Rightarrow \dot{y}(t=0) = c_3 = v_0 \cdot \sin \alpha$$

Aus (6) ergibt sich

$$y(t) = -\frac{g}{2} \cdot t^2 + v_0 \cdot \sin \alpha \cdot t + c_4 \Rightarrow y(t=0) = c_4 = h_0$$

Die Lösung der Bewegung in y-Richtung ist somit

$$y(t) = -\frac{g}{2} \cdot t^2 + v_0 \cdot \sin \alpha \cdot t + h_0.$$

Die gesamte Lösung der Bewegungsgleichung des schiefen Wurfes lautet somit:

$$\vec{x}(t) = \begin{pmatrix} v_0 \cdot \cos \alpha \cdot t \\ -\frac{g}{2} \cdot t^2 + v_0 \cdot \sin \alpha \cdot t + h_0 \end{pmatrix}$$

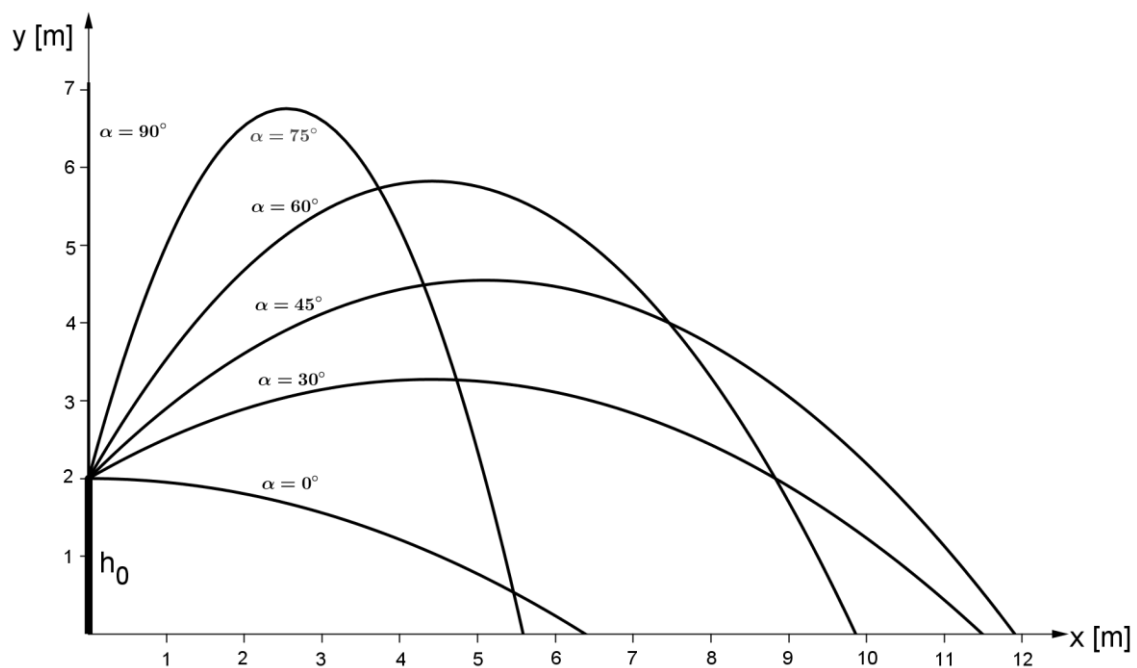


Abbildung 86: Wurfbahnen mit verschiedenen Abschusswinkeln

In Abbildung 86 sind einige ausgewählte Wurfbahnen für verschiedene Anfangsbedingungen gezeichnet. Die Anfangsgeschwindigkeit beträgt immer dieselbe und beträgt $10\text{ m}\cdot\text{s}^{-1}$, die Abschusswinkel sind allerdings unterschiedlich und betragen $\alpha = 0^\circ, 30^\circ, 45^\circ, 60^\circ, 75^\circ, 90^\circ$.

Für den Winkel $\alpha = 90^\circ$ erhält man den Spezialfall des senkrechten Wurfes. Da $\cos 90^\circ = 0$ und $\sin 90^\circ = 1$ ist, ergibt sich für den senkrechten Wurf folgende Lösung:

$$\vec{x}(t) = \begin{pmatrix} 0 \\ -\frac{g}{2} \cdot t^2 + v_0 \cdot t + h_0 \end{pmatrix}$$

In der Abbildung 87 sind ebenfalls wie in Abbildung 86 einige ausgewählte Wurfbahnen für verschiedene Anfangsbedingungen gezeichnet. Hier ist aber der Abschusswinkel immer derselbe mit $\alpha = 45^\circ$. Die Anfangsgeschwindigkeiten sind dagegen verschieden mit $v_0 = 0\text{ m}\cdot\text{s}^{-1}, 5\text{ m}\cdot\text{s}^{-1}, 10\text{ m}\cdot\text{s}^{-1}, 15\text{ m}\cdot\text{s}^{-1}, 20\text{ m}\cdot\text{s}^{-1}$.

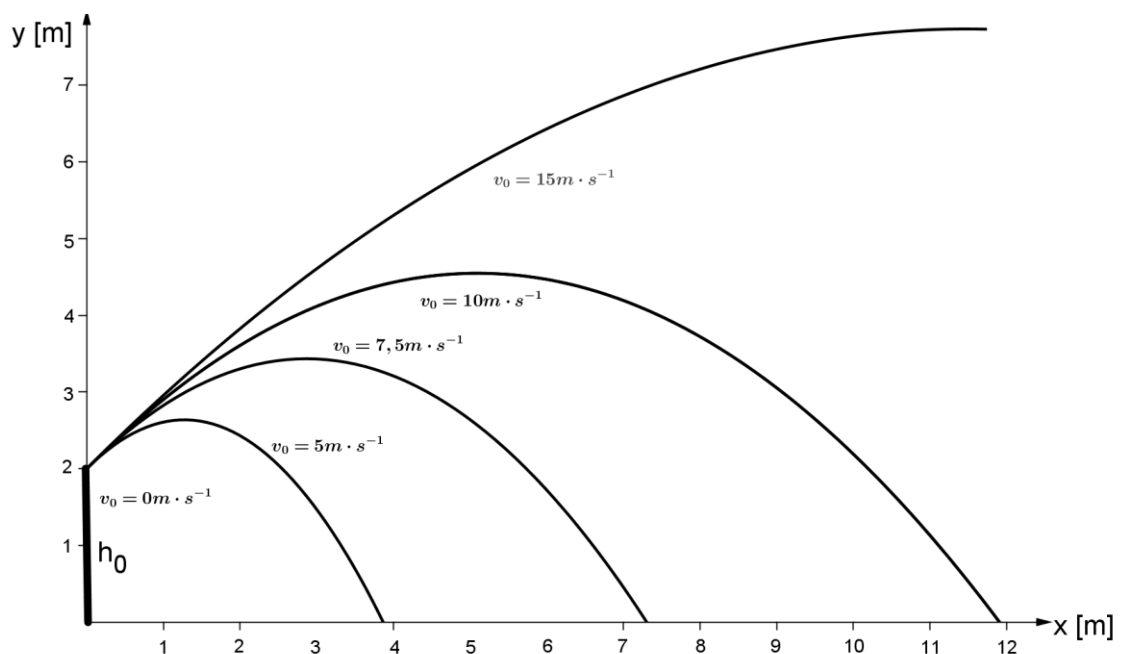


Abbildung 87: Wurfbahnen mit verschiedenen Anfangsgeschwindigkeiten

Für die Anfangsgeschwindigkeit $v_0 = 0 \text{ m} \cdot \text{s}^{-1}$ erhält man den Spezialfall des freien Falles:

$$\vec{x}(t) = \begin{pmatrix} 0 \\ -\frac{g}{2} \cdot t^2 + h_0 \end{pmatrix}$$

Beispiel 3: Aufstellen und Lösen der Exponentialgleichung für den radioaktiven Zerfall

Von den bekannten etwa 3300 Isotopen der chemischen Elemente sind ungefähr 250 stabil, die restlichen sind radioaktiv und zerfallen nach einer gewissen Zeit. Es zerfallen aber nun nicht alle Kerne gleichzeitig, sondern einzeln innerhalb eines bestimmten Zeitintervalls. Man kann jedoch nicht mit Sicherheit sagen, dass ein bestimmter Kern innerhalb dieses Zeitintervalls zerfällt, sondern nur mit einer gewissen Wahrscheinlichkeit. Der radioaktive Zerfall ist also ein Zufallsprozess.

Man kann nun aber bemerkenswerterweise trotzdem aufgrund der Tatsache, dass eine makroskopische Probe immer eine überaus große Anzahl von Atomen¹⁶⁵ enthält, sehr genau angeben, welcher Anteil der Probe in einer gewissen Zeit radioaktiv zerfallen wird.

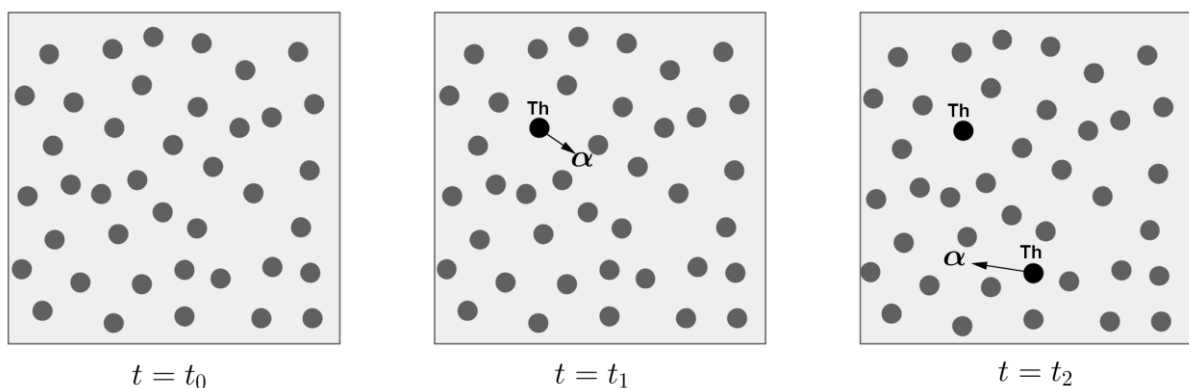


Abbildung 88: Der radioaktive Zerfall als Zufallsprozess

¹⁶⁵ Uran-238 etwa hat eine relative Atommasse von ungefähr 238. Das heißt, dass 238 g Uran-238 1 mol Uran-238-Atome enthält, das entspricht wiederum etwa der Avogadro'schen Zahl $N_A \approx 6 \cdot 10^{23}$ an Kernen. Nimmt man nur 1/1000 g U-238, so sind das immer noch $6 \cdot 10^{23} / 238 \approx 2,5 \cdot 10^{21}$ Atomkerne in 1/1000 g U-238, eine immense Zahl, so dass die Wahrscheinlichkeitsrechnung hier ungemein genaue Aussagen machen kann.

Abbildung 88 soll dies verdeutlichen. Die grauen Kreise sollen U-238-Kerne darstellen, die schwarzen mit Th beschrifteten Kreise Th-234-Kerne, die aus dem Urankernen durch einen α -Zerfall entstehen. Zur Zeit $t = t_0$ sei in der Probe eine gewisse Anzahl von Urankernen vorhanden, in der Abbildung im ersten Bild sind das 40. Zu einem bestimmten Zeitpunkt $t = t_1$ später zerfällt ein nicht genauer bestimmbarer Urankern in einen Thoriumkern. Es sind dann in der Probe um einen Urankern weniger enthalten, dafür um einen Thoriumkern mehr (im Bild 2 aus Abbildung 88 sind das noch 39 Urankerne und ein Thoriumkern). Zu einem anderen Zeitpunkt $t = t_2$ später zerfällt ein weiterer Urankern in einen Thoriumkern, die Anzahl der Urankerne verringert sich wieder um einen und die Anzahl der Thoriumkerne vergrößert sich um einen (im Bild 2 aus Abbildung 88 sind das jetzt noch 38 Urankerne und zwei Thoriumkerne). Wichtig ist, dass ein Thoriumkern nicht wieder in einen Urankern zurückzerfallen kann, das bedeutet, dass die Anzahl der noch vorhandenen Urankerne sukzessive mit der Zeit abnimmt.¹⁶⁶

Man kann sich nun gedanklich vorstellen, dass jedem Kern in der Probe die gleiche Zerfallswahrscheinlichkeit zugeordnet wird. Man macht nun die Näherung, dass man annimmt, dass in einem sehr kurzen Zeitintervall Δt die Zahl der Zerfälle proportional zu Δt und proportional zur Anzahl der noch nicht zerfallenen Kerne N ist.¹⁶⁷

Dies kann man mathematisch anschreiben als

$$\Delta N = -\lambda \cdot N \cdot \Delta t.$$

Das Minuszeichen kommt daher, weil die Zahl der Kerne mit der Zeit abnimmt. λ ist eine Konstante, die die „Geschwindigkeit“ angibt, mit der die Anzahl der Kerne in der Probe abnimmt, sie ist für jedes radioaktive Nuklid charakteristisch.

¹⁶⁶ Der Thoriumkern kann natürlich selber weiter zerfallen, aber nicht mehr in Uran. Er zerfällt nach einer gewissen Zeit in einen Pa-234-Kern, dieser Protactiniumkern zerfällt dann wieder einen anderen Kern und so weiter. Die Uranprobe wird sich mit der Zeit so entwickeln, dass immer weniger Uran enthalten ist und eine immer größere Anzahl der verschiedensten Kerne in der Probe enthalten sein wird. Die entstandenen neuen Kerne werden so lange weiter zerfallen, bis sie bei einem Kern ankommen, der stabil ist, was im Allgemeinen ein Bleikern sein wird. Auf diese Weise entstehen die sogenannten Zerfallsreihen.

¹⁶⁷ Man nimmt also etwa an, dass in einer der doppelt so großen (kurzen) Zeitspanne, also $2 \cdot \Delta t$ auch die doppelte Anzahl von Zerfällen stattfindet, genauso, wenn die Zahl der vorhandenen Kerne doppelt so groß ist, also $2 \cdot N$. Diese Annahme darf man treffen, und sie ist umso genauer je kleiner das Zeitintervall Δt ist, weil sich im sehr kurzen Zeitintervall die Zahl der Kerne nur sehr wenig ändert im Vergleich zur Gesamtzahl der vorhandenen Kerne.

Wenn man das Zeitintervall Δt immer kleiner werden lässt und schließlich den Grenzübergang $\Delta t \rightarrow 0$ bildet, kann man diesen Zusammenhang als Grenzübergang schreiben als

$$\lim_{\Delta t \rightarrow 0} \frac{\Delta N(t)}{\Delta t} = \frac{dN(t)}{dt} = -\lambda \cdot N(t)$$

Man erhält somit die Differentialgleichung für den radioaktiven Zerfall

$$\frac{dN(t)}{dt} = -\lambda \cdot N(t)$$

Diese Differentialgleichung gilt es nun zu lösen, um zu einer Gleichung für die Anzahl der nach der Zeit t noch vorhandenen Kerne zu gelangen. Man schreibt die Gleichung formal in folgender Form an (Leibniz'sche Schreibweise, in der der Differentialquotient wie ein Bruch behandelt wird):

$$dN(t) = -\lambda \cdot N(t) \cdot dt$$

Man formt die Gleichung noch folgend um

$$\frac{dN}{N} = -\lambda \cdot dt$$

und integriert jetzt beide Seiten

$$\int \frac{dN}{N} = -\lambda \cdot \int dt$$

Die Integration ist elementar ausführbar und liefert als Ergebnis

$$\ln N = -\lambda \cdot t + c$$

wobei c eine beliebige Konstante ist, die durch eine Anfangsbedingung bestimmt werden kann. Hieraus kann leicht nach N aufgelöst werden

$$e^{\ln N} = N = e^{-\lambda \cdot t + c} = e^c \cdot e^{-\lambda \cdot t} = c' \cdot e^{-\lambda \cdot t}$$

e^c wird als neue Konstante c' bezeichnet und es bleibt nun nur noch die Konstante c' zu bestimmen. Sei die Anzahl der Kerne, die zur Zeit $t=0$ vorhanden waren N_0 , so ergibt sich

$$N(t=0) = c' = N_0$$

Damit lässt sich das Zerfallsgesetz in der endgültigen Form schreiben als

$$N(t) = N_0 \cdot e^{-\lambda \cdot t}$$

Es soll noch angemerkt werden, dass man die Integration der Differentialgleichung auch als bestimmtes Integral hätte durchführen können. Man schreibt dann

$$\int_{N_0}^N \frac{dN}{N} = -\lambda \cdot \int_0^t dt$$

Das liefert als Lösung

$$\begin{aligned} \ln N \Big|_{N_0}^N &= -\lambda \cdot t \Big|_0^t \\ \ln N - \ln N_0 &= -\lambda \cdot t \\ \ln \frac{N}{N_0} &= -\lambda \cdot t \end{aligned}$$

Und damit für N

$$\frac{N}{N_0} = e^{-\lambda \cdot t}$$

Man erhält wieder dieselbe Gleichung für den radioaktiven Zerfall wie zuvor

$$N(t) = N_0 \cdot e^{-\lambda \cdot t}$$

3.8.3 Modelle und Modellbildung

Dieses letzte Kapitel kann sicherlich so nicht als normaler Schulstoff unterrichtet werden, sondern kann für besonders interessierte oder begabte SchülerInnen, etwa im Bereich eines Wahlpflichtfaches oder in einer besonders interessierten und leistungsstarken Klasse behandelt werden. Ich habe ein Beispiel gewählt, das schön zeigt, wie die Vorgangsweise in der Physik ist und welche Rolle dabei die Mathematik spielt und wie beides Hand in Hand geht. Zuerst das Modellieren des Problems, worauf danach die Mathematisierung folgt. Die Lösung der Gleichungen ist dann wieder eine rein „innermathematische“ Aufgabe, wobei aber letztlich immer physikalische Überlegungen miteinbezogen werden müssen, um eventuelle physikalisch unsinnige Lösungen auszuschließen. Mit diesem Beispiel ist aber wohl die Grenze der Schulmathematik und Schulphysik endgültig erreicht oder überschritten und eine weitere Vertiefung ist einem Universitätsstudium vorbehalten. Zumindest kann sich die SchülerIn durch dieses letzte, mathematisch sicher anspruchsvolle Beispiel ein Bild von der Physik machen, auch dahingehend, was sie bei der Wahl eines Studiums der Physik erwarten würde, wobei natürlich ein Studium mathematisch noch weit anspruchsvoller sein würde.

Ich möchte noch anmerken, dass dieses Beispiel sehr ausführlich behandelt wird, da auch gezeigt werden soll, dass mit den Mitteln der Schulmathematik, wenn diese wirklich ausgereizt werden, durchaus interessante Probleme analysiert und behandelt werden können. Ich könnte mir durchaus vorstellen, dass ein Beispiel dieser Art mit dem hier gezeigten theoretischen Hintergrund in Kombination mit entsprechenden Experimenten sehr gut auch als Thema für ein Wahlpflichtfach herangezogen werden könnte.

Beispiel: Oberflächenform einer Flüssigkeit im rotierenden Eimer

Das letzte Beispiel geht weder von einer bekannten physikalischen Differentialgleichung aus, die es zu lösen gilt noch wird ein physikalisches Gesetz wie das Zerfallsgesetz des radioaktiven Zerfalls hergeleitet. Vielmehr liegt ein Problem vor, das es in gewisser Weise zu modellieren gilt, wobei in diesem Fall eine Differentialgleichung zur Lösung führen wird¹⁶⁸.

Das Problem, das es zu lösen gilt, ist die Bestimmung der Form der Oberfläche einer Flüssigkeit, die in einem Eimer rotiert. Es soll angenommen werden, dass der Eimer und damit die darin enthaltene Flüssigkeit mit einer konstanten Winkelgeschwindigkeit ω um eine Achse, die die Symmetrieachse des Eimers sein soll, rotiert.

Die Oberfläche einer Flüssigkeit nimmt immer die Form an, sodass jeder Punkt der Oberfläche normal auf die auf diesen Punkt wirkende Gesamtkraft ist.¹⁶⁹

Für eine waagrechte Oberfläche ergibt sich folgender vektorieller Zusammenhang (siehe Abbildung 89):

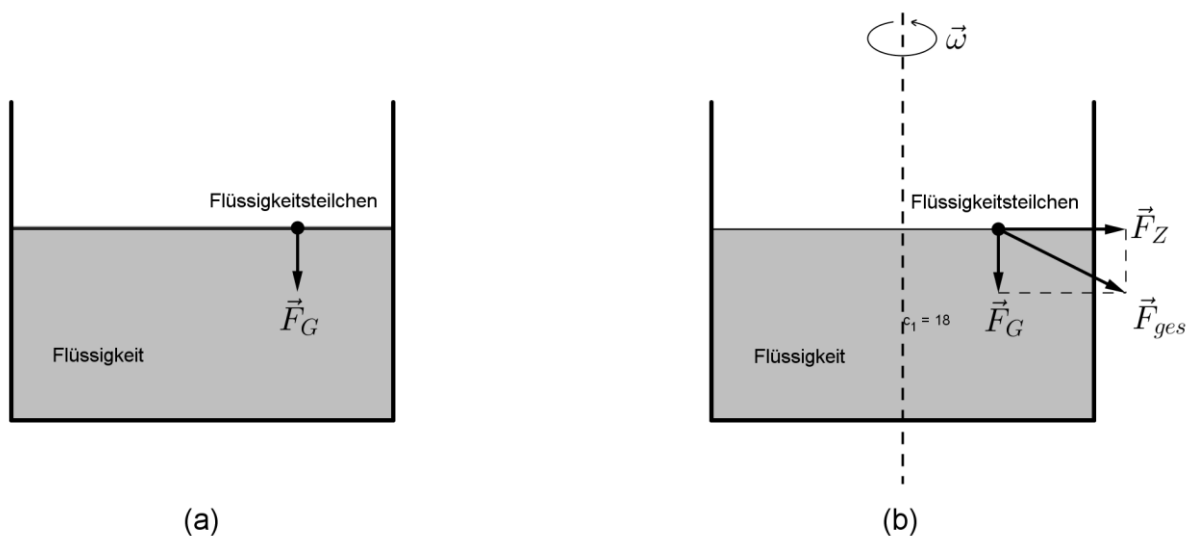


Abbildung 89: Kräfte auf eine waagrechte Flüssigkeitsoberfläche

¹⁶⁸ Das ist eine in der Praxis der Physik durchaus gängige Vorgangsweise. Man leitet eine das Problem beschreibende Gleichung her, in vielen Fällen ist das eine Differentialgleichung, und löst dann diese, was oft sehr viel mathematisches Geschick erfordert. Aus dieser allgemeinen Lösung ist dann durch die Einbeziehung der Anfangs- oder Randbedingungen die Lösung dem speziellen Problem anzupassen.

¹⁶⁹ Da die Schwerkraft senkrecht nach unten gerichtet ist und die Zentrifugalkraft immer radial nach außen, liegen die Kräfte immer in einer Ebene durch die Rotationsachse, wenn diese senkrecht steht, das heißt genauer in Richtung Erdmittelpunkt weist. Da der rotierende Eimer zylindersymmetrisch ist, kann man das Problem in dieser Ebene, in der die Kräfte und die Rotationsachse liegen, behandeln. Die räumliche Oberflächenform der Flüssigkeit ergibt sich dann durch Rotation der erhaltenen zweidimensionalen Kurve um die Rotationsachse.

In Abbildung 89a befindet sich der Eimer in Ruhe ($\vec{\omega} = \vec{0}$). Die Schwerkraft $\vec{F}_G$ ist die einzige vorkommende Kraft, sie zeigt senkrecht nach unten und steht somit normal auf die Flüssigkeitsoberfläche. Die Oberfläche ist eine waagrechte ebene Fläche, insbesondere, da hier durch eine fehlende Kraftkomponente parallel, also tangential zur Oberfläche, keine Strömung tangential zur Oberfläche stattfindet, die das Teilchen an der Oberfläche verschieben könnte.

Anders ist die Situation, wenn der Eimer nun rotiert. Die Beschleunigungsphase, die der Eimer und damit die Flüssigkeit ausführt, wenn die Winkelgeschwindigkeit von Null auf einen konstanten Wert ω erhöht wird, wird hier außer Acht gelassen. Jedes Flüssigkeitsteilchen erfährt nun eine zusätzliche Kraft nach außen, die Zentrifugalkraft $\vec{F}_Z$, die sich mit der Gewichtskraft $\vec{F}_G$ vektoriell zu einer resultierenden Gesamtkraft $\vec{F}_{ges}$ summiert. Diese steht nun nicht mehr normal auf die Oberfläche, so dass sich eine Tangentialkomponente zur Oberfläche ergibt. Diese bewirkt eine Strömung, was zur Folge hat, dass die Oberfläche keine ebene Fläche mehr ist, sondern eine andere Form annimmt. Diese Form nun gilt es zu bestimmen.

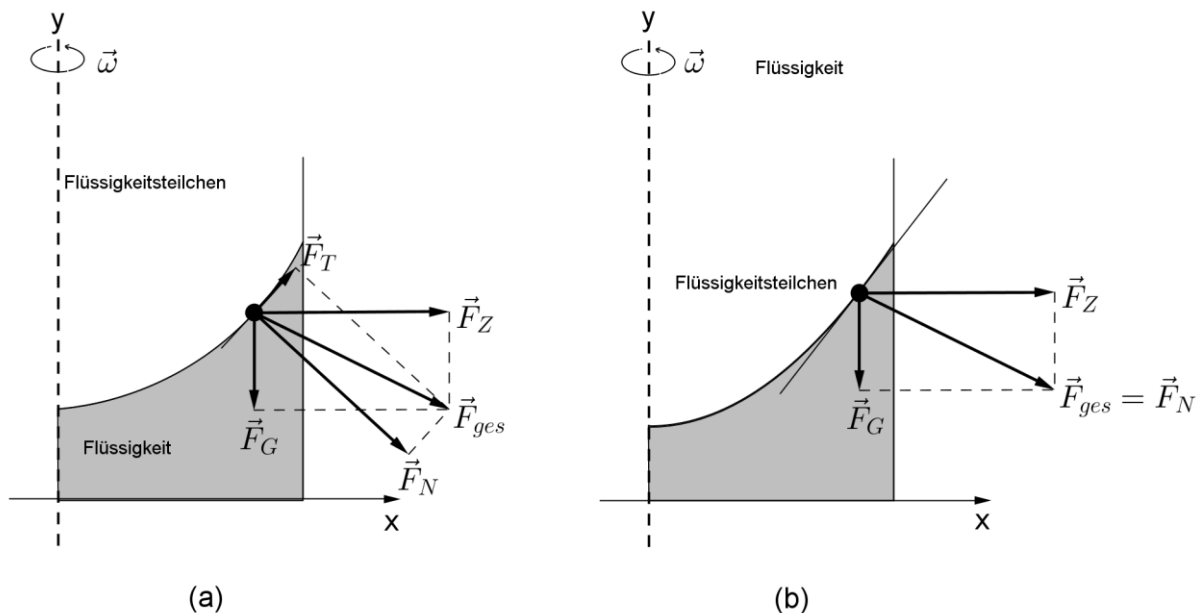


Abbildung 90: Kräfte auf eine rotierende Flüssigkeitsoberfläche

In Abbildung 90a ist ein Zustand gezeichnet, an dem die Flüssigkeit sich aus der Waagrechten entfernt hat. Man sieht, dass der resultierende Kraftvektor $\vec{F}_{\text{ges}}$ noch nicht normal auf die Flüssigkeitsoberfläche steht. Man kann $\vec{F}_{\text{ges}}$ in eine Kraftkomponente $\vec{F}_N$ normal zur Oberfläche und eine Komponente $\vec{F}_T$ tangential zur Flüssigkeitsoberfläche zerlegen. Die tangentielle Komponente bewirkt eine Strömung der Flüssigkeit in Richtung dieses Vektors, sodass sich die Oberflächenform weiter verändert, bis die Tangentialkraft verschwindet.

In Abbildung 90b ist der Zustand gezeichnet, an dem schließlich der resultierende Kraftvektor $\vec{F}_{\text{ges}}$ normal auf die Flüssigkeitsoberfläche steht. Es gibt jetzt keine tangentielle Kraftkomponente mehr, die eine Strömung der Flüssigkeit bewirken würde, nur mehr eine Normalkomponente $\vec{F}_N$. Diese Komponente zeigt ins Innere der Flüssigkeit und erzeugt einen Druck, der von den Wänden des Eimers aufgenommen wird und somit verhindert, dass sich die Flüssigkeitsoberfläche nach innen bewegt. Dieser Zustand ist der stationäre Zustand, der jetzt allein aus geometrischen Verhältnissen bestimmt werden kann.

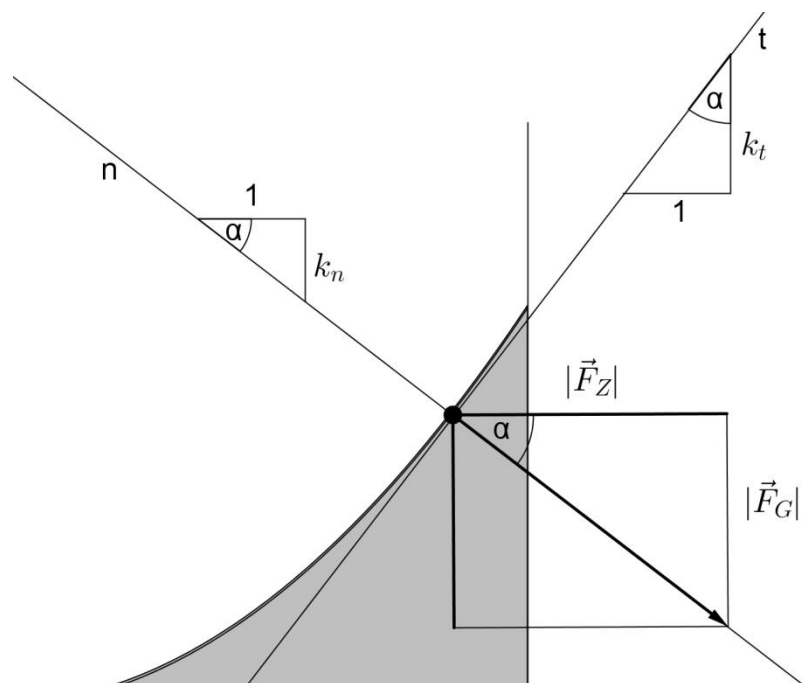


Abbildung 91: Verhältnisse zur Bestimmung des stationären Zustands der Oberfläche

Abbildung 91 zeigt noch einmal die geometrischen Verhältnisse, wobei auch die beiden Steigungsdreiecke der Tangente und Normalen eingezeichnet sind.

Man erkennt unmittelbar, dass

$$\tan \alpha = -\frac{|\vec{F}_G|}{|\vec{F}_Z|} = k_n ,$$

wobei k_n die Steigung der Normalen auf die Flüssigkeitsoberfläche ist.

Für die Steigung der Tangente ist

$$k_t = -\frac{1}{k_n} = \frac{|\vec{F}_Z|}{|\vec{F}_G|}$$

Für die beiden wirkenden Kräfte, die Gewichtskraft und die Zentrifugalkraft, gilt

$$\vec{F}_G = -m g \cdot \hat{e}_y$$

und

$$\vec{F}_Z = m \omega^2 x \cdot \hat{e}_x .$$

Daraus ergibt sich für die Steigung der Tangente

$$k_t = \frac{|\vec{F}_Z|}{|\vec{F}_G|} = \frac{\omega^2 x}{g}$$

Es ist nun die 1. Ableitung einer Funktion gleich der Steigung der Funktion, also folgt daraus die Differentialgleichung für die Kurve

$$\boxed{y'(x) = \frac{\omega^2 x}{g}} \quad (1)$$

Die Lösung ergibt sich durch einfache Integration beider Seiten zu

$$y(x) = \frac{\omega^2 x^2}{2g} + c \quad (2)^{170}$$

mit einer noch zu bestimmenden Integrationskonstanten c . Man erkennt, dass die Form der Kurve die einer Parabel ist, räumlich gesehen ist die Flüssigkeitsoberfläche ein Paraboloid.

Da hier keine zeitabhängige Funktion vorliegt, sind jetzt keine Anfangsbedingungen vorgegeben, sondern stattdessen sogenannte Randbedingungen, also Bedingungen, die das System geometrisch einschränken beziehungsweise festlegen. Es sind grundsätzlich mehrere verschiedene Randbedingungen möglich, am einfachsten und sinnvollsten ist es, das Volumen der Flüssigkeit zu nehmen, da dieses konstant bleibt.¹⁷¹ Es ist zu diesem Zweck das Volumen des entstehenden Raumes auszurechnen, den die Flüssigkeit bei der Rotation einnimmt und daraus dann die Integrationskonstante zu bestimmen, denn obwohl sich die Form des Flüssigkeitsvolumens ändert, bleibt natürlich das Volumen V , den die Flüssigkeit bei der Nullrotation eingenommen hat, auch weiterhin konstant.

Am einfachsten für die Volumsberechnung nutzt man die Zylindersymmetrie und nimmt infinitesimale Zylinderschalen, die aufintegriert werden. Eine Zylinderschale mit dem Radius x hat das infinitesimale Volumenelement $2\pi x \cdot f(x) \cdot dx$. Somit ergibt sich für das Volumen

$$V = \int_0^R 2\pi x \cdot y(x) dx = 2\pi \cdot \int_0^R \frac{\omega^2 x^3}{2g} + c \cdot x dx = 2\pi \cdot \left(\frac{\omega^2 x^4}{8g} + \frac{c \cdot x^2}{2} \right) \Big|_0^R = \frac{\omega^2 \pi R^4}{4g} + c \cdot \pi R^2$$

Daraus ergibt sich für die Integrationskonstante

$$c = \frac{V}{\pi R^2} - \frac{\omega^2 R^2}{4g}$$

Somit lautet das Ergebnis für die Kurve der Flüssigkeitsoberfläche

¹⁷⁰ Die Aufstellung und Lösung dieser Differentialgleichung steht in vielen Standardlehrbüchern der Physik, die Anpassung an die Randbedingungen ist dagegen selten dargelegt.

¹⁷¹ Flüssigkeiten sind in sehr guter Näherung inkompressibel, behalten ihr Volumen also bei. Sollte allerdings die Winkelgeschwindigkeit so groß werden, dass ein Teil der Flüssigkeit „überschwappt“, so wäre das Volumen der Flüssigkeit nicht mehr konstant und es müssten die Randbedingungen angepasst werden. Es wird hier aber angenommen, dass der Eimer immer so hoch ist, dass dieser Fall des Überschwapens nicht auftritt.

$$y(x) = \frac{\omega^2 x^2}{2g} + \frac{V}{\pi R^2} - \frac{\omega^2 R^2}{4g}$$

Damit ist aber das Problem noch nicht vollständig gelöst, denn die Flüssigkeitshöhe an der Rotationsachse kann Null werden, dann sind die Grenzen der Integration nicht mehr gültig. Es ist

$$y(0) = \frac{V}{\pi R^2} - \frac{\omega^2 R^2}{4g}$$

Dieser Ausdruck wird Null, wenn die Bedingung $\pi \omega^2 R^4 = 4gV$ erfüllt ist, anders gesagt, die erhaltene Gleichung für die Oberfläche ist nur dann anwendbar, wenn $\pi \omega^2 R^4 \leq 4gV$ ist.

$$y(x) = \frac{\omega^2}{4g} (2x^2 - R^2) + \frac{V}{\pi R^2} \quad , \text{ falls } \pi \omega^2 R^4 \leq 4gV \quad (3)$$

Die Grenzwinkelgeschwindigkeit bei konstantem Radius des Zylinders und konstantem Volumen der Flüssigkeit ergibt sich aus der Bedingung von (3) indem man statt der Ungleichung die Grenzgleichung anschreibt, also $\pi \omega^2 R^4 = 4gV$, zu

$$\omega^2 = \frac{4gV}{\pi R^4} = \frac{4gH}{R^2}$$

(hier wurde für das Volumen ohne Rotation, das einem Zylinder mit der Höhe H entspricht, $V = R^2 \pi H$ eingesetzt). Also ergibt sich für die Grenzwinkelgeschwindigkeit, bei der die Flüssigkeitshöhe an der Rotationsachse die Höhe null erreicht

$$\omega_{\text{Grenz}} = \frac{\sqrt{4gH}}{R}$$

Diese Ergebnisse sollen noch grafisch veranschaulicht werden. Es soll ein Eimer, der Zylinderform hat, rotieren, wobei die Winkelgeschwindigkeit erhöht wird, der Eimer aber seine Abmessungen beibehält. Der Eimer soll auch so hoch sein, dass das Wasser nie überschwappen kann, da sich ansonsten natürlich die Randbedingung, dass das

Volumen des Wassers konstant bleibt, verändern würde und die Gleichung an diese neue Randbedingung angepasst werden müsste.

Als Zahlenwerte soll der Eimer einen Radius $R = 1\text{LE}$ haben und das Wasser soll im nichtrotierenden Eimer $H = 1\text{LE}$ hoch stehen. Damit ergibt sich das Volumen des Wassers im Eimer zu $V = R^2 \pi H L E^3 = \pi L E^3$.

Damit ergibt sich als Bedingung $\pi \omega^2 R^4 \leq 4gV$

$$\omega^2 \leq 4g$$

$$\text{Also } \omega \leq \sqrt{4g} \approx 6,26\text{s}^{-1}$$

Für alle $\omega \leq 6,26\text{s}^{-1}$ gilt

$$y(x) = \frac{\omega^2}{4g} \cdot (2x^2 - 1) + 1$$

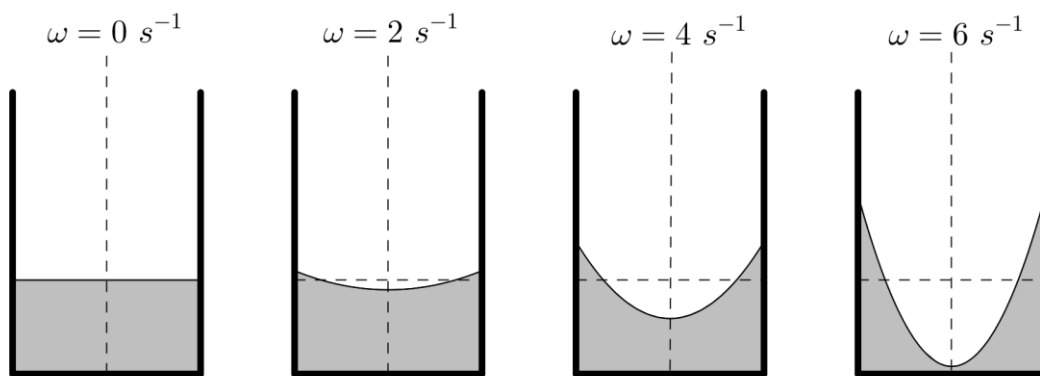


Abbildung 92: Rotation unterhalb der Grenzwinkelgeschwindigkeit

In der Abbildung 92 ist dies für die Winkelgeschwindigkeiten $\omega = 0\text{s}^{-1}$ (das entspricht keiner Rotation), $\omega = 2\text{s}^{-1}$, $\omega = 4\text{s}^{-1}$, $\omega = 6\text{s}^{-1}$ veranschaulicht. $\omega = 6\text{s}^{-1}$ liegt sehr nahe bei der Grenzwinkelgeschwindigkeit $\omega = 6,26\text{s}^{-1}$, so dass bei $x = 0$ die Flüssigkeitshöhe fast null wird.

Ein interessantes Ergebnis kann man erhalten, wenn man die Frage stellt, wie hoch bei einem rotierenden Zylinder die Flüssigkeit am Zylinderrand steht, wenn der Zylinder gerade mit der Grenzwinkelgeschwindigkeit rotiert, bei der die Flüssigkeitshöhe am Ort der Rotationsachse Null ist.

Für die Flüssigkeitshöhe am Rand gilt

$$y(R) = \frac{\omega^2 R^2}{2g} + \frac{V}{\pi R^2} - \frac{\omega^2 R^2}{4g} = \frac{\omega^2 R^2}{4g} + \frac{V}{\pi R^2}$$

und die Grenzwinkelgeschwindigkeit ω_{Grenz} kennen wir bereits zu

$$\omega_{\text{Grenz}}^2 = \frac{4gH}{R^2}$$

und damit

$$y(R) = \frac{4gH}{R^2} \cdot \frac{R^2}{4g} + \frac{R^2 \pi H}{\pi R^2} = 2 \cdot H$$

Das Ergebnis ist somit, dass bei einem rotierenden Zylinder die Flüssigkeitshöhe am Zylinderrand genau doppelt so hoch ist wie die ursprüngliche Flüssigkeitshöhe, wenn der Zylinder nicht rotiert.

Wir wollen uns noch den 2. Fall, wenn die Bedingung $\pi \omega^2 R^4 \geq 4gV$ erfüllt ist, ansehen. Dann müssen in der Integration zur Bestimmung der Integrationskonstanten die Grenzen anders bestimmt werden. Die untere Grenze ist dann die Entfernung von der Achse, an der die Flüssigkeitshöhe gleich Null ist. Aus der ursprünglichen Lösung der Differentialgleichung (2) sind die Stellen x_0 zu bestimmen, an denen $y(x_0) = 0$ ist.

$$y(x) = \frac{\omega^2 x^2}{2g} + c = 0$$

Daraus folgt

$$x_0 = \sqrt{-\frac{2g \cdot c}{\omega^2}}$$

Hier geht die Integrationskonstante c als Teil der unteren Integrationsgrenze mit ein.

Es muss das Integral $\int_{x_0}^R \left(\frac{\omega^2 x^2}{2g} + c \right) \cdot 2\pi x \, dx$ gleich dem Volumen V der Flüssigkeit

sein.

Die Auswertung des Integrals ergibt

$$\begin{aligned}
 \int_{x_0}^R \left(\frac{\omega^2 x^2}{2g} + c \right) \cdot 2\pi x \, dx &= \int_{x_0}^R \frac{\pi \omega^2 x^3}{g} + 2\pi c x \, dx \\
 &= \frac{\pi \omega^2 x^4}{4g} + \pi c x^2 \Big|_{x_0}^R \\
 &= \frac{\pi \omega^2 R^4}{4g} + \pi c R^2 - \frac{\pi \omega^2 x_0^4}{4g} - \pi c x_0^2
 \end{aligned}$$

Hier kann jetzt der Ausdruck für das oben erhaltene x_0 eingesetzt und das Ergebnis gleich V gesetzt werden.

$$\begin{aligned}
 \frac{\pi \omega^2 R^4}{4g} + \pi c R^2 - \frac{\pi \omega^2}{4g} \cdot \left(-\frac{2g \cdot c}{\omega^2} \right)^2 - \pi c \cdot \left(-\frac{2g \cdot c}{\omega^2} \right) &= V \\
 \frac{\pi \omega^2 R^4}{4g} + \pi c R^2 - \frac{g \pi c^2}{\omega^2} + \frac{2\pi g \cdot c^2}{\omega^2} &= V \\
 \frac{g \pi}{\omega^2} \cdot c^2 + \pi R^2 \cdot c + \frac{\pi \omega^2 R^4}{4g} - V &= 0
 \end{aligned}$$

Das ist eine quadratische Gleichung für c , die gelöst werden muss. Es ergibt sich zuerst durch Umformung

$$c^2 + \frac{\omega^2 R^2}{g} \cdot c + \frac{\omega^4 R^4}{4g^2} - \frac{V \omega^2}{g \pi} = 0$$

mit den Lösungen

$$\begin{aligned}
 c_{1,2} &= -\frac{\omega^2 R^2}{2g} \pm \sqrt{\left(\frac{\omega^2 R^2}{2g} \right)^2 - \frac{\omega^4 R^4}{4g^2} + \frac{V \omega^2}{g \pi}} \\
 &= -\frac{\omega^2 R^2}{2g} \pm \sqrt{\frac{\omega^4 R^4}{4g^2} - \frac{\omega^4 R^4}{4g^2} + \frac{V \omega^2}{g \pi}} \\
 &= -\frac{\omega^2 R^2}{2g} \pm \sqrt{\frac{V \omega^2}{g \pi}}
 \end{aligned}$$

Hier kann die Lösung mit dem negativen Vorzeichen vor dem Wurzel Ausdruck ausgeschlossen werden, da sie mit den Randbedingungen nicht kompatibel ist. Somit ist die Integrationskonstante bestimmt zu

$$c = -\frac{\omega^2 R^2}{2g} + \sqrt{\frac{V\omega^2}{g\pi}}$$

und die Lösung lässt sich schreiben als

$$y(x) = \frac{\omega^2}{2g} \cdot (x^2 - R^2) + \sqrt{\frac{V\omega^2}{g\pi}} \quad , \text{ falls } \pi\omega^2 R^4 \geq 4gV \quad (4)$$

Dass die zweite Lösung der quadratischen Gleichung ausgeschlossen werden muss, kann man folgendermaßen erkennen:

$$c_2 = -\frac{\omega^2 R^2}{2g} - \sqrt{\frac{V\omega^2}{g\pi}} \text{ ist sicher kleiner als } -\frac{\omega^2 R^2}{2g} .$$

Die Nullstelle x_0 für den Wert $c = -\frac{\omega^2 R^2}{2g}$ ergibt $x_0 = \sqrt{-\frac{2g}{\omega^2} \cdot \left(-\frac{\omega^2 R^2}{2g}\right)} = R$, also

genau den Radius des Zylinders. Da c_2 kleiner ist, ergäbe sich ein Wert für die Nullstelle, der größer als der Radius des Zylinders ist. Das ist von den Randbedingungen nicht erlaubt, also ist diese Lösung zu verwerfen.

Setzen wir in (4) für ω die Grenzwinkelgeschwindigkeit ein, so ergibt sich

$$\begin{aligned} y(x) &= \frac{\omega_{\text{Grenz}}^2}{2g} \cdot (x^2 - R^2) + \sqrt{\frac{R^2 \omega_{\text{Grenz}}^2 H}{g}} \\ &= \frac{x^2}{2g} \cdot \frac{4gH}{R^2} - \frac{R^2}{2g} \cdot \frac{4gH}{R^2} + \sqrt{\frac{R^2 H}{g} \cdot \frac{4gH}{R^2}} \\ &= \frac{2H}{R^2} \cdot x^2 - 2H + 2H \\ &= \frac{2H}{R^2} \cdot x^2 \end{aligned}$$

Machen wir dasselbe für (3), so erhalten wir

$$\begin{aligned}
y(x) &= \frac{\omega_{\text{Grenz}}^2 x^2}{2g} - \frac{\omega_{\text{Grenz}}^2 R^2}{4g} + \frac{R^2 \pi H}{\pi R^2} \\
&= \frac{x^2}{2g} \cdot \frac{4gH}{R^2} - \frac{R^2}{4g} \cdot \frac{4gH}{R^2} + \frac{R^2 \pi H}{\pi R^2} \\
&= \frac{2x^2 H}{R^2}
\end{aligned}$$

Man sieht, dass die beiden Lösungen (3) und (4) für diesen Grenzfall identisch sind, was natürlich auch der Fall sein muss!

Leiten wir noch die Abstände her, bei denen die Flüssigkeitshöhe Null ist, das entspricht

dem hergeleiteten $x_0 = \sqrt{-\frac{2g \cdot c}{\omega^2}}$

$$x_0 = \sqrt{-\frac{2g}{\omega^2} \cdot \left(-\frac{\omega^2 R^2}{2g} + \sqrt{\frac{V \omega^2}{g \pi}} \right)} = \sqrt{R^2 - \frac{2g}{\omega^2} \cdot \sqrt{\frac{V \omega^2}{g \pi}}} = \sqrt{R^2 - \sqrt{\frac{4gV}{\omega^2 \pi}}}$$

Damit ergibt sich für den Grenzwert $\omega \rightarrow \infty$ der Abstand, wenn die Winkelgeschwindigkeit immer mehr anwächst

$$x_0 = \lim_{\omega \rightarrow \infty} \sqrt{R^2 - \sqrt{\frac{4gV}{\omega^2 \pi}}} = R,$$

was auch erwartet wird, das heißt, die Flüssigkeit wird immer mehr an die Wand des rotierenden Behälters gepresst.

Diese Ergebnisse sollen ebenfalls noch grafisch veranschaulicht werden. Es seien dieselben Maße für den rotierenden Zylinder angenommen, wie beim Fall 1, also $R = 1\text{LE}$, $H = 1\text{LE}$ und damit das Volumen des Wassers im Eimer zu $V = R^2 \pi H \text{LE}^3 = \pi \text{LE}^3$.

Als Bedingung $\pi \omega^2 R^4 \geq 4gV$ ergibt sich

$$\omega^2 \geq 4g$$

Also $\omega \geq \sqrt{4g} \approx 6,26 \text{ s}^{-1}$

Für alle $\omega \geq 6,26 \text{ s}^{-1}$ gilt

$$y(x) = \frac{\omega^2}{2g} \cdot (x^2 - 1) + \sqrt{\frac{\omega^2}{g}} \quad \text{mit den „Nullstellen“} \quad x_0 = \sqrt{1 - \frac{2}{\omega} \sqrt{g}}.$$

Für die Höhe der Flüssigkeit am Rand erhält man allgemein aus (4)

$$y(R) = \sqrt{\frac{V\omega^2}{g\pi}} = \sqrt{\frac{R^2 H \omega^2}{g}} = R\omega \cdot \sqrt{\frac{H}{g}}$$

Und speziell für $y(1) = \frac{\omega}{\sqrt{g}}$

So ergibt sich etwa für $\omega = 15 \text{ s}^{-1}$ eine Randhöhe von $y(1) \approx 4,8 \text{ LE}$, also das knapp fünffache der ursprünglichen Wasserhöhe.

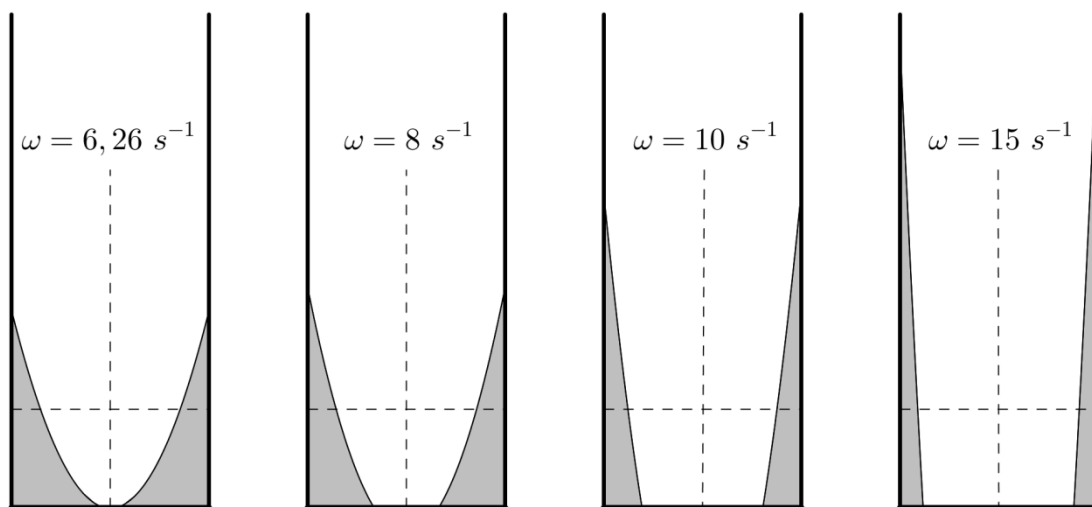


Abbildung 93: Rotation oberhalb der Grenzwinkelgeschwindigkeit

In der Abbildung 93 ist dies für die Winkelgeschwindigkeiten $\omega = 6,26 \text{ s}^{-1}$ (das entspricht der Grenzwinkelgeschwindigkeit), $\omega = 8 \text{ s}^{-1}$, $\omega = 10 \text{ s}^{-1}$ und $\omega = 15 \text{ s}^{-1}$ veranschaulicht.

Was an diesem Beispiel klar herausgekommen sein sollte, ist, dass sehr oft die Lösung der Differentialgleichung, die einem Problem zugrunde liegt, nicht die Hauptschwierigkeit ausmacht, sondern die Anpassung der Lösung an das spezielle Problem, das heißt die Randbedingungen beziehungsweise Anfangsbedingungen. Ein Problem ist auch oft die Koordinatenwahl. In dem hier gezeigten Beispiel sind Zylinderkoordinaten die geeigneten Koordinaten, um das Problem zu beschreiben. Vereinfacht wurde das Beispiel beträchtlich, weil ein Zylinder rotiert, somit auch die Randbedingungen leicht in Zylinderkoordinaten beschrieben werden können. Es könnte aber durchaus auch etwa ein Quader rotieren, der in kartesischen Rechteckkoordinaten angegeben ist. Die Lösungen (3) und (4) gelten natürlich für den Quader genauso, nur würde aus dem Paraboloid ein Rechteck (von oben gesehen) herausgestanzt werden, auf dem dann die Randbedingungen festgelegt werden.

Als Schlussbemerkung sei noch erwähnt, dass dieses Modell einer rotierenden Flüssigkeit durchaus auch von technischem Interesse ist. Man stellt etwa parabelförmige Spiegel her, indem man eine Grundform rotieren lässt und eine flüssige Substanz darin rotieren und erstarren lässt. Eine weitere geniale Variante ist, dass man die Flüssigkeit nicht einmal mehr erstarren lässt, sondern direkt als rotierenden Spiegel verwendet. Es existiert auch ein solches Teleskop mit einem Spiegel mit 6 Meter Durchmesser, der aus rotierendem Quecksilber besteht.¹⁷² Die Vorteile sind, dass die Spiegeloberfläche nicht mehr von Unebenheiten der Trägeroberfläche abhängt und somit die parabolische Spiegelform sehr genau und einfach erhalten werden kann. Der Nachteil ist, dass die Rotationsachse genau zum Massenmittelpunkt der Erde zeigen muss, somit das Teleskop nicht mehr schwenkbar ist. Es können nur Objekte beobachtet werden, die sich im Zenit befinden, wobei durch die jahreszeitliche Veränderung der Richtung der Erdachse zwar immer noch eine größere Himmelsregion beobachtbar ist, der Ausschnitt aber nicht mehr frei wählbar ist.

¹⁷² Das Large Zenit Teleskope der University of British Columbia; die Kosten sind nur etwa 10% der Kosten eines Spiegels aus Glas; siehe deren Website, www.astro.ubc.ca/LMT/lzt/, [Stand 11.1.2015]

Quellenverzeichnis

- [1] BIFIE. 2014. *Mathematik*. <https://www.bifie.at/node/80> [Stand: 14.11.2014]
- [2] BIFIE. *Haupttermin 2013/14 – Mathematik (AHS)*. <https://www.bifie.at/node/2633> [Stand: 14.11.2014]
- [3] BMBF. *Lehrplan Mathematik Unterstufe*.
https://www.bmbf.gv.at/schulen/unterricht/lp/ahs14_789.pdf?4dzgm2 [Stand: 14.11.2014]
- [4] BMBF. *Lehrplan Mathematik Unterstufe*.
https://www.bmbf.gv.at/schulen/unterricht/lp/lp_neu_ahs_07_11859.pdf?4dzgm2 [Stand: 14.11.2014]
- [5] BMBF. *Lehrplan Physik Unterstufe*.
https://www.bmbf.gv.at/schulen/unterricht/lp/ahs16_791.pdf?4dzgm2 [Stand: 14.11.2014]
- [6] BMBF. *Lehrplan Physik Oberstufe*.
https://www.bmbf.gv.at/schulen/unterricht/lp/lp_neu_ahs_10_11862.pdf?4dzgm2 [Stand: 14.11.2014]
- [7] *Das ist Mathematik 1* von Reichel, Humenberger (Hrsg.), Litschauer, Groß, Aue
1. Auflage 2011; Schulbuchnummer 135311
- [8] *Das ist Mathematik 2* von Reichel, Humenberger (Hrsg.), Litschauer, Groß, Aue
1. Auflage 2011; Schulbuchnummer 140399
- [9] *Das ist Mathematik 3* von Reichel, Humenberger (Hrsg.), Litschauer, Groß, Aue,
Neuwirth; 1. Auflage 2012; Schulbuchnummer 145076
- [10] *Das ist Mathematik 4* von Reichel, Humenberger (Hrsg.), Litschauer, Groß, Aue,
Neuwirth; 1. Auflage 2010; Schulbuchnummer 150100
- [11] *Mathematik verstehen 5* von Malle, Woschitz, Koth, Salzger;
1. Auflage 2010; Schulbuchnummer 150344

- [12] *Mathematik verstehen 6* von Malle, Woschitz, Koth, Salzger;
1. Auflage 2010; Schulbuchnummer 150616
- [13] *Mathematik verstehen 7* von Malle, Woschitz, Koth, Salzger;
1. Auflage 2011; Schulbuchnummer 155145
- [14] *Mathematik verstehen 8* von Malle, Woschitz, Koth, Salzger;
1. Auflage 2012; Schulbuchnummer 160187
- [15] *Physik 2* von Gollenz, Breyer, Tentschert, Reichel;
1. Auflage 2012; Schulbuchnummer 160186
- [16] *Physik 3* von Gollenz, Breyer, Tentschert, Reichel;
1. Auflage 2012; Schulbuchnummer 160173
- [17] *Physik 4* von Gollenz, Breyer, Tentschert, Reichel;
1. Auflage 2013; Schulbuchnummer 165242
- [18] *Big Bang 5* von Martin Apolin;
1. Auflage 2007; Schulbuchnummer 135061
- [19] *Big Bang 6* von Martin Apolin;
1. Auflage 2008; Schulbuchnummer 135062
- [20] *Big Bang 7* von Martin Apolin;
1. Auflage 2008; Schulbuchnummer 140177
- [21] *Big Bang 8* von Martin Apolin;
1. Auflage 2008; Schulbuchnummer 140424
- [22] *Physik compact Basiswissen 5 RG* von Jaros, Nussbaumer, Nussbauer, Kunze;
1. Auflage 2004; Schulbuchnummer 115924
- [23] *Mathematik 7* von Götz, Reichel Hrsg.), Müller, Hanisch;
1. Auflage 2011, Schulbuchnummer 155152
- [24] *Einführung in die Astronomie und Astrophysik* von Arnold Hanslmeier;
2. Auflage 2007; Spektrum akademischer Verlag
- [25] *Physik für Wissenschaftler und Ingenieure* von Paul A. Tipler und Gene Mosca;
6. Auflage 2009; Spektrum akademischer Verlag

- [26] Faszination Physik 1+2 von Bruno Putz;
4. Auflage 2007; Veritas Verlag
- [27] Mathematik 8 von Götz, Reichel Hrsg.), Müller, Hanisch;
1. Auflage 2013, Schulbuchnummer 160197
- [28] Physik, Lehr- und Übungsbuch von Douglas C. Giancoli;
3., erweiterte Auflage, Pearson Studium
- [29] Prisma Physik 2
1. Auflage 2006; Schulbuchnummer 140352

Abbildungsverzeichnis

Alle in dieser Diplomarbeit vorkommenden Abbildungen wurden vom Autor der Arbeit selbsttätig mit dem allgemein zugänglichen Programm GeoGebra erstellt.

Abbildung 1: Zeitliche Koordination von Mathematik und Physik	10
Abbildung 2: Ein Eisberg in Wasser	40
Abbildung 3: v-t-Diagramm eines gleichförmig bewegten Körpers.....	44
Abbildung 4: v-t-Diagramm des freien Falls	45
Abbildung 5: Spannungsarbeit einer gespannten Feder.....	46
Abbildung 6: Zweites Kepler'sches Gesetz	47
Abbildung 7: Flussüberquerung	56
Abbildung 8: Geschwindigkeiten vor und nach dem Stoß beim Billardspiel	57
Abbildung 9: Billardspiel, Zusammenhang der Geschwindigkeiten.....	57
Abbildung 10: s-t-Diagramm eines ruhenden Körpers	60
Abbildung 11: s-t-Diagramm eines Körpers mit gleichbleibenden Geschwindigkeiten	61
Abbildung 12: Vergleich zweier gleichförmiger Bewegungen	63
Abbildung 13: Eine nicht-gleichförmige Bewegung	65
Abbildung 14: s-t-Diagramm des freien Falls	66
Abbildung 15: Ein s-t-Diagramm mit einer Geschichte	66
Abbildung 16: v-t-Diagramm zu Abbildung 15	68
Abbildung 17: Regressionsgerade	75
Abbildung 18: Ebene Polarkoordinaten	78
Abbildung 19: Räumliche Polarkoordinaten oder Kugelkoordinaten	80
Abbildung 20: Zylinderkoordinaten	81
Abbildung 21: Das Zustandekommen der Sternparallaxe	84
Abbildung 22: Das Parsec als astronomische Längeneinheit.....	86
Abbildung 23: Stöße bei gleichen Massen	90
Abbildung 24: Eine doppelt so schwere Kugel stößt eine ruhende leichte Kugel	91
Abbildung 25: Eine halb so schwere Kugel stößt eine ruhende schwere Kugel	92

Abbildung 26: Eine leichte und eine doppelt so schwere Kugel bewegen sich aufeinander zu	92
Abbildung 27: Stoß gegen eine Wand.....	93
Abbildung 28: Senkrechter Wurf.....	96
Abbildung 29: Der schiefe Wurf, Vektorzerlegung der Geschwindigkeiten	98
Abbildung 30: Kräftezerlegung bei einem Pendel.....	101
Abbildung 31: Kräftezerlegung für das Federpendel	103
Abbildung 32: Näherung für kleine Winkel φ	103
Abbildung 33: Zerfallsgesetz für die Altersbestimmung von „Ötzi“	110
Abbildung 34: Periodendauer, Amplitude und Phase zweier harmonischen Schwingungen	114
Abbildung 35: Superposition einer Schwingung	115
Abbildung 36: Überlagerung zweier Sinusschwingungen mit verschiedener Phase	118
Abbildung 37: Konstruktive Interferenz	119
Abbildung 38: Destruktive Interferenz.....	120
Abbildung 39: Eine Schwebung.....	122
Abbildung 40: Drehmoment zweier Massen 1.....	123
Abbildung 41: Drehmoment zweier Massen 2.....	124
Abbildung 42: Überhang eines Buches	125
Abbildung 43: Überhang zweier Bücher	126
Abbildung 44: Überhang dreier Bücher	126
Abbildung 45: Zur Herleitung der Winkelgeschwindigkeit	131
Abbildung 46: Drehgeschwindigkeitsvektor	132
Abbildung 47: Der Bahndrehimpulsvektor $\vec{v}$	133
Abbildung 48: Der Drehmomentvektor $\vec{M}$	134
Abbildung 49: Impulsübertrag auf die Behälterwände.....	141
Abbildung 50: Zur Herleitung der Bernoulli-Gleichung.....	142
Abbildung 51: Reflexion eines Lichtstrahls	150
Abbildung 52: Zur Herleitung des Reflexionsgesetzes	151
Abbildung 53: Bestimmung der Extremstelle anhand des Funktionsverlaufs	154
Abbildung 54: Zur Herleitung des Brechungsgesetzes.....	156
Abbildung 55: Drehbewegung einer Punktmasse.....	164
Abbildung 56: Geschwindigkeit und Beschleunigung einer Drehbewegung	166

Abbildung 57: Eine Aufgabe zum Halley'schen Kometen	167
Abbildung 58: Die Ellipse	168
Abbildung 59: Die Bahnkurve des Halley'schen Kometen	169
Abbildung 60: Ein geometrischer Beweis zur Ellipse	170
Abbildung 61: Ein Ellipsenspiegel.....	171
Abbildung 62: Ein geometrischer Beweis zur Hyperbel	172
Abbildung 63: Ein Hyperbelspiegel	174
Abbildung 64: Ein geometrischer Beweis zur Parabel	175
Abbildung 65: Ein Parabelspiegel.....	176
Abbildung 66: Das Cassegrain-Teleskop.....	177
Abbildung 67: Sinusschwingung und Zeigerdiagramm	179
Abbildung 68: Komplexe Darstellung des Zeigermodells.....	181
Abbildung 69: Ein Widerstand in einem Wechselstromkreis.....	183
Abbildung 70: Phasenbeziehung beim Widerstand in einem Wechselstromkreis	184
Abbildung 71: Eine Spule in einem Wechselstromkreis.....	185
Abbildung 72: Phasenbeziehung an einer Spule in einem Wechselstromkreis	186
Abbildung 73: Ein Kondensator in einem Wechselstromkreis.....	187
Abbildung 74: Phasenbeziehung an einem Kondensator in einem Wechselstromkreis	188
Abbildung 75: Komplexe Widerstände am ohmschen Widerstand, einer Spule und einer Kapazität.....	189
Abbildung 76: Wechselstromkreise mit zwei Bauteilen in Serie	190
Abbildung 77: Impedanzen im Serienwechselstromkreis mit zwei Bauteilen	191
Abbildung 78: Ein Serienschwingkreis.....	191
Abbildung 79: Die komplexen Widerstände beim Serienschwingkreis	192
Abbildung 80: Massenmittelpunkt zweier Punktmassen	198
Abbildung 81: Massenmittelpunkt eines Dreiecks.....	200
Abbildung 82: Trägheitsmoment eines rotierenden Stabes	202
Abbildung 83: Trägheitsmoment einer Schwungscheibe	203
Abbildung 84: Momentan- und mittlere Wechselstromleistung.....	206
Abbildung 85: Der schiefe Wurf zur Zeit $t = 0$	212
Abbildung 86: Wurfbahnen mit verschiedenen Abschusswinkeln	215
Abbildung 87: Wurfbahnen mit verschiedenen Anfangsgeschwindigkeiten.....	216

Abbildung 88: Der radioaktive Zerfall als Zufallsprozess	217
Abbildung 89: Kräfte auf eine waagrechte Flüssigkeitsoberfläche.....	222
Abbildung 90: Kräfte auf eine rotierende Flüssigkeitsoberfläche	223
Abbildung 91: Verhältnisse zur Bestimmung des stationären Zustands der Oberfläche	224
Abbildung 92: Rotation unterhalb der Grenzwinkelgeschwindigkeit	228
Abbildung 93: Rotation oberhalb der Grenzwinkelgeschwindigkeit	233

Tabellenverzeichnis

Tabelle 1: Ergebnisse und Auswertung des Versuchs zur Geschwindigkeitsbestimmung	25
Tabelle 2: Die 7 SI-Basisgrößen und SI-Basiseinheiten	29
Tabelle 3: Vorsilben	29
Tabelle 4: Abgeleitete physikalische Größen	30
Tabelle 5: Vollständige Messwertetabelle des Versuchs zur Geschwindigkeitsbestimmung .	36
Tabelle 6: Gegenüberstellung Translation - Rotation	136
Tabelle 7: Einige Koeffizienten a und b der Van-der-Waals-Gleichung	148

Anhang A: Die Reifeprüfung vom 9. Mai 2014

Es sind nur diejenigen Aufgaben angefügt, die einen naturwissenschaftlichen Bezug haben:

Typ-1-Aufgaben:

- Aufgabe 2: Punktladungen
- Aufgabe 7: Zerfallsprozess
- Aufgabe 15: Nikotin
- Aufgabe 18: Pflanzenwachstum

Typ-2-Aufgaben:

- Aufgabe 2: Zustandsgleichung idealer Gase
- Aufgabe 3: Chemische Reaktionsgeschwindigkeit
- Aufgabe 5: Sportwagen

Name:	
Klasse:	



Standardisierte kompetenzorientierte
schriftliche Reifeprüfung

Mathematik

9. Mai 2014

Teil-1-Aufgaben



Aufgabe 2

Punktladungen

Der Betrag F der Kraft zwischen zwei Punktladungen q_1 und q_2 im Abstand r wird beschrieben durch die Gleichung $F = C \cdot \frac{q_1 \cdot q_2}{r^2}$ (C ... physikalische Konstante).

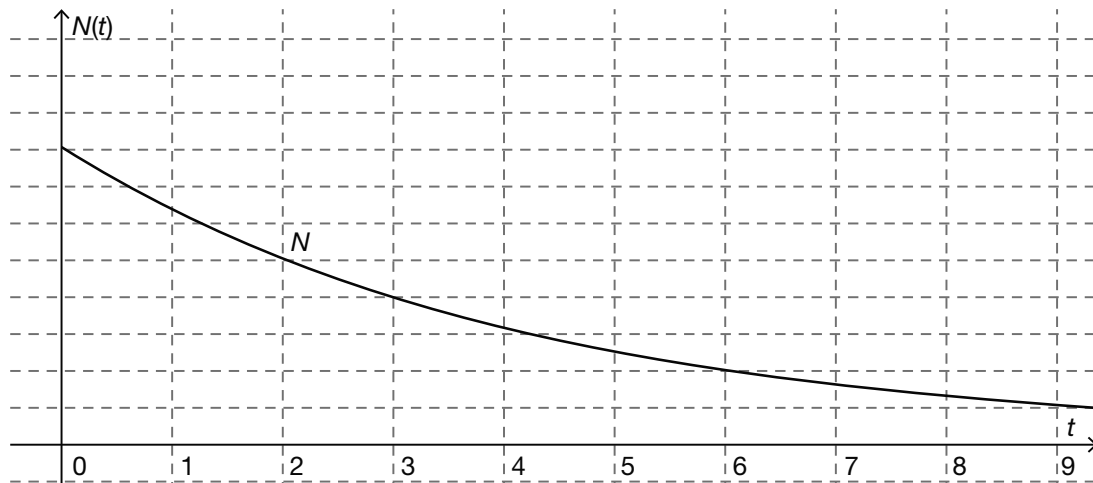
Aufgabenstellung:

Geben Sie an, um welchen Faktor sich der Betrag F der Kraft ändert, wenn der Betrag der Punktladungen q_1 und q_2 jeweils verdoppelt und der Abstand r zwischen diesen beiden Punktladungen halbiert wird!

Aufgabe 7

Zerfallsprozess

Der unten abgebildete Graph einer Funktion N stellt einen exponentiellen Zerfallsprozess dar; dabei bezeichnet t die Zeit und $N(t)$ die zum Zeitpunkt t vorhandene Menge des zerfallenden Stoffes. Für die zum Zeitpunkt $t = 0$ vorhandene Menge gilt: $N(0) = 800$.



Mit t_H ist diejenige Zeitspanne gemeint, nach deren Ablauf die ursprüngliche Menge des zerfallenden Stoffes auf die Hälfte gesunken ist.

Aufgabenstellung:

Kreuzen Sie die beiden zutreffenden Aussagen an!

$t_H = 6$	☐
$t_H = 2$	☐
$t_H = 3$	☐
$N(t_H) = 400$	☐
$N(t_H) = 500$	☐

Aufgabe 15

Nikotin

Die Nikotinmenge x (in mg) im Blut eines bestimmten Rauchers kann modellhaft durch die Differenzengleichung $x_{n+1} = 0,98 \cdot x_n + 0,03$ (n in Tagen) beschrieben werden.

Aufgabenstellung:

Geben Sie an, wie viel Milligramm Nikotin täglich zugeführt werden und wie viel Prozent der im Körper vorhandenen Nikotinmenge täglich abgebaut werden!

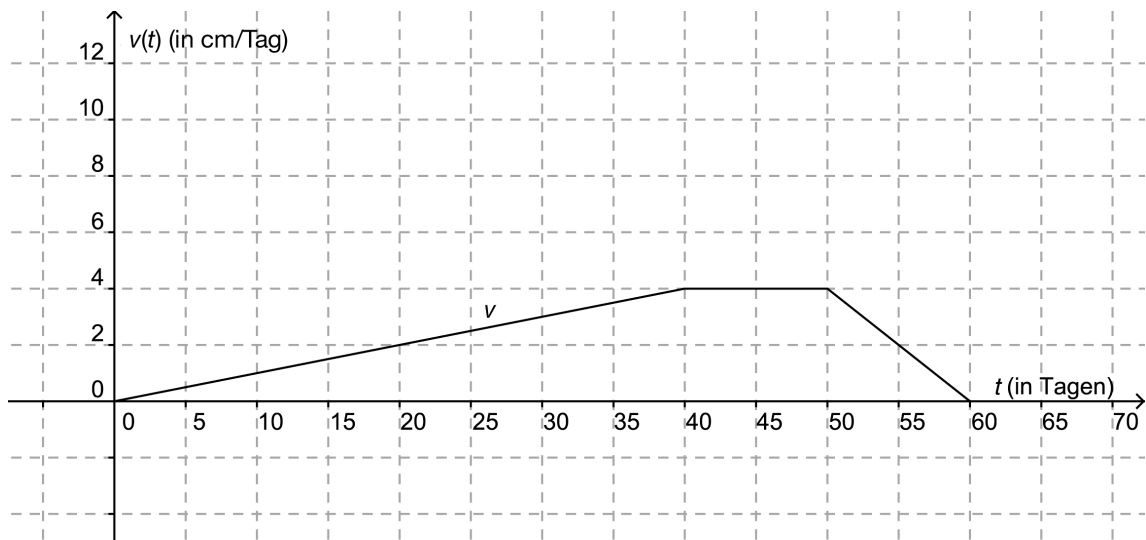
_____ mg

_____ %

Aufgabe 18

Pflanzenwachstum

Die unten stehende Abbildung beschreibt näherungsweise das Wachstum einer schnellwüchsigen Pflanze. Sie zeigt die Wachstumsgeschwindigkeit v in Abhängigkeit von der Zeit t während eines Zeitraums von 60 Tagen.



Aufgabenstellung:

Geben Sie an, um wie viel cm die Pflanze in diesem Zeitraum insgesamt gewachsen ist!

Name:	
Klasse:	



Standardisierte kompetenzorientierte
schriftliche Reifeprüfung

Mathematik

9. Mai 2014

Teil-2-Aufgaben



Aufgabe 2

Zustandsgleichung idealer Gase

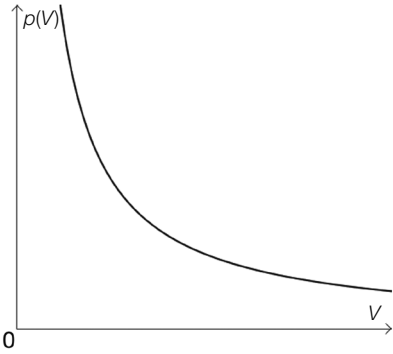
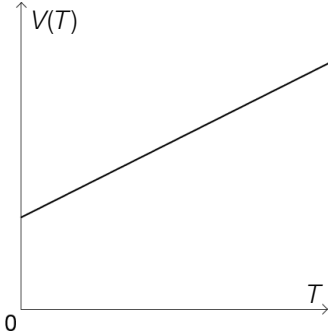
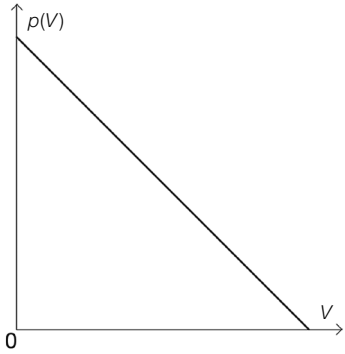
Die Formel $p \cdot V = n \cdot R \cdot T$ beschreibt modellhaft den Zusammenhang zwischen dem Druck p , dem Volumen V , der Stoffmenge n und der absoluten Temperatur T eines idealen Gases und wird *thermische Zustandsgleichung idealer Gase* genannt. R ist eine Konstante.

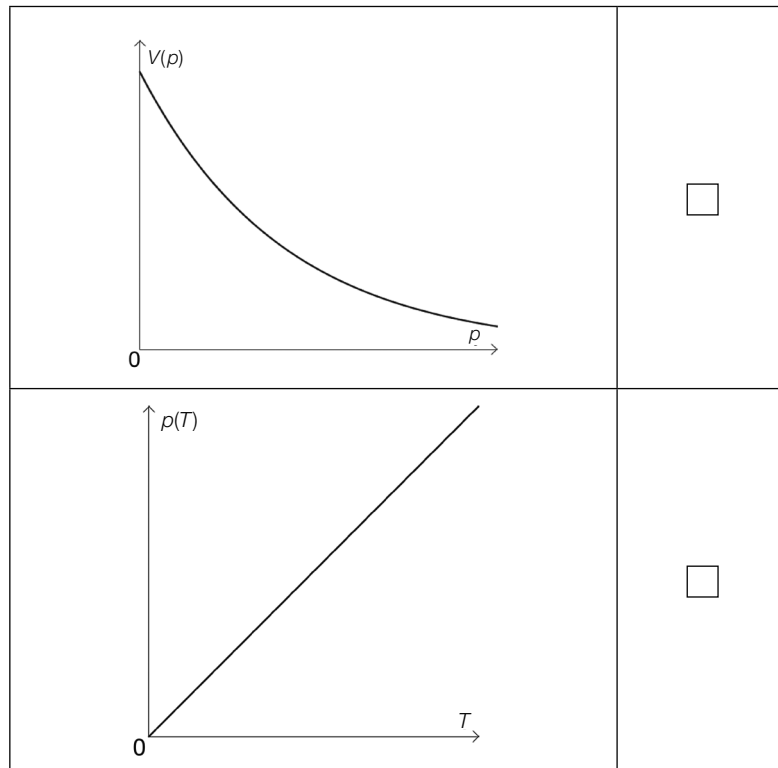
Das Gas befindet sich in einem geschlossenen Gefäß, in dem die Zustandsgrößen p , V und T verändert werden können. Die Stoffmenge n bleibt konstant.

Aufgabenstellung:

- a) Führen Sie alle Möglichkeiten an, die zu einer Verdopplung des Drucks führen, wenn jeweils eine der Zustandsgrößen verändert wird und die anderen Größen konstant bleiben!

Genau zwei der folgenden Graphen stellen die Abhängigkeit zweier Zustandsgrößen gemäß dem oben genannten Zusammenhang richtig dar. Kreuzen Sie diese beiden Graphen an! Beachten Sie: Die im Diagramm nicht angeführten Größen sind jeweils konstant.

	☐
	☐
	☐



- b) Bei gleichbleibender Stoffmenge und gleichbleibender Temperatur kann das Volumen des Gases durch Änderung des Drucks variiert werden.

Begründen Sie, warum die *mittlere* Änderung des Drucks in Abhängigkeit vom Volumen

$$\frac{p(V_2) - p(V_1)}{V_2 - V_1}$$

für jedes Intervall $[V_1; V_2]$ mit $V_1 \neq V_2$ ein negatives Ergebnis liefert!

Ermitteln Sie jene Funktionsgleichung, die die *momentane* Änderung des Druckes in Abhängigkeit vom Volumen des Gases beschreibt!

Aufgabe 3

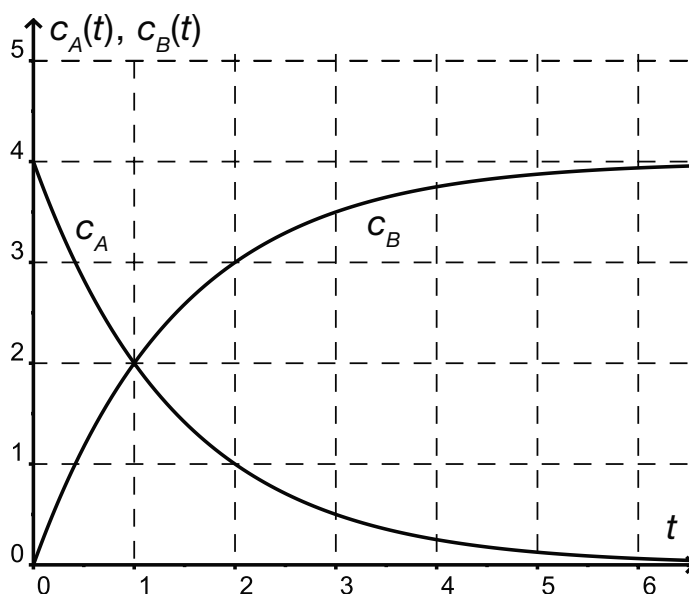
Chemische Reaktionsgeschwindigkeit

Die Reaktionsgleichung $A \rightarrow B + D$ beschreibt, dass ein Ausgangsstoff A zu den Endstoffen B und D reagiert, wobei aus einem Molekül des Stoffes A jeweils ein Molekül der Stoffe B und D gebildet wird.

Die Konzentration eines chemischen Stoffes in einer Lösung wird in Mol pro Liter (mol/L) angegeben. Die Geschwindigkeit einer chemischen Reaktion ist als Konzentrationsänderung eines Stoffes pro Zeiteinheit definiert.

Die unten stehende Abbildung zeigt den Konzentrationsverlauf der Stoffe A und B bei der gegebenen chemischen Reaktion in Abhängigkeit von der Zeit t .

$c_A(t)$ beschreibt die Konzentration des Stoffes A , $c_B(t)$ die Konzentration des Stoffes B . Die Zeit t wird in Minuten angegeben.



Aufgabenstellung:

- a) ☐ A Ermitteln Sie anhand der Abbildung die durchschnittliche Reaktionsgeschwindigkeit des Stoffes B im Zeitintervall $[1; 3]$!

Für die gegebene Reaktion gilt die Gleichung $c_A'(t) = -c_B'(t)$. Interpretieren Sie diese Gleichung im Hinblick auf den Reaktionsverlauf!

- b) Bei der gegebenen Reaktion kann die Konzentration $c_A(t)$ des Stoffes A in Abhängigkeit von der Zeit t durch eine Funktion mit der Gleichung $c_A(t) = c_0 \cdot e^{k \cdot t}$ beschrieben werden.

Geben Sie die Bedeutung der Konstante c_0 an!

Argumentieren Sie anhand des Verlaufs des Graphen von c_A , ob der Parameter k positiv oder negativ ist!

Leiten Sie eine Formel für jene Zeit τ her, nach der sich die Konzentration des Ausgangsstoffes halbiert hat! Geben Sie auch den entsprechenden Ansatz an!

Aufgabe 5

Sportwagen

Ein Sportwagen wird von 0 m/s auf 28 m/s (≈ 100 km/h) in ca. 4 Sekunden beschleunigt. $v(t)$ beschreibt die Geschwindigkeit in Metern/Sekunde während des Beschleunigungsvorganges in Abhängigkeit von der Zeit t in Sekunden. Die Geschwindigkeit lässt sich durch die Funktionsgleichung $v(t) = -0,5t^3 + 3,75t^2$ angeben.

Aufgabenstellung:

- a) ☐ A Geben Sie die Funktionsgleichung zur Berechnung der momentanen Beschleunigung $a(t)$ zum Zeitpunkt t an!
Berechnen Sie die momentane Beschleunigung zum Zeitpunkt $t = 2$!
- b) Geben Sie einen Ausdruck zur Berechnung des in den ersten 4 Sekunden zurückgelegten Weges an! Ermitteln Sie diesen Weg $s(4)$ (in Metern)!
- c) Angenommen, dieser Sportwagen beschleunigt – anders als ursprünglich angegeben – gleichmäßig in 4 Sekunden von 0 m/s auf 28 m/s. Nun wird mit $v_1(t)$ die Geschwindigkeit des Sportwagens nach t Sekunden bezeichnet.

Geben Sie an, welcher funktionale Zusammenhang zwischen v_1 und t vorliegt!
Ermitteln Sie die Funktionsgleichung für v_1 !

Anhang B: Abstract

Im Zuge der neuen kompetenzorientierten Reifeprüfung wird fächerübergreifendes Unterrichten immer mehr gefordert, dies im Speziellen natürlich in Hinblick auf die beiden Unterrichtsfächer Mathematik und Physik. Darüber hinaus sind die Mathematik und die Physik natürlich in ihrer geschichtlichen Entwicklung eng miteinander verflochten und beeinflussen sich seit jeher gegenseitig. Physik ohne Mathematik ist undenkbar und genauso gab die Physik der Mathematik immer wieder neue Impulse, woraus teilweise ganze Teilgebiete der Mathematik befruchtet beziehungsweise erst entstanden sind, man denke an die Vektorrechnung oder die Tensorrechnung. Dies berücksichtigend sollte natürlich auch die Mathematik in der Schule in den Physikunterricht einfließen und somit auch ein wichtiger Wesenszug der Physik, die ja heute sehr mathematisiert ist, den SchülerInnen nähergebracht werden. In dieser Arbeit soll dargelegt werden, inwieweit es möglich ist, die Schulmathematik gewinnbringend in den Physikunterricht einfließen zu lassen. Dies soll anhand von konkreten praxisnahen und direkt nutzbaren Beispielen geschehen und den gesamten Unterrichtszeitraum von der 1. Klasse bis zur 8. Klasse abdecken. Es soll aber auch insbesondere den MathematiklehrerInnen eine Hilfestellung in die Hand gegeben werden, um die Nützlichkeit der gelehrt mathematischen Methoden zu erkennen und im Unterricht Verwendung zu finden. Aber es sollen auch die PhysiklehrerInnen, die nicht Mathematik unterrichten, eine im Unterricht anwendbare Unterlage erhalten, um ausgewählte physikalische Problemstellungen mathematisch exakt im Unterricht lösen zu können.

Anhang C: Lebenslauf

Ausbildung:

1985: Matura am Bundesgymnasium Bregenz

Ab 1986: Studium der Technischen Physik an der TU-Wien

Ab 2009: Lehramtsstudium UF Physik und UF Mathematik an der Uni-Wien

Berufliche Tätigkeiten:

1990-1992: DATALINE-Datenverarbeitungsgesellschaft

1992-1993: WAVE Solutions Information Technology GmbH

1993: Präsenzdienst

1997-1998: SIEMENS

1998-1999: OBERKOFER & SCHWARZ GmbH

1999-2008: Selbstständig

Ab 2012: Lehrer für Mathematik und Physik am Evangelischen Gymnasium und Werkschulheim, 1110 Wien