

**Versuchsanleitung
zum
Experimentierkoffer
Teilchenfallen**

Inhaltsverzeichnis

1	Einleitung	4
2	Der Vortrag	5
3	Was ist eine Teilchenfalle?	6
3.1	Wozu braucht man eine Teilchenfalle?	6
4	Die Atomuhr - eine Anwendung	8
4.1	Theorie der Atomuhr	8
4.1.1	Energieniveaus	8
4.1.2	Linienbreite	10
4.2	Die Cäsiumuhr	13
4.3	Das Ion in der Falle - eine neue Uhr	14
4.4	Wozu braucht man so genaue Uhren?	15
4.5	Weitere Anwendungen von Teilchenfallen	16
5	Wie baut man eine Teilchenfalle?	17
5.1	Das Theorem von Earnshaw	18
6	Die Fallen im Koffer	19
6.1	Paulfalle	19
6.1.1	Theorie	19
6.1.2	Praxis	23
6.1.3	Der rotierende Sattel	25
6.2	Levitron	26
6.2.1	Theorie	26
6.2.2	Praxis	27
6.3	Diamagnetisches Schweben	28
6.3.1	Theorie	28
6.3.2	Praxis	30
6.4	Geregeltes Schweben	30
6.4.1	Theorie	30
6.4.2	Praxis	31
7	Internetlinks	33

1 Einleitung

Dieser Versuchskoffer entstand im Rahmen meiner Staatsexamensarbeit mit dem Titel „Konzeption und Aufbau einer mobilen Experimentiereinheit für Schülerpräsentationen zum Thema Teilchenfallen“ am 5. Physikalischen Institut der Universität Stuttgart. Bei der Konzeption wurde also hauptsächlich an Schüler gedacht, insbesondere an die im Lehrplan neu enthaltene „Gleichwertige Feststellung von Schülerleistungen“ (GFS).

Eine anderweitige Verwendung durch Lehrer und Schüler (Unterrichtseinheit, Vortrag, Physik-AG, Tag der offenen Tür, etc.) ist jedoch durchaus erwünscht. Solange dieser Koffer den Beteiligten die Physik näher bringt, erfüllt er meines Erachtens seinen Zweck.

Nachdem in Kapitel 3 kurz erklärt wird, was genau eine Teilchenfalle ist, und warum Physiker so gerne Teilchen einfangen, wird in Kapitel 4 mit der Ionenuhr eine wichtigste Anwendungsmöglichkeit besprochen, wobei natürlich ein wenig Theorie nicht zu vermeiden ist. Wer sich eher für Teilchenfallen, als für deren Anwendungen interessiert, kann dieses Kapitel auch überblättern. Es eignet sich jedoch durchaus als Thema für einen eigenständigen GFS-Vortrag.

Bevor im Kapitel 6 die im Koffer enthaltenen Fallen in Theorie und Praxis erläutert werden, widmet sich das Kapitel 5 dem wohl größten Problem beim Bau einer Teilchenfalle: der Unmöglichkeit des freien, stabilen Schwebens in statischen elektrischen und magnetischen Feldern.

Zu guter Letzt soll in Kapitel 7 mit den Links zu jedem Thema die hoffentlich geweckte Experimentier- und Bastellust befriedigt werden.

Ich wünsche allen, die mit diesem Koffer arbeiten, mindestens so viel Freude, wie mir seine Zusammenstellung bereitet hat.

Tobias Koch

Für Anmerkungen, Verbesserungsvorschläge und dergleichen bin ich sehr dankbar. Bitte an: t.koch@physik.uni-stuttgart.de

P.S. Wenn an den Experimenten etwas kaputt geht, dann sagt mir das bitte. Die meisten Sachen sind schnell repariert oder ersetzt. Eure Nachfolger werden es euch danken.

2 Der Vortrag

Um den GFS-Vortrag sowohl für die Vortragenden (in Form einer guten Note) als auch für die Zuhörer (in Form von Erkenntnisgewinn) erfolgreich zu gestalten, enthält dieses Kapitel einige Ratschläge, die meines Erachtens wesentlich für einen guten Vortrag sind.

- Versucht eurem Vortrag einen **roten Faden** zu geben, an dem sich die Zuhörer orientieren können. Dazu hilft es oft zu Beginn eine Übersicht aufzulegen (zum Beispiel auf dem Tageslichtprojektor), aus der hervorgeht, wie euer Vortrag strukturiert ist. An dieser Übersicht kann man sich dann während des Vortrags auch selbst „entlanghangeln“, und ab und zu kurz sagen, an welchem Punkt man sich gerade befindet.
- Überlegt euch, welche Teile der Versuchsanleitung für das Verständnis eures Vortrags sinnvoll sind und welche nicht. Es muss nicht alles, was in der Anleitung steht, auch im Vortrag vorkommen. Andererseits muss auch nicht alles, was im Vortrag vorkommt, in der Anleitung stehen; eine weitere Recherche, zum Beispiel im Internet, ist durchaus erwünscht.
- Redet bei eurem Vortrag langsam und erklärt auch scheinbar einfache Sachen (zum Beispiel Begriffe wie Potentialminimum). Eine kleine Skizze kann oft behilflich sein. Ihr dürft nicht vergessen, dass ihr euch, im Gegensatz zu euren Mitschülern, einige Stunden mit dem Thema beschäftigt habt. Zur Vorbereitung hilft es auch sehr, sich gegenseitig seinen Versuch zu erklären und „dumme“ Fragen zu stellen.
- Setzt nicht zwanghaft alle zur Verfügung stehenden Medien ein. Wenn ihr mit PC, Laptop und Beamer arbeitet, testet vorher, ob alle Komponenten kompatibel sind.

3 Was ist eine Teilchenfalle?

Eine Teilchenfalle ist ein Hilfsmittel des Experimentalphysikers, das ein Teilchen **frei** und **stabil** fängt. Etwas genauer ausgedrückt heißt das, ein Teilchen wird durch physikalische Wechselwirkungen, also Kräfte, auf möglichst kleinem Raum festgehalten, **ohne** es durch materielle Wände einzuschließen. Für die physikalische Grundlagenforschung werden dazu Elementarteilchen wie Elektronen, Protonen, Neutronen, Atome, Ionen oder Moleküle gefangen. Zur anschaulichen Erklärung der Funktionsweise von Teilchenfallen werden mit den Fallen im Koffer makroskopische Teilchen wie Staub oder noch größere „Teilchen“ gefangen.

3.1 Wozu braucht man eine Teilchenfalle?

Das Ziel der Physik ist das Verständnis der Vorgänge in der Natur und der Wechselwirkungen zwischen den beteiligten Partikeln. Auf dem Weg zu diesem Ziel arbeiten theoretische und experimentelle Physiker zusammen. Theorien werden ausgearbeitet und durch Messungen überprüft. Aufgrund neuer genauerer Messungen müssen Theorien dann verfeinert oder gar verworfen werden und neue Theorien erarbeitet werden. Im Rahmen dieses Entwicklungsprozesses muss natürlich mit immer größerer Genauigkeit gemessen werden. Insbesondere muss der Experimentator sicherstellen, dass in eine Messung nur Effekte eingehen, die die Theorie auch berücksichtigt. Externe Einflüsse auf das betrachtete physikalische System müssen eliminiert werden. (Ein einfaches Beispiel hierfür ist zum Beispiel die Kompensation des eine Messung beeinflussenden Erdmagnetfeldes.)

Um diese störenden Einflüsse (Kräfte) loszuwerden, ist eine Teilchenfalle das ideale Instrument: Bei einer „normalen“ Messung (zum Beispiel der Energieniveaus eines Atoms) wird, da kein einzelnes Atom zur Verfügung steht, immer an einer großen Anzahl von Atomen gemessen. Geschieht dies zum Beispiel in der Gasphase, so misst man bei einem Volumen von einem Kubikzentimeter, unter Normaldruck, an circa 10^{20} Atomen! Das hat zwei entscheidende Nachteile:

1. Die Atome sind nicht alle gleich. Sie haben zum Beispiel verschiedene Geschwindigkeiten und befinden sich an unterschiedlichen Orten. Hängt die zu messende Eigenschaft von Geschwindigkeit und Ort ab, so misst man also alle Werte, die mit den jeweiligen Geschwindigkeiten und Orten möglich sind, nicht die Eigenschaft eines Atoms.
2. Die Atome beeinflussen sich gegenseitig. Sie stoßen miteinander, verbinden sich eventuell zu einem Molekül, ionisieren sich gegenseitig...

Beides sind störende Effekte, die eine Theorie zwar berücksichtigen kann, aber nur statistisch. Das hat dann aber zur Folge, dass man eine statistische Verteilung der Messwerte erhält und nicht einen einzigen sehr genauen.

Beide Probleme lassen sich vermeiden, wenn man es schafft, ein einzelnes Atom in einer Falle zu isolieren. Es sitzt dann idealerweise in der Mitte der Falle (definierter Ort) und bewegt sich nicht mehr (Geschwindigkeit gleich null). Es gibt dann keine Störungen mehr, die von der Bewegung des Atoms oder von Wechselwirkungen zwischen mehreren Atomen herrühren.¹ Störungen dieses Systems (Atomkern + Elektronenhülle) können jetzt nur noch von äußeren Feldern oder Stößen mit dem Hintergrundgas herrühren. Kompensiert man die Felder und sorgt für ein möglichst gutes Vakuum, so hat man sich das Messobjekt Atom bestmöglich präpariert.

Außerdem verhält sich ein so präpariertes Atom in jedem Labor der Welt gleich. Das ist einer der Vorteile der im nächsten Kapitel vorgestellten Atomuhr.

¹Im realen Experiment schafft man es natürlich nie, das Teilchen ganz zur Ruhe zu bekommen. Dies ist eine direkte Folge der Heisenbergschen Unschärferelation $\Delta p \cdot \Delta x \geq \hbar$. Das Teilchen hat immer eine nicht verschwindende Nullpunktsenergie.

4 Die Atomuhr - eine Anwendung

Was ist eine Uhr? Das wichtigste an jeder Uhr ist immer ein periodischer Taktgeber, dessen Schwingungsperiode möglichst konstant sein sollte, da davon die Genauigkeit der Uhr abhängt. Im Laufe der Jahrtausende wurden sie deshalb immer weiter entwickelt. Der erste Taktgeber war derjenige, der das Leben der Menschen am meisten beeinflusst hat: die Rotation der Erde. Diese ist zwar in ihrer Periodendauer relativ konstant, aber diese Periodendauer ist auch sehr lang. Das macht es schwierig, kurze Zeiten zu messen, wie zum Beispiel die fünf Minuten bis ein Ei weichgekocht ist. Deshalb wurden bald schnellere Taktgeber benutzt: Ein langsam tropfendes Gefäß (Wasseruhr), Pendel, Schwingquartze... Die Genauigkeit einer Uhr hängt davon ab, wie genau die Frequenz des Taktgebers bestimmt ist. Außerdem darf die Frequenz natürlich nicht von der Zeit oder äußeren Einflüssen abhängen.

4.1 Theorie der Atomuhr

4.1.1 Energieniveaus

Ein Atom besteht aus einem positiv geladenen Kern (Protonen und Neutronen) und einer negativ geladenen Atomhülle, bestehend aus den Elektronen, die durch das vom Kern erzeugte elektrische Feld an diesen gebunden sind. Wenn man sich das Atom wie das Sonnensystem vorstellt (Die Elektronen umkreisen den Kern, wie die Planeten die Sonne.), dann sollte für die Elektronenbahnen prinzipiell jeder Radius möglich sein. Auf einer Bahn, die näher am Kern ist, muss das Elektron eben entsprechend schneller kreisen, damit es nicht in den Kern fällt (Zentrifugalkraft = Coulombkraft). Und weil jeder Radius einer gewissen Energie entspricht (Je näher ein Elektron am Kern sitzt, desto stärker ist es an ihn gebunden.), sollte auch jede beliebige Energie möglich sein. **Das ist aber nicht so!**

Richtig beschrieben wird die Physik so kleiner Systeme wie eines Atoms durch die Quantenmechanik und deren grundlegende Gleichung, die Schrödinger-Gleichung. Aus der Schrödinger-Gleichung erhält man das für uns wichtige Resultat, dass die Elektronen im Feld des Atomkerns **nicht** beliebige Energiewerte annehmen können. Es gibt nur bestimmte, diskrete Energien E_n , ($n = 1, 2, 3, \dots$), die von den sogenannten Quantenzahlen abhängen. Die Elektronen können aber zwischen diesen Energieniveaus hin- und herspringen. Damit dabei der Energieerhaltungssatz nicht verletzt wird, tauscht das Atom mit der Umgebung Energie aus, und zwar in Form eines Photons/Lichtquants. Wird ein Elektron in ein höheres Energieniveau gehoben, so absorbiert (=aufnehmen) das Atom ein Photon. Fällt ein Elektron

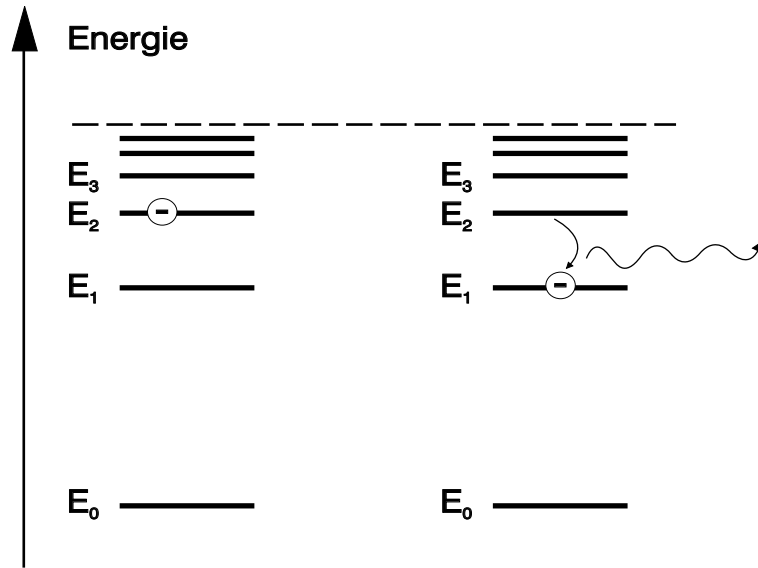


Abbildung 4.1: Die Energieniveaus eines Elektrons im Feld des Atomkerns. Beim Übergang in ein tieferes Niveau wird Energie in Form eines Photons frei. Beträgt die Energiedifferenz zwischen 2 und 5 eV, dann ist das ausgestrahlte Photon fürs menschliche Auge sichtbar.

Die gestrichelte Linie kennzeichnet die Ionisierungsenergie. Bekommt ein Elektron so viel Energie, dass es über diese Schwelle gehoben wird (zum Beispiel durch einen Stoß), so wird das Atom ionisiert.

in ein tieferes Energieniveau, emittiert (=aussenden) das Atom ein Photon. Die Frequenz f und die Wellenlänge λ dieses Photons ist dabei genau definiert. Es gilt

$$f = \frac{\Delta E}{h} = \frac{E_1 - E_2}{h} \quad \text{und} \quad \lambda = \frac{c}{f},$$

wobei $c = 2,998 \cdot 10^8 \frac{\text{m}}{\text{s}}$ die Lichtgeschwindigkeit, $h = 6,626 \cdot 10^{-34} \text{Js}$ das Plancksche Wirkumsquantum¹ und ΔE die Energiedifferenz zwischen Anfangszustand E_1 und Endzustand E_2 ist. Da die Energien E_n in jedem Atom durch die Schrödingergleichung exakt festgelegt sind, gibt es in jedem Atom auch nur bestimmte charakteristische Energiedifferenzen. Folglich kann eine Atomsorte auch nur Photonen bestimmter Wellenlängen/Frequenzen absorbieren, beziehungsweise emittieren. Das sind genau die Spektrallinien, die für jedes Atom charakteristisch sind.²

Und das ist auch schon der ganze Trick: Weil die Energieniveaus so exakt sind, ist auch die Frequenz eines Photons, das bei einem bestimmten Übergang entsteht, so exakt. Aus diesem Grund wurde 1967 die SI-Sekunde vom Internationalen Komitee für Maß und Gewicht neu definiert [9].

¹Sowohl c als auch h sind Naturkonstanten.

²Das Gesagte gilt natürlich nicht nur für Atome, sondern auch für Moleküle und Ionen.

”Die Sekunde ist das 9.192.631.770-fache der Periodendauer der dem Übergang zwischen den beiden Hyperfeinstrukturniveaus des Grundzustandes von Atomen des Nuklids ^{133}Cs entsprechenden Strahlung.”

Seit 1967 ist die Sekunde also über einen bestimmten Übergang im Cäsiumatom definiert³. Dass der heutigen Definition genau das 9.192.631.770-fache der Periodendauer zugrunde liegt, liegt einfach daran, dass man die „neue“ Sekunde der „alten“ möglichst genau angepasst hat. Diese war über die Erdrotation definiert, grob gesagt hat man einen Tag durch $24 \cdot 60 \cdot 60 = 86400$ geteilt, um auf eine Sekunde zu kommen. Diese Definition war jedoch nicht so genau, weil die Frequenz der Erdrotation aufgrund von Masseverschiebungen (Gezeiten, flüssiges Magma im Innern) nicht dauerhaft konstant ist.

Der große Vorteil einer Atomuhr besteht also in der dauerhaft konstanten und exakt definierten Frequenz ihres Taktgebers.

4.1.2 Linienbreite

Nach dem bisher Gesagten müsste es also genau eine Frequenz geben, die zu einem bestimmten Übergang gehört. Jedoch sind die Energieniveaus in der Realität immer etwas verbreitert, was dazu führt, dass auch die Übergangsfrequenz nicht mehr ganz exakt ist (Abbildung 4.2). Jede Linie hat somit eine gewisse endliche Linienbreite, die die Genauigkeit der Messung limitiert. Ein Maß für die Genauigkeit einer Schwingung ist die sogenannte Güte $Q = f/\Delta f$. Sie ist definiert als der Quotient aus der Frequenz der Schwingung f und der Genauigkeit Δf mit der diese bestimmbar ist. Je schmaler die Linie, desto exakter lässt sich deren Mitte bestimmen. Die Gründe für diese Linienbreite werden jetzt kurz erklärt:

Natürliche Linienbreite Jede Linie hat eine natürliche Linienbreite, die in der endlichen Lebensdauer eines Zustands begründet ist. Jeder energetische Zustand (außer dem Grundzustand) hat eine gewisse Lebensdauer τ , nach der er im Mittel zerfällt. Nach der Heisenbergschen Unschärferelation

$$\Delta E \cdot \Delta t \geq \hbar$$

ist eine Energie, in einem Zeitintervall Δt nur bis auf $\Delta E \geq \hbar/\Delta t$ genau bestimmbar. Und diese Zeit Δt , in der die Energie bestimmt werden kann ist natürlich durch die Lebensdauer des Zustands beschränkt. Die Genauigkeit, mit der die Energie und damit also auch die Frequenz bestimmt ist, hängt somit auch direkt mit der Zeit zusammen, in der das Messfeld (das eingestrahlte Lichtfeld) mit dem Messobjekt (dem Cäsiumatom) wechselwirken kann. Wie wir später sehen werden, ist auch das ein Problem, das mit Atomen in Fallen wegfällt. Durch Stöße mit dem Hintergrundgas ist die Lebensdauer eines Atoms in einer Teilchenfalle

³Die Sekunde ist seither die am genauesten definierte SI-Einheit. Sie ist so genau, dass man mittlerweile auch die Definition des Meters über die Sekunde festgelegt hat: Der Weg den Licht im Vakuum im 299.792.458ten Teil einer Sekunde zurücklegt, ist ein Meter.

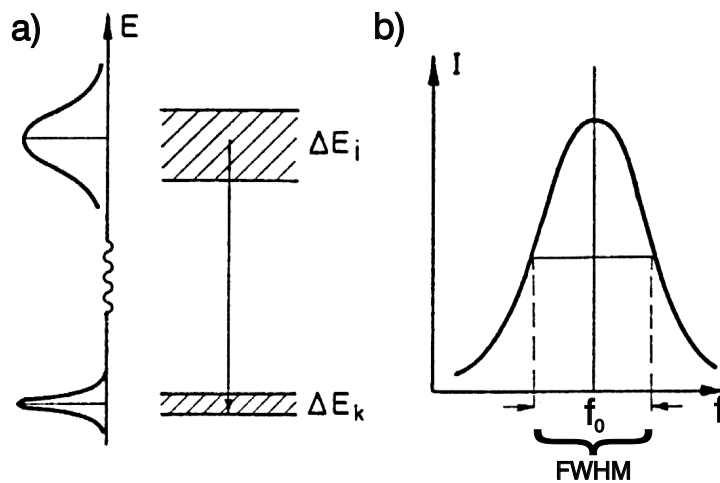


Abbildung 4.2: Weil die einzelnen Energieniveaus verbreitert sind (a), können auch Photonen mit größerer oder kleinerer Frequenz emittiert werden. Im Emissionsspektrum (b) sieht man dann eine um die Resonanzfrequenz f_0 verschmierte Linie. Je schmäler die beteiligten Energieniveaus sind, desto schmäler ist die Linie. Ein Maß für die Breite ist die volle Halbwertsbreite (engl: Full Width at Half Maximum), Grafik aus[11].

zwar auch begrenzt, aber in einem guten Vakuum trotzdem deutlich länger als die Lebensdauer des angeregten Zustands, was für eine Messung ausreicht.

Die natürliche Linienbreite ist eine intrinsische Eigenschaft des beobachteten Übergangs. Auch durch noch so geschicktes Messen wird die Linie nicht schmaler. Um eine scharfe Resonanz als Taktgeber für die Atomuhr zu bekommen, bleibt einzig und allein die Suche nach einem sogenannten Uhrenübergang mit möglichst langer Lebensdauer. Die Beobachtung der natürlichen Linienbreite ist jedoch meist nicht möglich, da sie durch andere um Größenordnungen größere Verbreiterungsmechanismen überdeckt wird.

Stoßverbreiterung Wenn man Atome in einem gewissen Volumen einsperrt, fliegen diese, abhängig von Druck und Temperatur, wild durcheinander. Dabei stoßen sie auch miteinander. Ist die mittlere Zeit zwischen zwei Stößen kleiner als die Lebensdauer des angeregten Zustands, so wird durch die Stöße die Messzeit verringert, was nach dem oben gesagten eine Verbreiterung der Linie mit sich bringt. Die Stoßverbreiterung lässt sich durch Absenken des Druckes verringern oder idealerweise durch eine Messung an einem einzelnen Teilchen, wie das in einer Teilchenfalle möglich ist, ganz eliminieren.

Dopplerverbreiterung Abhängig von der Temperatur bewegen sich die Teilchen in einem Gas in alle Raumrichtungen (Brownsche Molekularbewegung). Bei Raumtemperatur beträgt die mittlere Geschwindigkeit der Teilchen mehrere hundert

Meter pro Sekunde! Ganz analog zum akustischen Dopplereffekt (Der Ton einer Sirene eines Polizeiautos, das auf einen zukommt ist höher, als der eines weg-fahrenden.) gibt es dann auch den Dopplereffekt für Licht. Das heißt, ein Atom kann auch Licht mit größerer oder kleinerer Frequenz als der eigentlichen Resonanzfrequenz emittieren oder absorbieren. Zum Beispiel heißt das, dass ein sich bewegendes Atom ein entgegenkommendes Photon, dessen Frequenz für ein ruhendes Atom zu klein ist, trotzdem absorbieren kann. Das selbe gilt auch für die Emission von Photonen, wenn sich das Atom bewegt. Dadurch wird die beobachtete Linie breiter. Die Dopplerverbreiterung kann verringert werden indem man die mittlere Geschwindigkeit verringert, also das Atom kühlt. Nun kann man aber ein in einer Falle gefangenes Atom nicht einfach mit flüssigem Helium oder ähnlichem kühlen, denn das gefangene Atom soll ja die Wand nicht berühren!

Der Trick ist, dass man von allen Seiten Licht einstrahlt, dessen Frequenz unterhalb der Resonanzfrequenz liegt. Es kann also nur von Atomen absorbiert werden, die dem Lichtstrahl entgegenkommen. Dabei nimmt es den Impuls des Photons⁴ auf und wird ein wenig abgebremst. Wenn anschließend wieder ein Photon emittiert wird, bekommt das Atom zwar jedesmal wieder einen kleinen Rückstoß, da das aber statistisch verteilt in alle Raumrichtungen passiert bleibt als Nettoeffekt eine Abbremsung/Kühlung übrig⁵. Um eine gute Kühlleistung zu erzielen muss, das Atom möglichst viele Photonen pro Zeit absorbieren. Man benutzt deshalb zur Kühlung einen sehr kurzlebigen Übergang, den sogenannten Kühlübergang. Mit dieser Methode, man spricht auch von optischer Kühlung, sind Temperaturen unter einem Millikelvin ($= 0,001^\circ$) über dem absoluten Nullpunkt erreichbar.

Die beiden letzten Verbreiterungsmechanismen lassen sich weitgehend ausschalten, wenn man die Frequenz an einem einzelnen gekühlten Teilchen misst.

⁴Auch Photonen haben einen Impuls. Aus dem Impulserhaltungssatz folgt dann, dass das Atom langsamer wird, wenn es ein Photon absorbiert. (Inelastischer Stoß)

⁵Die absorbierten Photonen kommen alle aus der gleichen Richtung.

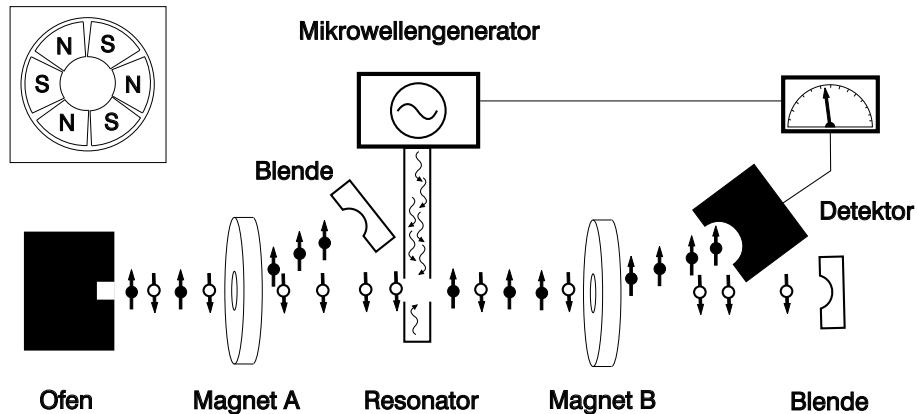


Abbildung 4.3: Das Schematische Aufbau einer Cäsium-Atomuhr. Die beiden Zustände sind hier durch $\uparrow$ und $\downarrow$ gekennzeichnet. Die Frequenz des Mikrowellengenerators wird mit dem Signal des Detektors gesteuert. Oben links: Um die beiden Zustände zu sortieren braucht man sogenannte Hexapolmagnete (griechisch: hexa = sechs).

4.2 Die Cäsiumuhr

Die Cäsiumuhr (Abbildung 4.3) besteht aus vier wichtigen Komponenten, deren Funktionsweise jetzt erläutert wird:

Der **Atomstrahlofen** dient zur Erzeugung eines Strahls von Cäsiumatomen. Das Cäsium (ein Metall, das bereits bei 28°C schmilzt) wird im Ofen verdampft, die Atome können dann durch ein kleines Loch austreten. Sie haben eine Geschwindigkeit von circa 200 m/s. Cäsium hat im Grundzustand zwei verschiedene sogenannte Hyperfeinstrukturniveaus (siehe Sekundendefinition in Abschnitt 4.1.1), die zwei verschiedenen Energien entsprechen. Im inhomogenen Magnetfeld von **Magnet A** werden die Atome verschieden stark abgelenkt. Durch eine Blende werden die Atome im energetisch höher liegenden Zustand aussortiert, nur die im tiefer liegenden gelangen in den **Mikrowellenresonator**. Hier wird jetzt die elektromagnetische Strahlung mit einer einstellbaren Frequenz von ca. 9192631770 Hz eingestrahlt. Durch Absorption der Mikrowellenstrahlung können die Atome vom tieferen in den höheren Zustand übergehen. Im **Magnet B** werden die Atome ein zweites Mal sortiert. Diesmal werden die Atome im energetisch tieferen Zustand aussortiert, die im höheren gelangen auf den **Detektor**, der zählt wieviele Atome ankommen. Ändert man jetzt die Frequenz der Strahlung, so ändert sich auch die Intensität am Detektor. Wenn das Signal maximal wird, hat man die Übergangsfrequenz optimal getroffen, denn es sind ja die meisten Photonen von Cäsiumatomen absorbiert worden. Die Frequenz wird dadurch stabilisiert, dass der Mikrowellengenerator an das Signal des Detektors gekoppelt wird. Die gesamte Anordnung befindet sich im Hochvakuum, um störende Effekte zu minimieren. Die Gangunsicherheit einer solchen Cäsiumuhr beträgt 10^{-14} , also eine Abwei-

chung von rund einer Sekunde in 3 Millionen Jahren.⁶

Mit einem anderen Aufbau, der sogenannten Cäsiumfontäne, kann die Genauigkeit auf eine Sekunde in 30 Millionen Jahren erhöht werden. Bei dieser Uhr wird eine lasergekühlte Wolke von Cäsiumatomen mit einem Laser senkrecht nach oben beschleunigt. Die Atomwolke beschreibt dann eine ballistische Flugbahn. Durch diese Technik wird die Wechselwirkungszeit auf circa eine Sekunde erhöht. Die Genauigkeit der Cäsiumuhr ist dann durch Stoßverbreiterung limitiert, ein Problem, dass in der unten beschriebenen Ionenuhr nicht auftaucht, denn dort gibt es nur ein Ion.

4.3 Das Ion in der Falle - eine neue Uhr

Die Ionen-Uhr basiert auf einem Uhrenübergang in einem einfach positiv geladenen Ion. Dieses wird in einer Paulfalle (siehe Abschnitt 6.1) gefangen und spektroskopiert. Der große Vorteil dieser Uhren ist die viel höhere Übergangsfrequenz des Uhrenübergangs, die im optischen Bereich liegt. Das heißt, man arbeitet bei einer optischen Frequenz von circa 10^{15} Hertz, im Gegensatz zu der Radiofrequenz von circa 10^{10} Hertz im Cäsium. Das wirkt sich natürlich maßgeblich auf die Güte $Q = f/\Delta f$ der Schwingung aus. Die erfolgversprechendsten Kandidaten für ein neues Zeitnormal auf Basis eines gefangenen Ions sind derzeit $^{199}\text{Hg}^+$, $^{171}\text{Yb}^+$, $^{88}\text{Sr}^+$, $^{115}\text{In}^+$ und $^{40}\text{Ca}^+$ [10].

Funktionsweise einer Ionenuhr

Wie in Abschnitt 4.1.2 schon kurz angesprochen, braucht man für eine genaue Uhr auf der Basis eines gefangenen Ions ein Energieniveauschema mit zwei völlig verschiedenen Übergängen aus dem selben Grundzustand (siehe Abb. 4.4). Der Uhrenübergang hat aufgrund seiner langen Lebensdauer zwar eine sehr geringe Linienbreite, für die Kühlung ist er aber völlig ungeeignet, da für eine effektive Kühlung möglichst oft pro Sekunde absorbiert werden muss. Deshalb benutzt man zur Kühlung einen kurzlebigen Zustand. Bei einer Lebensdauer von $0,44\mu\text{s}$ werden so circa 200000 Photonen pro Sekunde absorbiert. Zweitens braucht man den Kühlübergang auch für den Nachweis und die Stabilisierung des Uhrenübergangs. Das Ion wird über die Fluoreszenz des Kühlübergangs nachgewiesen. Die vom Ion absorbierten Photonen werden in alle Raumrichtungen reemittiert und können, da es circa 200000 pro Sekunde sind, detektiert werden. Mit einem zweiten Laser wird gleichzeitig die Frequenz des Uhrenübergangs eingestrahlt. Geht das Ion in den Zustand E_2 über, kann es kurzfristig keine „Kühlphotonen“ mehr absorbieren und folglich auch nicht reemittieren. Es wird für kurze Zeit unsichtbar. Die Frequenz des Uhrenübergangs ist also exakt eingestellt, wenn die Dunkelphasen der Fluoreszenz maximal werden. Mit gefangenen Ionen als Taktgeber wurden bereits

⁶Quelle: Physikalisch Technische Bundesanstalt

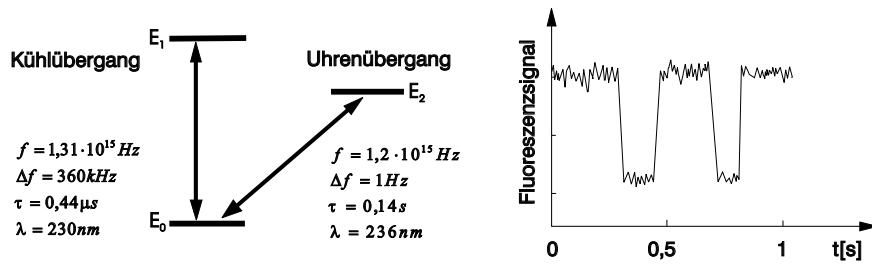


Abbildung 4.4: links: vereinfachtes Termschema des Indiumions, angegeben sind Frequenz, Lebensdauer, Linienbreite und Wellenlänge der beiden Übergänge, rechts: Dunkelphase im Fluoreszenzspektrum. Wird das Elektron in den Zustand E_2 gehoben wird es für kurze Zeit „unsichtbar“.

Uhren gebaut, die die Genauigkeit von Cäsiumfontänen erreichen. Theoretisch ist eine Genauigkeit von $\Delta f/f = 10^{-18}$ erreichbar. Hätte man eine solche Uhr zum Zeitpunkt des Urknalls (vor fast 14 Milliarden Jahren) gestartet, würde sie bis jetzt maximal eine Sekunde falsch gehen!

4.4 Wozu braucht man so genaue Uhren?

Die wohl bekannteste Anwendung von Atomuhren ist das GPS (Global Positioning System). Um die genaue Position auf der Erde mittels GPS bestimmen zu können wird sie von einem Netz von Satelliten umspannt, so dass von jedem Punkt der Erde immer mindestens drei sichtbar sind. Jeder dieser 24 Satelliten führt eine Atomuhr mit sich. Die Satelliten senden permanent Signale aus (elektro-magnetische Wellen), die stark vereinfacht lauten: „Ich bin Satellit A, meine Position ist B, die aktuelle Uhrzeit ist C.“ Der GPS-Empfänger empfängt solch ein Signal und kann dann berechnen wie lange es unterwegs war, denn die Absendezeit steht ja im Signal. Weil sich die Signale mit Lichtgeschwindigkeit c ausbreiten, kann der GPS-Empfänger aus der Laufzeit t auch die Entfernung x berechnen. (Dasselbe machen wir, wenn wir beim Gewitter die Sekunden zwischen Blitz und Donner zählen.) Um seine Position zu bestimmen, berechnet der Empfänger jetzt den Ort, der von den drei Satelliten jeweils x_1 , x_2 und x_3 entfernt ist (siehe Abbildung 4.5).

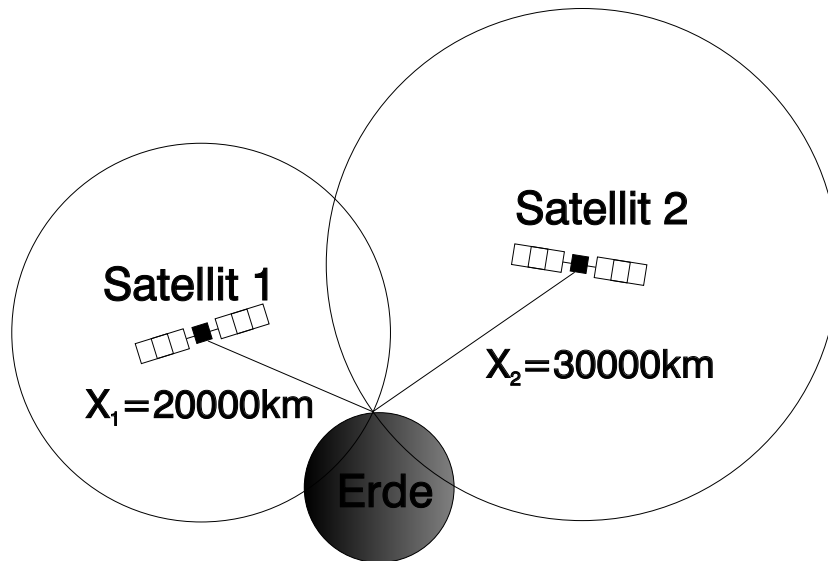


Abbildung 4.5: Positionsbestimmung beim GPS. Aus den Laufzeiten errechnet der Empfänger den Abstand vom jeweiligen Satelliten. Da die Erdoberfläche zweidimensional ist, wird noch ein dritter Satellit gebraucht.

4.5 Weitere Anwendungen von Teilchenfallen

Neben den jetzt ausführlich beschriebenen Möglichkeiten von Messungen mit höchster Präzision an einzelnen gefangenen und quasi wechselwirkungsfreien Teilchen, kommen Teilchenfallen auch bei Experimenten mit extremen Temperaturen zum Einsatz.

Eine der möglichen Energiequellen der Zukunft ist zum Beispiel die Kernfusion. Dabei werden zwei leichte Kerne zu einem schweren verschmolzen und es wird Energie frei. Leider braucht man zum Starten der Kernreaktion eine Temperatur von über einer Million Grad Celsius! Ein so heißes Fusionsplasma kann natürlich nicht in einem Gefäß mit materiellen Wänden eingeschlossen werden. Dabei würde das Plasma abkühlen und die Wand verdampfen! Deshalb wird das Plasma meist in einer Teilchenfalle namens Tokamak⁷ gefangen.

Aber auch am anderen Ende der Temperaturskala werden Teilchenfallen angewendet: Kühlt man bestimmte Atome oder Moleküle (Bosonen) auf Temperaturen von rund 100 nK (= 0,0000001°Celsius über dem absoluten Nullpunkt), so findet ein Phasenübergang statt. So wie Wasser bei 0°C von flüssig nach fest übergeht gehen die Atome/Moleküle in einen Zustand über, der Bose-Einstein-Kondensat⁸ genannt wird. Auch für die Erforschung solcher tiefkalter Gase werden Teilchenfallen eingesetzt.

⁷Torusförmige **K**ammer, die **ma**knetisch Teilchen fängt; Abkürzung aus dem Russischen

⁸Benannt nach Satyendra Nath Bose und Albert Einstein, die dieses Phänomen bereits 1924 theoretisch vorhergesagt haben. Die erste Herstellung eines Bose-Einstein-Kondensats gelang allerdings erst 1995.

5 Wie baut man eine Teilchenfalle?

Nach dem wir jetzt wissen, was man mit einer Teilchenfalle alles machen kann, sollen nun die grundlegenden Prinzipien erläutert werden.

Um ein Teilchen frei zum Schweben zu bringen, muss zum einen eine Kraft da sein, die das Teilchen gegen das Herunterfallen sichert, und zum anderen eine Kraft, die das Teilchen in der Falle hält. Es müssen also die folgenden Bedingungen erfüllt sein:

1. Die auf das Teilchen wirkende Gewichtskraft muss durch eine andere Kraft kompensiert werden. In der Fallenmitte muss die resultierende Kraft (=Summe aller auf das Teilchen wirkenden Kräfte) gleich null sein.
2. Wenn das Teilchen durch eine kleine Störung aus der Fallenmitte (Gleichgewichtslage) herausbewegt wird, muss es eine Kraft geben, die es wieder in die Gleichgewichtslage zurücktreibt. Man muss ein Potentialminimum erzeugen.¹

Wie wir im Folgenden sehen werden, ist die erste Bedingung recht leicht zu erfüllen, da die elektrischen und magnetischen Kräfte, die zum Fangen benutzt werden, sehr viel stärker als die Gravitation sind. Problematischer wird die Erfüllung der zweiten Bedingung. Als Beispiel hierfür kann der Milikanversuch² dienen: Durch das Einstellen der Spannung ist es kein großes Problem einen geladenen Öltropfen zwischen den zwei Kondensatorplatten zum Schweben zu bringen: Man erhöht die Spannung solange, bis sich die Gewichtskraft und die elektrische Kraft, die auf den Tropfen wirken, gerade aufheben. Das ist allerdings noch keine Teilchenfalle, denn jede noch so kleine Störung würde den Tropfen auf eine der Kondensatorplatten oder seitlich aus dem Kondensator treiben. (Der Tropfen ist ja kräftefrei und bewegt sich also einmal angestoßen geradlinig weiter.) Hier ist also die erste Bedingung (Kompensation der Gravitation) erfüllt, nicht jedoch die zweite (Potentialminimum).

¹Die Potentialfunktion $\Phi(x)$ gibt an wie groß die potentielle Energie/Lageenergie eines Teilchens am Ort x ist. Die Lageenergie eines Teilchens auf einer Oberfläche sieht zum Beispiel aus wie die Oberfläche selbst. Und weil ein Teilchen immer seine potentielle Energie minimieren will (Die Kraft auf ein Teilchen geht immer in die Richtung des steilsten Abfalls seiner potentiellen Energie), muss man mit elektrischen und/oder magnetischen Kräften ein Potentialminimum (Mulde) erzeugen, wenn man Teilchen fangen will.

²Beim Milikanversuch bringt man ein geladenes Öltröpfchen in einem waagrecht stehenden Plattenkondensator zum Schweben. Dadurch kann man die Elementarladung bestimmen.

5.1 Das Theorem von Earnshaw

Die scheinbar einfachste Möglichkeit, einen Körper zum freien Schweben zu bringen, ist die folgende: Man versucht, Ladungen so zu platzieren, dass ein geladenes Teilchen im Feld dieser Ladungsverteilung gefangen ist. Man könnte sich zum Beispiel eine homogen geladene Salatschüssel vorstellen, über der ein geladenes Styroporkügelchen schwebt. Das dies nicht funktionieren kann, wurde bereits 1839 von dem Physiker und Kleriker Samuel Earnshaw gezeigt. In einem Vortrag vor der Cambridge Philosophical Society bewies er das nach ihm benannte Theorem von Earnshaw [13]. In einer leicht abgewandelten umgangssprachlichen Formulierung lautet es:

Ein freies stabiles Schweben eines Körpers in einer beliebigen statischen Anordnung von elektrischen Ladungen und magnetischen Dipolen ist nicht möglich.

Dieses Theorem ist mit den Mitteln der Vektoranalysis sehr leicht zu beweisen. Da dieser streng mathematische Beweis jedoch sehr unanschaulich ist, soll der Sachverhalt an einem Spezialfall erläutert werden.[18]

Wir betrachten eine Punktladung in einem durch eine beliebige Ladungsverteilung erzeugten elektrischen Feld. Nehmen wir nun also an, die Ladung sitze in einem Potentialminimum (Fallenmitte) gefangen. Dann muss es ein kleines kugelförmiges Volumen um die Ladung geben, in dem sich keine andere Ladung befindet, denn die Ladung soll ja **frei** schweben. Auf der Oberfläche dieser Kugel muss die Kraft, die auf die Ladung wirkt, immer in die Kugel hinein zeigen, es soll ja eine **stabile** Schwebeposition sein. Also müssen auch die Feldlinien des elektrischen Felds immer in die Kugel hinein zeigen. Wenn also alle Feldlinien in das Volumen hinein zeigen, müssen sie sich irgendwo treffen und dort enden. Und wo enden Feldlinien? Natürlich dort wo Ladungen sitzen! Das ist aber ein Widerspruch, denn in unserem Volumen sind ja außer der schwebenden Punktladung keine Ladungen, und diese wirkt auf sich selbst keine Kraft aus!

Ganz so einfach ist es also nicht, eine Teilchenfalle zu bauen. Wie man das Theorem von Earnshaw überlisten kann, zeigen die Fallen im Koffer.

6 Die Fallen im Koffer

Nachdem jetzt schon sehr viel über die allgemeine Theorie von Teilchenfallen, sowie deren Anwendungen gesagt wurde, werden nun verschiedene Typen von Fallen für makroskopische Teilchen in Theorie und Praxis erklärt.

6.1 Paulfalle

Die Paulfalle hat ihren Namen von ihrem „Erfinder“ Wolfgang Paul (1913-1993), der die Theorie in den 50er Jahren entwickelte und 1955 die erste Paulfalle aufbaute; 1989 wurde er dafür mit dem Nobelpreis ausgezeichnet. Die Paulfalle ist eine elektrische Falle, das heißt, sie fängt Teilchen mit rein elektrischen Feldern. Damit ist auch sofort klar, dass mit einer Paulfalle nur elektrisch geladene Teilchen wie Ionen, Elektronen oder in unserem Fall geladener Staub gefangen werden können.

6.1.1 Theorie

Um ein Teilchen zu fangen, wollen wir eine Kraft erzeugen, die immer in Richtung Fallenzentrum zeigt. Der Einfachheit halber wollen wir, dass die Kraft F linear mit dem Abstand r vom Zentrum zunimmt; doppelter Abstand, doppelte Kraft.

$$\vec{F} = -const \cdot \vec{r} = -const \cdot \begin{pmatrix} x \\ y \\ z \end{pmatrix}$$

Da sich die Kraft als negativer Gradient (Ableitung) des Potentials Φ ergibt¹, sieht unser Potential in seiner allgemeinsten Form so aus.

$$\Phi(x, y, z) = ax^2 + by^2 + cz^2$$

Für ein elektrisches Potential gilt im ladungsfreien Raum (und den wollen wir ja im Zentrum der Falle) die Poisson-Gleichung, die direkt aus den Maxwell'schen Gleichungen² folgt.

$$\Delta\Phi = \frac{\partial^2}{\partial x^2}\Phi + \frac{\partial^2}{\partial y^2}\Phi + \frac{\partial^2}{\partial z^2}\Phi = 0$$

¹Die Kraft ergibt sich aus der Ableitung des Potentials; wo das Potential steiler ist, ist die Kraft auf die Kugel größer.

²Die Maxwell'schen Gleichungen sind vier Differentialgleichungen, die die grundlegenden Eigenschaften elektrischer und magnetischer Felder beschreiben.

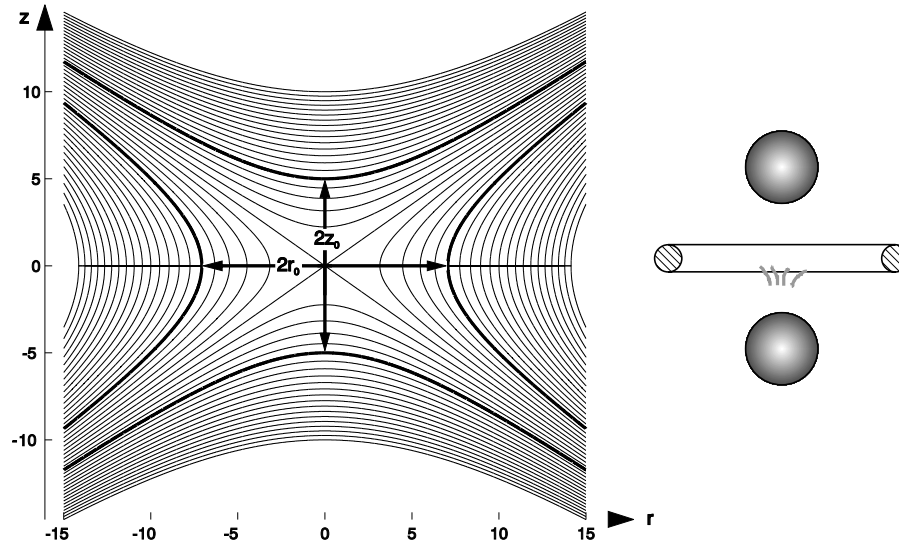


Abbildung 6.1: Links: Äquipotentiallinien des Potentials ($\Phi = \text{const}$). Wenn man die dicken Linien um die z -Achse ($r = 0$) rotiert erhält man die Form der Elektroden. Rechts: Die hyperbolische Form der Elektroden wird bei unserer Paulfalle durch zwei Kugeln und einen Ring angenähert. Es ist auch die ungefähre Position der Teilchen eingezeichnet.

Die Poissongleichung sagt übrigens genau das aus, was wir zum Beweis des Theorems von Earnshaw in Abschnitt 5.1 verwendet haben. Für unser Potential ergibt sich also $2a + 2b + 2c = 0$, die drei Koeffizienten können somit nicht allesamt das gleiche Vorzeichen haben. Wenn wir nun noch fordern, dass unsere Falle rotationssymmetrisch um die z -Achse ist ($a = b$) haben wir unser Potential bereits bis auf eine Konstante bestimmt. Es gilt dann $4a + 2c = 0$, also $c = -2a$ und unser Potential lautet

$$\Phi(x, y, z) = a(x^2 + y^2 - 2z^2).$$

Nutzen wir nun noch die Rotationssymmetrie $x^2 + y^2 = r^2$ aus (Satz des Pythagoras), dann lautet unser Potential

$$\Phi(r, z) = a(r^2 - 2z^2).$$

Jetzt wissen wir, wie die Äquipotentialflächen (Flächen auf denen sich das Potential nicht ändert) $\Phi(x, y, z) = \text{konstant}$ aussehen. Damit wissen wir aber auch, wie die Elektroden aussehen, die solch ein Feld erzeugen, denn dort ist ja gerade das Potential konstant. Für unser Potential ergeben sich als Äquipotentialflächen Hyperbeln, die um die z -Achse rotiert werden. Und so sehen die Elektroden einer idealen Paulfalle auch aus (Abbildung 6.1). Damit das Innere der Falle besser beobachtbar ist (eine ideale Paulfalle ist quasi geschlossen), sind die Elektroden unserer Falle zwei Kugeln und ein Ring. Leider haben wir so unsere anfangs gestellte Aufgabe nicht gelöst (Potentialminimum), denn wenn wir jetzt an unsere Elektroden eine Spannung (=elektrisches Potential) anlegen, wird ein geladenes

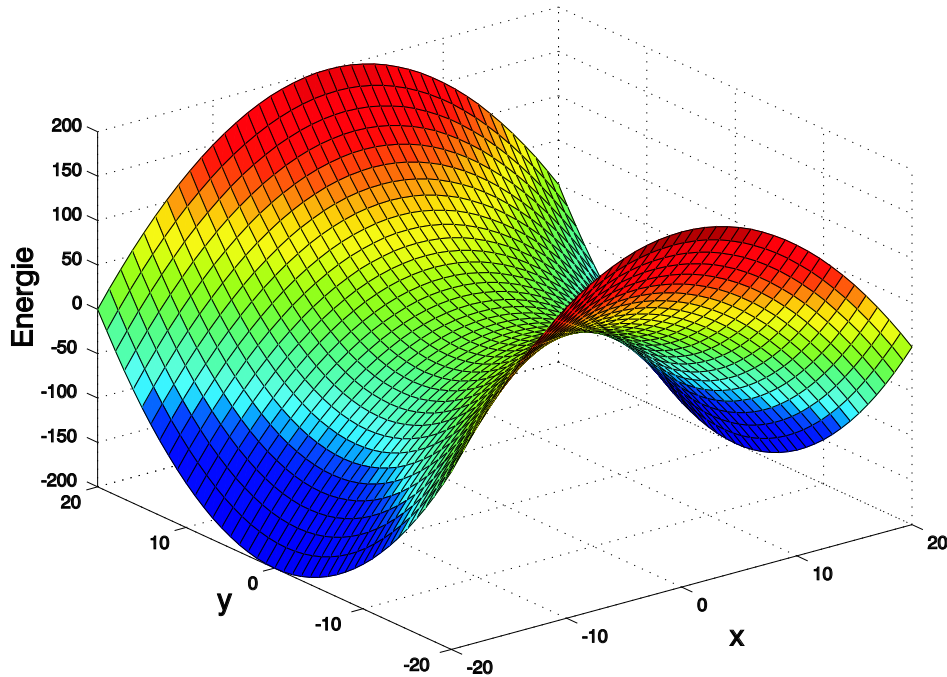


Abbildung 6.2: Rolzt eine Kugel auf dieser Sattelfläche, so ist sie in y-Richtung stabil (Kugel rollt zur Mitte), in x-Richtung jedoch instabil (Kugel rollt nach außen weg).

Teilchen nicht in der Mitte der Falle bleiben. In einer der beiden Richtungen (x oder z) wird das Teilchen von den Elektroden abgestoßen, und ist somit stabil, in der anderen wird es jedoch zwangsläufig angezogen, und bleibt nicht gefangen. (Ein positiv geladenes Teilchen wird zum Beispiel von den ebenfalls positiv geladenen Kugeln abgestoßen, vom negativ geladenen Ring jedoch angezogen.) Ein Potential wie unseres nennt man Sattelpotential (Abbildung 6.2). Legt man eine Kugel auf einen Sattel, so wird sie nicht vorn oder hinten hinausrollen, denn da geht es ja bergauf (stabile Richtung). Rechts und links jedoch geht es bergab (instabile Richtung). Dass es so (statisch) nicht funktionieren kann, war aber nach dem letzten Kapitel schon klar. Der Trick ist nun, anstatt einer Gleichspannung eine Wechselspannung anzulegen. Unser endgültiges Potential lautet also³:

$$\Phi(r, z, t) = \frac{U_0 \cdot \sin(\omega t)}{2r_0^2} (r^2 - 2z^2)$$

Das entspricht dann einem Sattel, bei dem sich die stabile und die instabile Richtung ständig abwechseln. Immer wenn die Kugel auf der einen Seite hinabrollen will, biegt sich diese Seite nach oben und treibt sie zurück in die Mitte. Die Kugel rollt dann in die andere (jetzt instabile) Richtung bis sich diese nach oben biegt

³Der Faktor $1/2r_0^2$ ist ein reiner Geometriefaktor, der die Größe der Falle beschreibt. Aus der Bedingung, dass die rücktreibende Kraft in allen drei Richtungen gleich vom Abstand abhängen soll, ergibt sich für die Größe der Falle in z-Richtung $z_0 = r_0/\sqrt{2}$ (siehe auch Abbildung 6.1).

und so weiter und so weiter. Man kann sich gut vorstellen, dass ab einer gewissen Frequenz der Wechselspannung das Teilchen ständig hin- und hergeschubst wird und nicht mehr aus der Falle kommt. Im zeitlichen Mittel spürt die Kugel ein immer anziehendes sogenanntes Pseudopotential. Um den Sachverhalt jedoch genauer zu verstehen, muss man sich die Bewegungsgleichungen des Teilchens in diesem Feld genauer anschauen. Dazu setzt man die Kraft, die ein Teilchen an einem gewissen Ort erfährt, gleich dem, was die Kraft mit dem Teilchen tut, sie beschleunigt es. (Erstes Newtonsches Axiom: $F = m \cdot a$) Dadurch erhält man einen Zusammenhang zwischen dem Ort des Teilchens und der zweiten Ableitung des Ortes nach der Zeit, nämlich der Beschleunigung. Durch das Lösen dieser Differentialgleichung kann man die Bahn (Ort in Abhängigkeit von der Zeit) des Teilchens berechnen.

In unserem Fall (konservative Kräfte) ergibt sich die Kraft aus der negativen Ableitung des Potentials multipliziert mit der Ladung q des Teilchens. Die Bewegungsgleichungen lauten damit

$$F(r) = -q \cdot \frac{\partial \phi}{\partial r} = -q \frac{U_0 \cdot \sin(\omega t)}{r_0^2} \cdot r = m \cdot a = m \cdot \frac{\partial^2 r}{\partial t^2}$$

$$F(z) = -q \cdot \frac{\partial \phi}{\partial z} = q \frac{2U_0 \cdot \sin(\omega t)}{r_0^2} \cdot z = m \cdot a = m \cdot \frac{\partial^2 z}{\partial t^2}$$

oder auf das Wichtigste reduziert

$$\ddot{r}(t) = -\frac{qU_0}{mr_0^2} \cdot \sin(\omega t) \cdot r(t)$$

$$\ddot{z}(t) = \frac{2qU_0}{mr_0^2} \cdot \sin(\omega t) \cdot z(t)$$

Durch die Lösung dieser Differentialgleichung kann man jetzt die Bahn eines geladenen Teilchens in der Falle berechnen. Wie diese Bahn aussieht, hängt natürlich noch von den Parametern ab. Diese sind die Masse m und die Ladung q des Teilchens, die Amplitude U_0 und die Frequenz ω der angelegten Wechselspannung und der Geometrieparameter r_0 . Diese Mathiesche Differentialgleichung⁴ besitzt im allgemeinen keine analytische Lösung, die Bahn des Teilchens ist also nicht durch elementare Funktionen (trigonometrische, exponentielle, polynomiale, etc.) darstellbar. Wichtig ist für uns nur die Existenz eines Stabilitätsparameters $\tilde{s}$ für den gelten muss:

$$\tilde{s} = \frac{4 \cdot q \cdot U_0}{m \cdot r_0^2 \cdot \omega^2} < 0,908 \quad (6.1)$$

Das heißt, nur wenn man die Parameter so einstellt, dass $\tilde{s}$ kleiner als 0,908 ist, führt das Teilchen eine stabile Bewegung aus und bleibt in der Falle gefangen.

⁴Differentialgleichungen der Form $\ddot{f}(t) = (a + b \sin t)f(t)$, benannt nach dem französischen Mathematiker Emile Leonard Mathieu, *1835 †1890.

Anderenfalls wird sich das Teilchen immer weiter vom Fallenzentrum entfernen und schließlich gegen eine der Elektroden prallen.

Im Fall einer stabilen Bahn kann die Lösung näherungsweise als Summe zweier Sinusschwingungen beschrieben werden. Die Mikrobewegung $\delta(t)$ hat eine kleine Amplitude und dieselbe Frequenz wie die angelegte Wechselspannung, die Makrobewegung $\Delta(t)$ hat eine größere Amplitude und kleinere Frequenz.

$$z(t) = \delta(t) + \Delta(t) = \delta_0 \sin(\omega t) + \Delta_0 \sin(\Omega t) \quad z_0 < Z_0 \quad \omega > \Omega$$

Da sich die Differentialgleichungen in z- und r-Richtung nur um einen Faktor -2 unterscheiden, ist auch die Bahn $r(t)$ eine Überlagerung einer Makro- und einer Mikrobewegung.

Das Theorem von Earnshaw wird hier nicht verletzt, da es sich ja nicht um statische Felder handelt.

6.1.2 Praxis

AN DEN ELEKTRODEN DER PAULFALLE LIEGT EINE WECHSELSPANNUNG VON 1500 VOLT! NIEMALS OHNE FESTGESCHRAUBTEN DECKEL UND IMMER NUR IM BEISEIN EINES LEHRERS BETREIBEN!

Bevor mit der Falle experimentiert werden kann, muss für eine **ausreichende Beleuchtung** gesorgt werden, denn die gefangenen Teilchen sind sehr klein⁵. Ohne Beleuchtung kann es deshalb passieren, dass man etwas gefangen hat ohne es zu sehen. Damit ihr eine Vorstellung davon bekommt, wie klein die Teilchen sind, und wo sie sich in der Falle befinden, sind auf der beigefügten CD, mehrere Filme von gefangenen Teilchen.

Zum Beleuchten eignet sich jede lichtstarke Lampe, deren Strahl sich fokussieren lässt. Den Strahl der Lampe richtet man schräg von oben auf das Zentrum der Falle, also genau zwischen die beiden Kugeln, denn dort werden die Teilchen gefangen. Wenn möglich sollten Reflexionen an den Kugeln vermieden werden. Ist die Lampe arretiert, kann mit dem Einbringen der Teilchen begonnen werden. Dazu steckt man das Netzteil ein und drückt die RESET-Taste am Stecker. Mit dem Drehregler stellt man die Spannung auf circa 1200 Volt ein. Die Teilchen, die durch Reibungseffekte immer ein wenig geladen sind, lädt man mit dem Kabelbinder (schmaler Plastikstreifen) in die Falle. Dazu bringt man auf dessen Spitze eine kleine (!) Menge des gewünschten Stoffs (Für den Anfang eignet sich das zerstoßelte Kochsalz.) und bugsiert die Spitze durch das kleine seitliche Loch in die Mitte der Falle. Dort zerstäubt man den Staub durch ein leichtes Klopfen auf den Kabelbinder. Wer jetzt enttäuscht die nächste Ladung einbringt, weil nichts gefangen wurde, sollte eventuell nochmal genau hinsehen. Die Teilchen sollten als kleine Striche zu sehen sein; ein dunkler Hintergrund ist hier von Vorteil! Wenn

⁵Ideal ist die Abbildung der Teilchen mit einer Kamera und Makroobjektiv.

tatsächlich nichts da sein sollte, wiederholt man die Prozedur eben noch einmal, bis man etwas gefangen hat.

Aber warum sind die gefangenen Teilchen eigentlich kleine Striche, wo wir doch punktförmige Staubteilchen eingebracht haben? Das liegt an der oben erwähnten Bewegung, die die Teilchen in der Falle ausführen. Wegen der Luftreibung wird die Makrobewegung sehr schnell ausgedämpft (circa eine Sekunde).⁶ Übrig bleibt nur die Mikrobewegung und die oszilliert mit der Frequenz der angelegten Spannung, also 50 Hertz. Weil unser Auge diese schnelle Bewegung nicht verfolgen kann, sehen wir eben Striche. Wer die Teilchen als Punkte sehen will, muss sie mit einem Stroboskop beleuchten. Wenn man sie nämlich 50 Mal pro Sekunde anblitzt, dann sieht man sie immer an der selben Stelle und folglich als Punkte.

Durch Drehen am Spannungsregler kann man die Bahnen der Teilchen beeinflussen. Verringert man die Spannung, so werden die Oszillationen und die Abstände der einzelnen Teilchen größer. Wenn man die Spannung immer weiter verringert, verliert man irgendwann auch Teilchen, da diese durch die geringere Spannung nicht mehr so stark an die Falle gebunden sind und durch Störungen (zum Beispiel Luftbewegung) herausfallen.

Wenn man nur wenige Teilchen in der Paulfalle gefangen hat⁷, kann man die Bildung von sogenannten Coulombkristallen beobachten. Die einzelnen Teilchen werden vom Feld der Elektroden in die Mitte der Falle gedrückt. Sind jedoch mehrere in der Falle, so stoßen sie sich aufgrund ihrer Ladung auch ab. Das hat zur Folge, dass sich ein Gleichgewicht zwischen den ins Zentrum gerichteten Kräften und den abstoßenden Coulombkräften einstellt. Daraus ergibt sich eine Positionierung der Teilchen mit gleichmäßigen Abständen, wie in einem Kristallgitter. Durch Erhöhen der Spannung kann die „Gitterkonstante“⁸ kleiner gemacht werden. Leider sind die Coulombkristalle aufgrund der unterschiedlichen Masse und Ladung der Teilchen nicht 100%ig symmetrisch. Die regelmäßige Anordnung von Teilchen lässt sich mit den länglichen Elektroden besser beobachten.

Um diese einzubauen stellt man den Spannungsregler auf null und **steckt das Netzgerät aus**. Dann kann man die drei Schrauben, die den Deckel halten lösen und diesen abnehmen. Um die Elektroden zu tauschen löst man die Muttern (Bitte keine Gewalt, die Gewinde sind Feingewinde!) und entfernt die Ringelektrode. Sie wird durch die längliche Elektrode ersetzt. Um auch die obere und untere Elektrode „länglich“ zu bekommen werden die Hülsen über die Kugeln gesteckt. Die drei Elektroden sollten jetzt so eingestellt werden, dass die mittlere möglichst waagrecht und die beiden anderen möglichst symmetrisch zur mittleren sind. Dann wird die Abdeckung wieder angebracht. Die Mutter dabei so fest anziehen, dass der Deckel zwar fest, aber noch verschiebbar ist. Nun bringt man wieder wie oben beschrieben Teilchen in die Falle. Diese ordnen sich nun in einer linearen Kette

⁶Die Luftreibung wurde in den Bewegungsgleichungen nicht berücksichtigt.

⁷Das lässt sich recht leicht erreichen, wenn man die Spannung verringert und wieder erhöht und so immer mehr Teilchen aus der Falle entfernt.

⁸Dieser Vergleich ist mit Vorsicht zu genießen. Unsere Abstände bewegen sich im Millimeterbereich, in Kristallen sind die Abstände zwischen den Atomen rund $1\text{\AA}=10^{-10}m$.

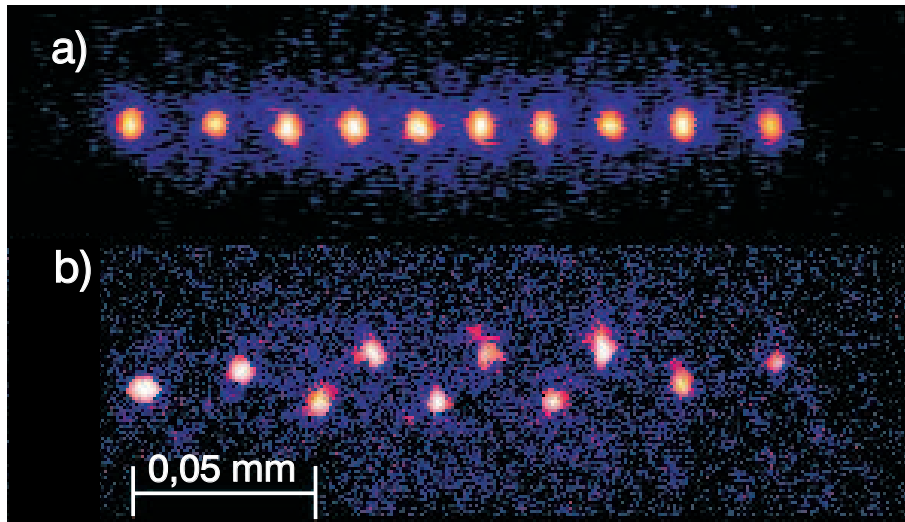


Abbildung 6.3: Coulombkristalle mit $^{40}\text{Ca}^+$ -Ionen in einer linearen Paulfalle. Abhängig von der Potentialtiefe in r- und z-Richtung bildet sich eine lineare Kette (a) oder eine Zick-Zack-Struktur (b). Die Bilder stammen von einer Arbeitsgruppe des ETAP an der Universität Mainz.

(man könnte auch sagen eindimensionaler Coulombkristall) an. Durch Verschieben des Deckels kann man die Teilchen kollektiv in eine Ecke der Elektrode schieben und dann zusehen, wie sich die energetisch günstigste Kette wieder ausbildet. In solchen linearen Paulfallen wurden auch die ersten Experimente mit Quantencomputern gemacht (Abbildung 6.3).

6.1.3 Der rotierende Sattel

Der rotierende Sattel verdeutlicht auf eindrucksvolle Weise das Funktionsprinzip der Paulfalle. Legt man an eine Paulfalle eine Gleichspannung an, so sieht das elektrische Potential aus wie ein Sattel. In einer Richtung wird das Teilchen in die Mitte der Falle gedrückt (zum Beispiel von den Kugeln weg), in der anderen aber aus der Mitte heraus geschoben (zum Ring hin). Das gleiche passiert mit einer Kugel im Sattel.

Legt man aber an die Paulfalle eine Wechselspannung an, so bleiben die Teilchen gefangen. Beim Sattel kann man die Kugel fangen, wenn man ihn in Rotation versetzt. Dann wechseln sich die anziehende und die abstoßende Richtung ständig ab und die Kugel bleibt in der Mitte. Wenn man den Motor an ein regelbares Gleichspannungsnetzteil anschließt, kann der Sattel über den Gummiriemen angetrieben werden. Legt man nun die Stahlkugel in die Mitte, so bleibt diese gefangen, sobald die Drehzahl hoch genug ist. Das passiert ungefähr bei einer Spannung von 7,5 Volt. Die maximale Spannung beträgt 10 Volt! So wie bei der Paulfalle die Frequenz und Amplitude der Wechselspannung bestimmen, ob die Teilchen gefangen werden, gehen hier die Rotationsgeschwindigkeit und Höhe des Sattels in den

Stabilitätsparameter ein.

Wenn der Sattel über einen längeren Zeitraum nicht betrieben wird, bitte den Gummiriemen abnehmen, damit er nicht ausleiert.

6.2 Levitron

Das Levitron ist ein Spielzeug, das von der Firma Fascinations vertrieben wird. Ein magnetischer Kreisel schwebt stabil über einer ebenfalls magnetischen Bodenplatte. Die Erfindung dieses Spielzeugs geht auf Roy Harrigan zurück, der es sich im Jahr 1983 patentieren lies. Im Jahr 1995 lies sich Edward W. Hones eine zweite sehr ähnliche Erfindung patentieren und hat diese seither mehrere Millionen Mal unter dem Namen Levitron verkauft.

6.2.1 Theorie

Die Kraft, die hier die Gewichtskraft des Kreisels kompensiert ist die magnetische Kraft. Diese Wechselwirkung unterscheidet sich von der elektrischen generell darin, dass es keine magnetischen Monopole gibt. (Während sich elektrische Ladungen in positive und negative trennen lassen, wird man nie einen einzelnen Süd- oder Nordpol finden. Diese treten immer nur zu zweit auf, man nennt sie deshalb auch Dipole.) Beim Levitron bestehen sowohl der Kreisel, als auch die Bodenplatte aus Permanentmagneten. Sind diese so orientiert, dass zum Beispiel Nordpol gegen Nordpol gerichtet ist, stoßen sie sich ab, und der Kreisel kann schweben.⁹ Der Gleichgewichtspunkt ist dort, wo die magnetische Kraft gerade gleich der konstanten Gewichtskraft ist.

Die Rotation des Kreisels sorgt jetzt dafür, dass der Kreisel nicht umkippt. Der Nordpol in der Bodenplatte übt ja eine abstoßende Kraft auf den Nordpol des Kreisels und eine anziehende auf dessen Südpol aus. Diese beiden Kräfte ergeben zusammen ein Drehmoment, das auf den Kreisel wirkt. Ist der Kreisel in Ruhe, wird er umgedreht, magnetische und Gravitationskraft zeigen in eine Richtung und er fällt herunter. Dreht er sich aber, so bewirkt das Drehmoment, dass der Kreisel eine Präzessionsbewegung ausführt. Das heißt, dass sich die Spitze des Kreisels auf einer Kreisbahn um die Senkrechte bewegt. Dies kann man auch an einem normalen Kreisel beobachten, dessen Achse nicht senkrecht steht. Durch die Präzessionsbewegung wird der Kreisel also gegen das Umkippen stabilisiert. Je langsamer die Drehbewegung des Kreisels ist, desto schneller ist die Präzession. Dies lässt sich auch am Levitron beobachten, vor allem kurz bevor der Kreisel herunterfällt: Wenn die Rotation wegen des Luftwiderstands langsamer wird, erhöht sich die Präzessionsfrequenz (sichtbar durch ein schnelles Zittern der Kreiselspitze), bis der Kreisel schließlich kippt und abstürzt.

⁹Diese Abstoßung resultiert aus einer Summe von vier Kräften. Zwei abstoßende (Süd-Süd und Nord-Nord) minus zwei anziehende (zweimal Süd-Nord). Dass eine abstoßende Gesamtkraft resultiert, liegt daran, dass die Kräfte mit dem Abstand abnehmen.

Die Stabilität in horizontaler Richtung lässt sich leider auf Schulniveau nicht erklären. Sie hängt jedoch damit zusammen, dass das Magnetfeld der Bodenplatte leicht von der Senkrechten abweicht, wenn man sich aus der Mitte herausbewegt. Der Kreisel orientiert sich dann an diesem leicht verkippten Feld, was eine rücktreibende Kraft bewirkt.¹⁰

Wer sich eine Weile mit dem Levitron beschäftigt, merkt, wie klein der stabile Bereich ist. Deshalb ist auch die Einstellung des Gewichts so sensibel. (Die leichteste Scheibe wiegt nur 0,3% vom Gewicht des Kreisels. Aus den Stabilitätsbetrachtungen folgt auch, dass es eine obere Grenze für die Rotationsgeschwindigkeit des Kreisels gibt. Ist er zu schnell, dann wird er radial (seitlich) instabil und fällt aus der Gleichgewichtsposition.[15] Man kann dieses obere Limit beobachten, wenn man den schwebenden Kreisel beschleunigt, zum Beispiel durch einen Luftstrom.

Das Theorem von Earnshaw wird hier dadurch überlistet, dass wir es nicht mit einer statischen Anordnung zu tun haben. Der Kreisel rotiert und reagiert dynamisch auf das äußere Feld, somit ist dies kein statisches System. Mit einem ähnlichen Prinzip lassen sich auch Neutronen einfangen. In einem geschickt gewählten Magnetfeld orientieren sich die magnetischen Neutronen an den Feldlinien und können so gefangen werden.

6.2.2 Praxis

Wenn man den Kreisel in die Hand nimmt, kann man damit sehr gut das Magnetfeld der Bodenplatte spüren und auch grob feststellen, wo der Gleichgewichtspunkt ist.

Einstellung des Gewichts Zuerst wird das Gewicht des Kreisels richtig eingestellt. Dies geschieht mit verschiedenen schweren Plastik- und Metallscheiben.¹¹ Zum Starten des Kreisels legt man die durchsichtige Startplatte auf die Bodenplatte und dreht den Kreisel in deren Mitte an. Dies kann entweder mit der Hand (braucht etwas Übung) oder mit dem Starter geschehen. Dazu steckt man den Kreisel in den Starter und setzt ihn in die Mitte der Startplatte. Durch drücken auf den Knopf kann man den Kreisel in Rotation versetzen. Es ist dabei **nicht** notwendig, den Kreisel nach unten zu drücken. Wenn sich die Drehzahl nicht mehr ändert (nach ein paar Sekunden) zieht man den Starter senkrecht nach oben und hält den Knopf dabei gedrückt. Nun wird die Startplatte mitsamt dem Kreisel **ganz langsam** angehoben bis dieser von der Platte abhebt. Wenn er nicht abhebt, muss man etwas Gewicht abnehmen, wenn er von der Platte „springt“ muss Gewicht hinzugefügt werden. Um das Gewicht zu ändern, muss man den schwarzen Gummiring

¹⁰Ein immer senkrecht stehender Kreisel kann nicht stabil sein. Das wäre ein Widerspruch zu den Maxwellschen Gleichungen.

¹¹Ein guter Ausgangszustand sind eine große und zwei kleine Metallscheiben plus eine rote Plasticscheibe.

abziehen und danach wieder fest machen, damit sich die Scheiben nicht drehen können.

Einstellung des Winkels Nachdem das Gewicht richtig justiert ist, wird das Levitron in der Regel noch nicht funktionieren, da der Kreisel seitlich ausbricht. Dies wird durch die drehbaren Füße der Bodenplatte behoben. Auf der Seite wo der Kreisel ausbricht, muss die Platte angehoben beziehungsweise auf der anderen abgesenkt werden.

Nachdem alles gut justiert ist, ist es faszinierend, dem Kreisel beim Schweben zuzusehen. Er schwebt so mehrere Minuten lang - bis die Rotationsgeschwindigkeit zu langsam wird und die immer schneller werdende Präzessionsbewegung den Kreisel nicht mehr senkrecht halten kann. Man kann allerdings auch sehen wie klein der stabile Bereich des Levitrons ist. Schon durch eine kleine Störung (Lufthauch, kleiner Magnet) kann man den Kreisel aus seiner stabilen Position werfen. Aber auch andere Faktoren können das System stören. So kann zum Beispiel eine kleine Änderung der Temperatur eine Änderung des Magnetfelds von Bodenplatte und/oder Kreisel bewirken. Dann müssen Gewicht und/oder Winkel neu eingestellt werden.

6.3 Diamagnetisches Schweben

Das diamagnetische Schweben beruht, wie der Name schon sagt, auf dem Phänomen des Diamagnetismus. Es ist das einzige Experiment, das es schafft ohne Energiezufuhr von außen ein Teilchen unendlich lange schweben zu lassen!

6.3.1 Theorie

Para- und ferromagnetische Stoffe sind dadurch charakterisiert, dass sie mikroskopische magnetische Momente besitzen. Diese haben ihren Ursprung in den Bahnen der Elektronen. Ein Elektron, das, in einem vereinfachten Bild, um den Kern kreist ist ein Kreisstrom und Kreisströme erzeugen, wie in einer Spule, Magnetfelder. Das heißt, jedes Atom/Molekül ist wie ein kleiner Magnet. Bringt man diese in ein Magnetfeld, so orientieren sie sich an den Feldlinien (wie eine Kompassnadel) und sie werden angezogen.

Bei diamagnetischen Stoffen ist das anders. Dort „kreisen“ zwar auch Elektronen, aber immer gleich viele in jeder Richtung. In der Summe ist das resultierende Magnetfeld gleich null; sie besitzen keine atomaren magnetischen Momente. Bringt man einen Diamagneten aber in ein Magnetfeld, so heben sich die Kreisströme nicht mehr auf. Es werden wie beim Transformator Ströme induziert, die ein entgegengesetztes Magnetfeld erzeugen (Lenzsche Regel). Es wird also vom Diamagnet ein Magnetfeld erzeugt, das genau in die andere Richtung zeigt wie das erzeugende. Da sich diese beiden Felder abstoßen, kann das diamagnetische Graphitplättchen über den Magneten schweben.

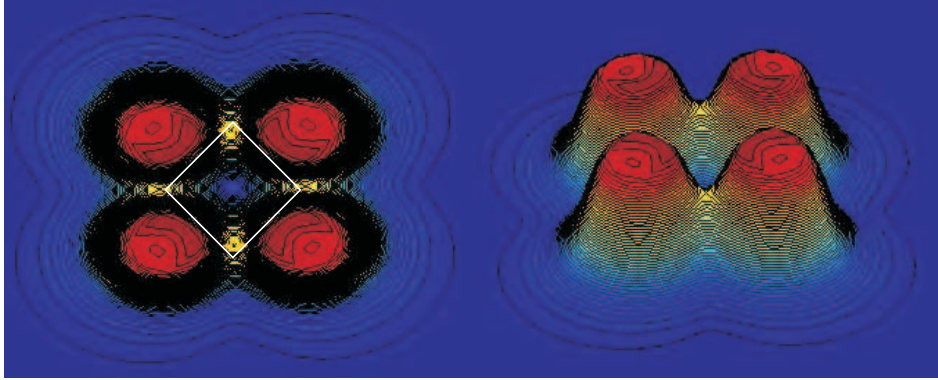


Abbildung 6.4: links: Der Betrag des Magnetfelds in der Schwebehöhe als Falschfarbenbild. Das Graphitplättchen sucht den Ort, wo es am wenigsten vom Magnetfeld durchdrungen wird. rechts: Dreidimensionale Darstellung von $|B|$.

Um die seitliche Stabilität zu erklären, muss man sich die potentielle Energie des Plättchens im Magnetfeld und den Feldverlauf genauer ansehen: Die Energie eines magnetischen Dipols $\vec{\mu}$ im Feld $\vec{B}$ ist $E = -\vec{\mu} \cdot \vec{B}$. Zusammen mit der potentiellen Energie im Gravitationsfeld der Erde erhält man.

$$E = -\vec{\mu} \cdot \vec{B} + m \cdot g \cdot z$$

Da der Dipol vom äußeren Magnetfeld induziert ist gilt $\vec{\mu} = \chi \vec{B}$, wobei $\chi < 0$ ist. Eingesetzt ergibt sich

$$E = -\vec{\mu} \cdot \vec{B} + m \cdot g \cdot z = -\chi \cdot \vec{B} \cdot \vec{B} + m \cdot g \cdot z = |\chi| \cdot |\vec{B}|^2 + m \cdot g \cdot z$$

und diese Energie wird von dem Graphitplättchen minimiert. Da die Energie mit $|B|^2$ steigt, richtet es sich so aus, dass es von dem Magnetfeld möglichst wenig spürt. Dazu müssen wir uns das Magnetfeld der vier goldenen Würfel genauer ansehen.

Jeder der vier Würfel ist ein sehr starker Magnet. Die Magnete sind so angeordnet, dass diagonal gegenüberliegende die gleiche Orientierung haben. Es zeigen also zwei Nord- und zwei Südpole nach oben. Das Magnetfeld B dieser Magneten wurde numerisch berechnet und sein Betrag in Abbildung 6.4 graphisch dargestellt. Anhand des Bildes wird auch klar, warum die um 45° verdrehte Position ein stabiles Minimum für das Plättchen ist. Jede Auslenkung aus dieser Position bewirkt, dass es insgesamt „mehr“ vom äußeren Magnetfeld spürt. Das bewirkt eine Erhöhung der potentiellen Energie und folglich eine rücktreibende Kraft.

Das Theorem von Earnshaw wird auch hier nicht verletzt. Erstens kannte Earnshaw den Diamagnetismus noch nicht¹², zweitens gilt Earnshaws Theorem für einzelne elektrischen Ladungen und magnetische Dipole nicht für Festkörper und

¹²Der Diamagnetismus ist ein sehr schwacher Effekt, der nur mit sehr starken Magnetfeldern und sehr stark diamagnetischen Materialien sichtbar wird. In unserem Versuch werden Neodymmagnete und sogenanntes pyrolithisches Graphit verwendet.

deren Materialeigenschaften.¹³

Der mittlerweile oft gezeigte Versuch des schwebenden Supraleiters basiert übrigens auf genau dem selben Effekt. Da der Supraleiter ein perfekter Diamagnet ist ($\chi = -1$)¹⁴ erzeugt er sogar ein Gegenfeld, das genau so groß ist wie das äußere.

6.3.2 Praxis

Bitte vorsichtig mit dem Graphitplättchen sein. Es ist sehr zerbrechlich und ziemlich teuer. Keine Permanentmagnete oder ferromagnetische Stoffe in die Nähe der Magnetbasis bringen. Die Neodymmagnete sind sehr spröde und können splintern. Das Experimentieren mit dem diamagnetischen Schwebexperiment gestaltet sich denkbar einfach: Man zieht vorsichtig den magnetischen Streifen vom Graphitplättchen - fertig. Mit einem (nicht ferromagnetischen!) Stift o.ä. kann das Plättchen ein wenig aus seiner Gleichgewichtsposition verschoben/verdrehen werden; man sieht dann gut, wie es sich wieder in die stabile Position zurückschiebt/dreht.

6.4 Geregelt Schweben

Eine weitere Möglichkeit, Dinge zum Schweben zu bringen, ist ihre aktuelle Lage durch irgendeinen Mechanismus (zum Beispiel eine Lichtschranke) zu messen und mit diesem Signal die Kraft auf das Teilchen zu steuern. Auf diesem Prinzip basiert der schwebende Globus.

6.4.1 Theorie

Wie bereits gesagt, wird die Kraft auf das zu fangende Teilchen durch dessen aktuelle Position bestimmt. In anschaulichen Worten entspricht das in etwa dem Potential, das man erzeugt, wenn man einen Tischtennisball auf einem Schläger balanciert. Man hebt den Schläger immer gerade da an, wo der Ball hinunterrollen will. Der Sensor, der hier die Position bestimmt, ist das Auge. Mit diesem Signal (Vorsicht Ball zu weit rechts!) wird die Hand gesteuert (Rechts hoch!).

Unser schwebender Globus (Abbildung 6.5) wird mit magnetischen Kräften zum Schweben gebracht. In der Schwebeposition wird ein Großteil der Gewichtskraft des Globus durch die Anziehung zwischen dem Permanentmagneten des Globus und der Eisenschraube kompensiert. Um die noch fehlende Kraft aufzubringen, wird mit der Spule ein Magnetfeld erzeugt, das den Magnet ebenfalls anzieht. Wie stark dieses Magnetfeld sein muss, hängt natürlich von der aktuellen Position des Globus ab. Um diese zu bestimmen, ist unter der Eisenschraube eine Hallsonde angebracht, die das Magnetfeld des Permanentmagnets „spürt“. Das Ausgangssignal der Hallsonde ist also von der Position des Globus abhängig. Mit diesem

¹³In dem Bild der kreisenden Elektronen kann man auch argumentieren, dass die Elektronen ja dynamisch auf das äußere Feld reagieren und der Schwebemechanismus also gar nicht statisch ist.

¹⁴Die Suszeptibilität „normaler“ diamagnetischer Stoffe ist rund -10^{-6}

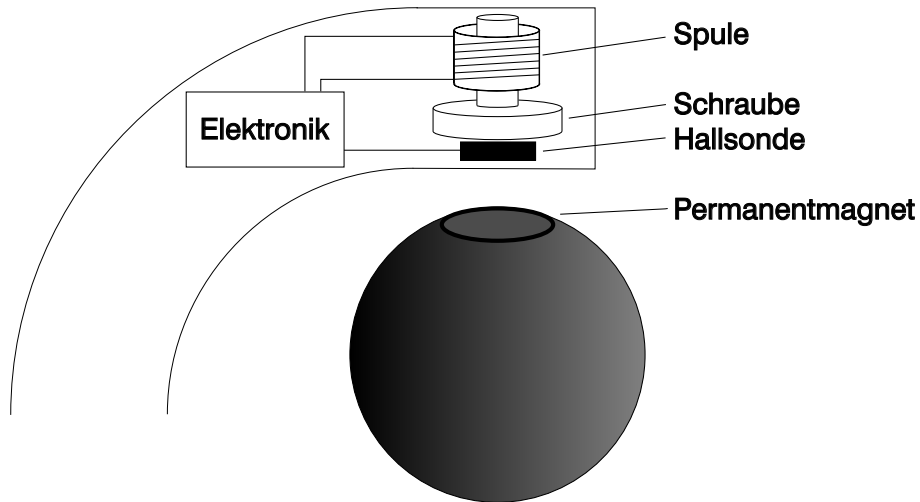


Abbildung 6.5: Schematischer Aufbau des schwebenden Globus. Der Globus wird durch den Permanentmagneten schon fast gehalten. Die restliche Kraft wird von der Spule erzeugt, die mit dem Signal der Hallsonde gesteuert wird.

Feedback wird nun der Strom in der Spule so gesteuert, dass die Schwebeposition stabil ist. Earnshaws Theorem kann auch hier nicht angewendet werden, das Magnetfeld ändert sich ja ständig.

Atome kann man so allerdings nicht fangen, da die Geschwindigkeit, mit der man das Atom detektieren und das Feld ändern müsste technisch schlichtweg nicht realisierbar ist.

6.4.2 Praxis

Der Betrieb dieser „Teilchenfalle“ ist sehr einfach. Hält man den Globus mit dem (geographischen) Nordpol nach oben in das silberne C, so kann man auch ohne eingestecktes Netzteil schon grob die Schwebeposition erfühlen. Sie ist ungefähr da, wo sich die magnetische und die Gewichtskraft gerade kompensieren ($F_{mag} \approx -F_G$). Steckt man nun das Netzteil ein und hält den Globus in die Nähe der Schwebeposition, dann beginnt die Leuchtdiode zu leuchten. In diesem Moment fängt auch die Spule des Elektromagneten an zu arbeiten. Das kann man als leichte Vibration spüren, wenn man den Globus in dieser Position festhält. Lässt man ihn nun langsam los, so beginnt er zu schweben. Durch leicht Störungen (anblasen, antippen oder mit einem Magneten) kann man den Globus ins Pendeln bringen und dann beobachten, wie sich die Lage langsam wieder stabilisiert. Ist die Störung jedoch zu groß, so kann sich das System nicht stabilisieren und der Globus fällt aus der Falle.

Hält man die beigegefügt Permanentmagnete in das C, so kann man sehen, dass auch mit diesen der Stabilisierungsmechanismus aktiviert werden kann (LED geht an. Dabei ist auf die Polarisation der Magneten zu achten). Wenn man nun die

Magnete noch beschwert indem man etwas zwischen sie klemmt, kann man auch diese schweben lassen.

7 Internetlinks

- Unter <http://www.quantenphysik-schule.de/Dokumente/jufopaul.pdf> beschreiben Bernhard Dörling und Peter Fritz, wie sie im Rahmen von „jugend forscht“ eine eigene Paulfalle konstruiert haben. (15.09.2004)
- Wie man sich selbst ein Levitron bauen kann, erklären Denis Faupel und Moritz Groba hier:
<http://users.aol.com/gykophys/levitron/levitron.htm> (1.10.2004)
- Ein anderes, auf dem Diamagnetismus beruhendes, Schwebeexperiment, das mit relativ einfachen Mitteln selbst gebaut werden kann, ist hier beschrieben:
<http://www.hcrs.at/DIAMAG.HTM> (8.10.2004)
- Hier erklärt die Physik-AG des Gymnasiums Korschenbroich, wie man ein Schwebeexperiment mit Lichtschranken aufbaut:
<http://www.physiktreff.de/material/menzel/skugel.htm> (1.10.2004)
- Eine umfassende Seite auf der man so gut wie alles zu GPS erfahren kann:
<http://www.kowoma.de/gps/> (11.10.2004)

Literaturverzeichnis

- [1] Verordnung des Kultusministeriums über die Jahrgangsstufen sowie über die Abiturprüfung an Gymnasien der Normalform und Gymnasien in Aufbauform mit Heim (Abiturverordnung Gymnasien der Normalform, NGVO), <http://www.leu.bw.schule.de/bild/NVGO.pdf>, (23.11.2004).
- [2] Bildungsstandards für Physik, herausgegeben vom Ministerium für Kultus, Jugend und Sport Baden-Württemberg, http://www.bildung-staerktemenschen.de/service/downloads/Bildungsstandards/Gymnasium/Gymnasium_Physik_Bildungsstandard.pdf, (23.11.2004).
- [3] H. Winter und H.W. Ortjohann, Simple demonstration of storing macroscopic particles in a "Paul trap", Am. J. Phys, 59 (9), 1991.
- [4] R.I. Thompson, T.J. Harmon und M.G. Ball, The rotating-saddle trap: a mechanical analogy to RF-electric-quadrupole ion trapping?, Can. J. Phys, 80, 1433-1448, 2002.
- [5] N.W. McLachlan, Theory and Applications of Mathieu functions, Oxford University Press, 1947.
- [6] L.D. Landau und E.M. Lifschitz, Lehrbuch der theoretischen Physik I, Akademie-Verlag, 1976.
- [7] R.F. Wuerker und R.V. Langmuir, J. Appl. Phys, 30, 342, 1959.
- [8] Hermann Haken, Atom- und Quantenphysik, Springer, 1996.
- [9] Compte rendus des Séances de la treizième Conférence des Poids et Mesures, Paris: Oktober 1967, Gauthier-Villars, 1968, S. 103.
- [10] P. Gill et al, Measurement Science and Technology, 14, 1174-1186, 2003.
- [11] Wolfgang Demtröder, Laserspektroskopie: Grundlagen und Techniken, Springer, 2000.
- [12] Herbert Walther, in: Das Atom in der Falle - Eine neue Uhr, Beiträge der akademischen Festveranstaltung zur Verleihung des Ernst Hellmut Vits-Preises, Aschendorff, 2000.

- [13] Samuel Earnshaw, On the nature of the molecular forces which regulate the constitution of the luminiferous ether, Trans. Cambridge Phil. Soc. 7, 97-112, 1842.
- [14] M.V. Berry, The Levitron: an adiabatic trap for spins, Proc. R. Soc. London A, 452, 1207-1220, 1996.
- [15] Martin Simon, Spin stabilized magnetic levitation, Am. J. Phys, 65 (4), 1997.
- [16] Tony Jones, Splitting the second - The story of atomic time, Institute of Physics Publishing, 2001.
- [17] Wolfgang Paul, Electromagnetic traps for charged and neutral particles, Rev. Mod. Phys, 62 (3), 1990.
- [18] Felix Voigt und Klaus Hirsch, Abgehoben - Schweben in elektromagnetischen Feldern, Phys. Unserer Zeit, 34 (5), 2003.
- [19] F.G. Major, The Quantum Beat - The Physical Principles of Atomic Clocks, Springer, 1998.
- [20] M.V. Berry und A.K. Geim, Of flying frogs and levitrons, Eur. J. Phys. 18, 307-313, 1997.