

Faszination Quantentheorie:

Die paradoxen Vorhersagen der Theorie und ihre experimentelle Bestätigung

Peter Schmüser, Department Physik, Univ. Hamburg

The “paradox” is only a conflict between reality and your feeling of what reality “ought to be”

Richard P. Feynman

1 Einleitung

Die Quantentheorie ist schon 90 Jahre alt, sie bereitet aber immer wieder große Verständnisschwierigkeiten, weil viele ihrer Aussagen unserer Erfahrung und Intuition widersprechen. Wenn man jedoch heute diese Theorie lehren oder lernen will, befindet man sich in einer viel besseren Lage als die Professoren und Studenten vor 90 oder auch 50 Jahren. Viele der von Bohr, Schrödinger, Heisenberg, Einstein und anderen vorgeschlagenen “Gedankenexperimente”, die damals heftig und kontrovers diskutiert wurden, sind durch enorme Fortschritte in der Experimentiertechnik inzwischen Wirklichkeit geworden. Messungen können dazu dienen, eine Entscheidung zwischen verschiedenartigen theoretischen Ideen und Interpretationen zu treffen, was um 1925 - 1930 kaum möglich war. Die fremdartigen, der Anschauung zuwiderlaufenden Aussagen der Theorie lassen sich Stück für Stück durch ausgeklügelte Experimente bestätigen: Teilchen-Welle-Komplementarität, das eigentümliche Verhalten von Bosonen und Fermionen, die Vorhersage der Antiteilchen, Vakuumfluktuationen, Verschränkung, Nichtlokalität der Quantenmechanik (auf dieses aktuelle Thema wird aus Platzgründen hier nicht eingegangen). Auf mich üben die modernen Experimente zur Quantenmechanik eine große Faszination aus, und ich hoffe, diese Faszination auch dem Leser vermitteln zu können [1]. Hinzu kommt, dass die Quantentheorie die Grundlage fast aller neuen Technologien ist und somit außer ihrem erkenntnistheoretischen Wert auch eine enorme praktische Bedeutung hat.

2 Teilchen-Welle-Dualismus

Die Quantentheorie ist nötig, um das Verhalten von Teilchen und Licht auf einer atomaren Größenskala zu beschreiben. Dabei treten fremdartige und paradoxe Phänomene auf, die unserer täglichen Anschauung eklatant widersprechen: Lichtwellen verhalten sich manchmal so, als ob sie aus Teilchen (Quanten) aufgebaut seien, und Teilchen zeigen gelegentlich Welleneigenschaften. Der Teilchen-Welle-Dualismus (oder etwas präziser: die Teilchen-Welle-Komplementarität) ist ein charakteristisches Merkmal der mikroskopischen Physik, das an einem Experiment verdeutlicht werden soll, für das es im Rahmen der klassischen Physik keine befriedigende Erklärung gibt: es ist das Doppelspaltexperiment. Historisch gesehen hat dies Experiment eine wesentliche Rolle gespielt um zu entscheiden, ob Licht aus Teilchen besteht oder eine Wellenerscheinung ist. Die Beobachtung von Interferenzstreifen und von vielen weiteren Beugungserscheinungen etablierte im 19. Jahrhundert die Vorstellung von der Wellennatur des Lichtes.

DOI: 10.3204/PUBDB-2026-00731

© 2015 Peter Schmüser. CC-BY-4.0 license

2.1 Wellenverhalten der Elektronen

Doppelspaltexperimente mit Elektronen sind heute keine “Gedankenexperimente” mehr, sondern können tatsächlich durchgeführt werden. Der Nachweis der Elektronen kann mit einem Film oder einem Pixeldetektor erfolgen. Was beobachtet man, wenn viele monoenergetische Elektronen die Doppelspaltapparatur durchlaufen? Überraschenderweise findet man genau wie beim Licht ein Interferenzbild mit hellen und dunklen Streifen, siehe Abb. 1 (“hell” bedeutet, dass viele Elektronen auftreffen). Es folgt daraus: Elektronen haben Welleneigenschaften. Mit einem Rastertunnel-

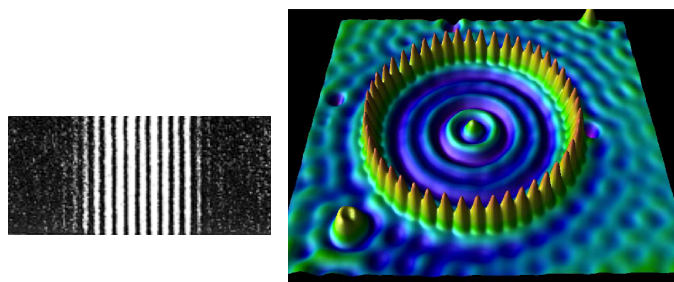


Abbildung 1: Die Wellennatur der Elektronen. Links: Doppelspalt-Interferenzen mit Elektronen. (G. Möllenstedt und Mitarbeiter, Universität Tübingen 1962). Rechts: Eine stehende Elektronenwelle innerhalb eines Ringes von 48 Eisenatomen auf einer Kupferoberfläche. (Wiedergabe mit freundlicher Genehmigung von D. Eigler, IBM-Forschungslabor Almaden. Image originally created by IBM Corporation.)

mikroskop kann man die Wellennatur der Elektronen direkt “sichtbar” machen.

Die Welleneigenschaften des Elektrons beschreiben wir durch eine Wellenfunktion $\psi(x)$. Dies ist eine komplexwertige Funktion, die in etwa der Amplitude einer klassischen Welle entspricht. Ihr Absolutquadrat $|\psi(x)|^2 = \psi^*(x)\psi(x)$ ist aber keineswegs identisch mit der beobachteten Intensitätsverteilung. Das ist schon aus dem Grund nicht möglich, weil Interferenzen auch dann auftreten, wenn sich immer nur ein Elektron zur Zeit im Apparat befindet. Das bedeutet, dass jedes Elektron mit sich selbst interferiert und nicht mit anderen Elektronen, ganz anders als Licht.

Was ist die Bedeutung der Wellenfunktion in der Quantenmechanik? Wir folgen hier der allgemein akzeptierten Wahrscheinlichkeitsinterpretation, die von Max Born eingeführt und insbesondere von Niels Bohr mit großer Überzeugung vertreten wurde (man spricht daher auch von der Kopenhagener Deutung der Quantentheorie). Die Wahrscheinlichkeit, das Elektron zwischen x und $x + dx$ auf dem Schirm zu finden, ist $\rho(x)dx = |\psi(x)|^2 dx$. Da das Elektron an irgendeiner Stelle auftreffen wird, muss die Integration über den gesamten Bereich den Wert 1 ergeben. Die Wellenfunktion ist in der Kopenhagener Deutung eine *Wahrscheinlichkeitsamplitude*.

Wie wird das Interferenzmuster aufgebaut?

Jedes Elektron trifft als ganzes, unteilbares Quant auf den Schirm und macht genau einen “Punkt” im Interferenzdiagramm, wobei die Wahrscheinlichkeit, den Treffer am Ort x zu erhalten, proportional zu $|\psi(x)|^2$ ist. Die Verteilung

wird durch viele Teilchen sukzessive aufgebaut. Schematisch

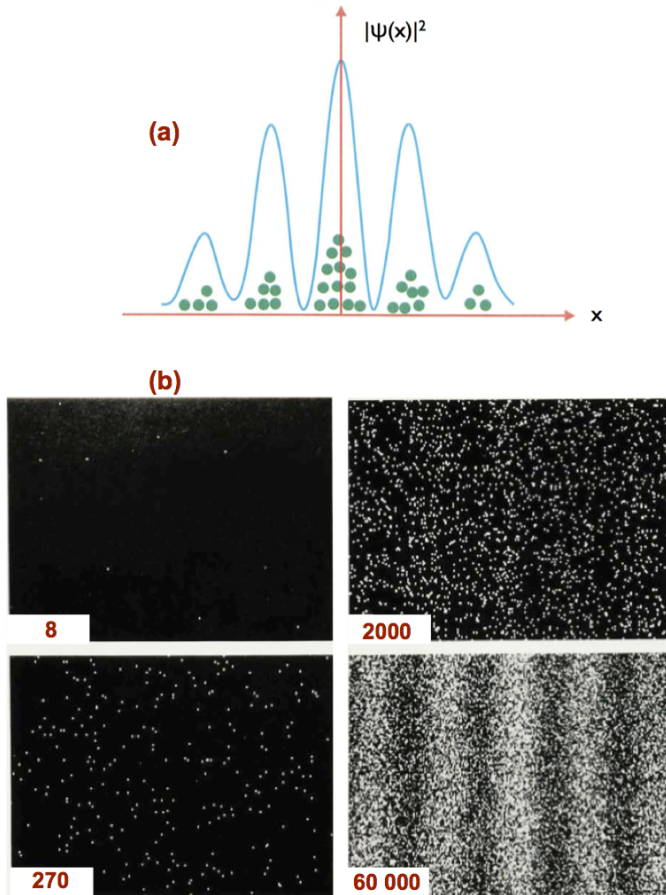


Abbildung 2: (a) Entstehung des Interferenzmusters beim Doppelspaltexperiment mit Elektronen. Jedes Elektron macht genau einen “Punkt” im Interferenzdiagramm, wobei die Wahrscheinlichkeit durch das Absolutquadrat der Wellenfunktion gegeben ist. (b) Beobachtung von Elektroneninterferenzen mit 8, 270, 2000 und 60 000 Elektronen [2]. Wiedergabe mit freundlicher Genehmigung von A. Tonomura und des Springer-Verlags.

ist dieser Vorgang in Abb. 2a dargestellt. Die vier Fotos in Abb. 2b sind ein experimenteller Beleg dafür, dass es in der Tat so abläuft. In den beiden ersten Bildern mit 8 oder 270 Teilchen sind überhaupt keine Interferenzstreifen erkennbar, dazu braucht man mehr als 10 000 Elektronen.

2.2 Welchen Weg wählt das Teilchen?

Die Elektronen treffen als praktisch punktförmige Teilchen auf den Beobachtungsschirm (fotografische Schicht oder Pixeldetektor), siehe Abb. 2. Demnach erscheint es vernünftig, sie generell als punktförmig anzusehen und folgende Behauptung aufzustellen: jedes Elektron fliegt entweder durch Spalt 1 oder durch Spalt 2, aber nicht gleichzeitig durch beide Spalte.

Diese einleuchtende Behauptung erweist sich als nicht haltbar. Sobald man eine Messung vornimmt, mit der geprüft werden kann, ob ein Elektron durch Spalt 1 oder Spalt 2 fliegt, verschwindet das Interferenzmuster.

Werner Heisenberg hat ein Gedankenexperiment konstruiert, in dem festgestellt werden könnte, durch welchen Spalt

das Elektron geht. Die Idee ist, kurzwelliges Licht zur Beobachtung zu benutzen. Der Haken ist jedoch, dass das Elektron bei der Streuung des Photons eine Richtungsänderung erleidet, die das Interferenzmuster stört, wobei die Stärke der Störung von der Wellenlänge des Photons abhängt.

(1) Bei sehr kleiner Wellenlänge $\lambda \ll d$ (d ist der Spaltabstand) kann man eindeutig den richtigen Spalt erkennen, die Störung ist jedoch stark, und es ergibt sich eine Addition der Wahrscheinlichkeitsdichten ohne Interferenzterm

$$\rho_{12}(x) = \rho_1(x) + \rho_2(x) = |\psi_1(x)|^2 + |\psi_2(x)|^2. \quad (1)$$

(2) Bei großer Wellenlänge ($\lambda \gg d$) kann man die beiden Spalte nicht mehr optisch auflösen, die Störung ist schwach, und man beobachtet das voll ausgeprägte Interferenzmuster

$$\begin{aligned} \rho_{12}(x) &= |\psi_1(x) + \psi_2(x)|^2 \\ &= \rho_1(x) + \rho_2(x) + \underbrace{2 \Re[\psi_1^*(x) \psi_2(x)]}_{\text{Interferenzterm}}. \end{aligned} \quad (2)$$

Durch eine sorgfältige Analyse des Doppelspaltexperiments gelangte Werner Heisenberg zu seiner berühmten Unbestimmtheitsrelation.

Ein interessanter neuer Gesichtspunkt wurde 1963 von Richard Feynman in seinen “Lectures on Physics” [3] diskutiert. Wenn die Wellenlänge des Photons und der Spaltabstand etwa gleich groß sind, kann man nur ungefähre Aussagen machen, wie dies auch in der Fuzzy-Logik der Fall ist. Nehmen wir an, die Wahrscheinlichkeit für die richtige Zuordnung des Spaltes sei w , die für die falsche Zuordnung $(1 - w)$, so ergibt sich eine Verteilung der Form (siehe Feynman, Band 3)

$$\begin{aligned} \rho_{12}(x) &= \left| \sqrt{w} \cdot \psi_1(x) + \sqrt{1-w} \cdot \psi_2(x) \right|^2 \\ &+ \left| \sqrt{w} \cdot \psi_2(x) + \sqrt{1-w} \cdot \psi_1(x) \right|^2. \end{aligned} \quad (3)$$

Für $w = 1$ tritt der obige Fall (1) ein, für $w = 0,5$ der Fall (2). Wenn aber beispielsweise $w = 0,75$ ist, was bedeutet, dass die Wahrscheinlichkeit für die richtige Zuordnung des Spaltes dreimal so groß ist wie Wahrscheinlichkeit für die falsche Zuordnung, ergibt sich ein Interferenzmuster mit reduziertem Kontrast. Dieser Übergangsbereich zwischen der klassischen Physik (gar keine Interferenz bei Teilchen) und der Quantenphysik (voll ausgeprägte Teilcheninterferenz) war den Gründungsvätern der Quantentheorie nicht bekannt und ist erst in den 1990er Jahren erforscht worden.

Der gleitende Teilchen-Wellen-Übergang

Der Übergangsbereich ist in einem Experiment [4] erforscht worden, das nur schematisch beschrieben werden soll (siehe Abb. 3). Die Teilchen sind hier Rubidium-Atome in einem hochangeregten Zustand, die Hauptquantenzahl beträgt $n = 51$. Die mit A_{51} bezeichneten Atome durchlaufen zwei hintereinander angeordnete Hochfrequenz-Resonatoren R_1 und R_2 , in denen die Atome durch Einstrahlung einer Frequenz von $f \approx f_0 = 51,099$ GHz vom Zustand ($n = 51$) in den Zustand ($n = 50$) überführt werden können. Am Detektor wird die Rate der A_{50} -Atome gemessen. Dies ist eine moderne Variante des Doppelspaltexperiments, denn man kann nicht unterscheiden, ob der Übergang $A_{51} \rightarrow A_{50}$ in R_1 oder R_2 erfolgt. Daher sind auch die typischen Interferenzen messbar, wenn man die eingestrahlte Hochfrequenz (HF) f ein wenig um f_0 moduliert.

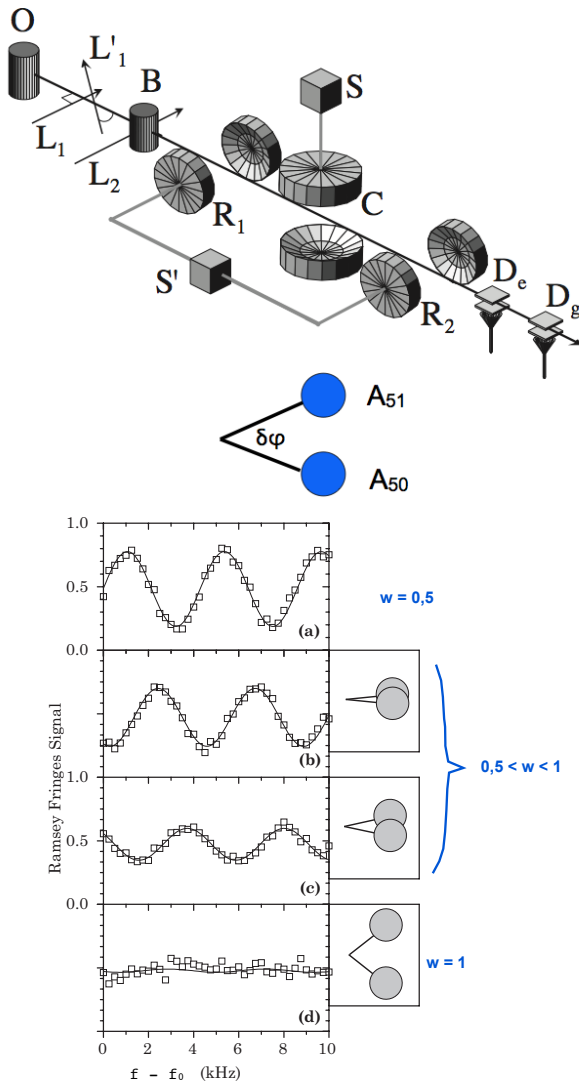


Abbildung 3: Oben: Aufbau zur Untersuchung des Grenzbereichs zwischen der Mikrowelt mit voll ausgeprägten Teilcheninterferenzen und der Makrowelt ohne Teilcheninterferenzen [4]. Darunter werden die Phasenverschiebungen $\pm\delta\varphi/2$ im Mikrowellenresonator C durch die Atome A_{51} und A_{50} gezeigt. Da nur wenige Mikrowellenphotonen in C vorhanden sind, gibt es einen Unschärfebereich für Amplitude und Phase des HF-Feldes, der durch die kreisförmigen Scheiben angedeutet wird. Unten: Die beobachteten Interferenzmuster für verschiedene präzise Identifikation des HF-Resonators R_1 oder R_2 , in dem der Übergang $A_{51} \rightarrow A_{50}$ erfolgt. Bild (a): gar keine Identifikation ($w = 0,5$) ergibt maximalen Kontrast. Bilder (b) und (c): teilweise Identifikation (partielle Überlappung der Unschärfe-Scheiben) ergibt Interferenzen mit reduziertem Kontrast. Bild (d): 100% sichere Identifikation (Scheiben getrennt) lässt die Interferenzen völlig verschwinden. Nachdruck mit freundlicher Genehmigung von J.-M. Raimond und S. Haroche. Figures adapted with permission from [4]. © 1996 American Physical Society.

Das Besondere an dem Aufbau ist nun ein supraleitender Mikrowellenresonator C zwischen R_1 und R_2 , der als Messinstrument wirkt und die hoch angeregten Atome A_{50} und A_{51} unterscheiden kann. Diese Atome sind riesig verglichen

mit Atomen im Grundzustand (der Radius der n -ten Bohrschen Bahn ist $r_n = n^2 r_1 \approx 125 \text{ nm}$ für $n = 50$) und besitzen ein sehr großes elektrisches Dipolmoment. Beim Flug durch den Mikrowellenresonator wirken sie wie eine kleine, schnell bewegte dielektrische Kugel, die den Resonator verstimmt und die Phase φ der Mikrowelle verschiebt. Die Frequenz des Resonators $f_C = f_0 + \Delta f$ wird so eingestellt, dass die Phase φ durch die A_{51} -Atome um $+\delta\varphi/2$ verschoben wird und durch die A_{50} -Atome um $-\delta\varphi/2$ (siehe Abb. 3). Der Grad der Verschiebung kann variiert werden, indem man Δf verändert. Damit erhält man ein einstellbares Messgerät zur mehr oder weniger deutlichen Unterscheidung der Zustände A_{50} und A_{51} . Übertragen auf unser Doppelspaltexperiment heißt das: man kann die Wahrscheinlichkeit w in Gl. (3) kontinuierlich zwischen $w = 0,5$ (keine Unterscheidungsmöglichkeit der beiden Spalte) und $w = 1$ (genaue Kenntnis des richtigen Spalts) variieren. Das Ergebnis des Experiments ist ebenfalls in Abb. 3 gezeigt. Man kann sehr deutlich die Veränderungen des Kontrasts im Interferenzmuster erkennen. In der Tat findet man maximalen Kontrast für $w = 0,5$ (keine Entscheidungsmöglichkeit), verminderten Kontrast für $0,5 < w < 1$ und Verschwinden des Musters für $w = 1$ (genaue Kenntnis des Spalts).

Ein hochinteressanter Aspekt dieses Experiments ist, dass es gar nicht darauf ankommt, die Phasenverschiebung im Mikrowellenresonator tatsächlich zu messen. Allein die Tatsache, dass dies im Prinzip möglich wäre, reicht aus, das Interferenzmuster zu beeinflussen oder sogar zu zerstören. Diese Erkenntnis geht weit über die Argumentation beim "Heisenberg-Mikroskop" hinaus: dort ist es der vom Photon auf das Elektron übertragene Rückstoßimpuls, der das Interferenzmuster zum Verschwinden bringt.

2.3 Neutron- und Molekül-Interferenzen

Mit Neutronen sind viele verschiedene Interferenzexperimente durchgeführt worden, siehe [5]. Aus einem Silizium-Einkristall wurde ein sog. Perfektkristall-Interferometer hergestellt (Abb. 4). Bragg-Reflexion wird benutzt, um die Teil-

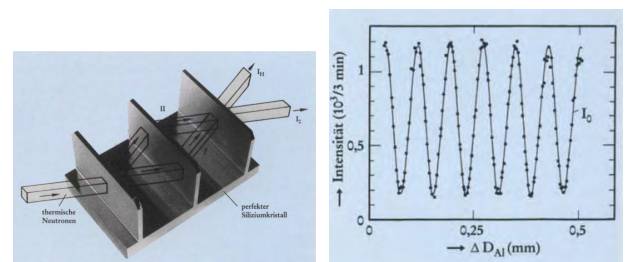


Abbildung 4: Das Perfektkristall-Neutroneninterferometer und die gemessenen Interferenzstreifen als Funktion der Dicke der Aluminiumplatte in einem Strahlweg [5]. Wiedergabe mit freundlicher Genehmigung von H. Rauch. © Wiley VCH-Verlag, reproduced with permission.

chenwelle in zwei Teilwellen aufzuspalten, die auf getrennten Pfaden durch das Interferometer geführt und wieder rekombiniert werden. Bringt man in einen Strahlengang eine dünne Aluminiumplatte, so wirkt diese ähnlich wie eine Glasplatte in einem Michelson-Interferometer. Die Neutronen haben verschiedene *de Broglie*-Wellenlängen in Aluminium und Luft, und daher ergibt sich ein Interferenzmus-

ter, bei dem die Intensität periodisch von der Dicke der Al-Platte abhängt (Abb. 4). In diesem Experiment entstehen die Neutronen durch Kernspaltung in einem Reaktor. Die Strahlintensität ist so niedrig, dass sich immer nur ein Neutron zur Zeit in der Apparatur befindet und das nachfolgende Neutron noch gar nicht entstanden ist. Wie bei den Elektronen beobachten wir also die Interferenz von Teilchen mit sich selbst und nicht mit anderen Teilchen. Bei Fermionen sind Interferenzen zwischen verschiedenen Teilchen auch gar nicht möglich, da sie aufgrund des Pauli-Prinzips verschiedene Quantenzustände einnehmen müssen und daher nicht kohärent sein können.

Als Fazit der Doppelspaltexperimente können wir festhalten:

Es ist nicht sinnvoll, von einer Bahnkurve des Elektrons im "Mikrokosmos" zu sprechen.

Aber wie steht es mit der Bahnkurve des Skifahrers in Abb. 5?



Abbildung 5: Die intellektuelle Herausforderung des Doppelspaltexperiments [5]. Cartoon von R.J. Buchelt, Abdruck mit freundlicher Genehmigung von H. Rauch.

Die Wellennatur der Materie ist für kleine Objekte - Elektronen oder Neutronen - durch Interferenzexperimente eindrucksvoll bestätigt worden, wie wir gerade gesehen haben. Bei großen Objekten - Billardkugeln oder Fußbällen - sind noch nie Interferenzen beobachtet worden. Ist die Quantenmechanik auf solche makroskopischen Objekte überhaupt anwendbar und macht es Sinn, von der *de Broglie*-Wellenlänge eines Fußballs zu sprechen? In der Kopenhagener Deutung wird dies verneint, dort wird eine deutliche Unterscheidung gemacht zwischen dem mikroskopischen Bereich der Elektronen, Atome und Moleküle, in dem die Gesetze der Quantentheorie anzuwenden sind, und dem makroskopischen Bereich unserer täglichen Erfahrung, der mit der klassischen Physik beschrieben wird. Eine spannende Frage ist: wie groß dürfen Objekte werden, um noch Interferenzen zu zeigen? Oder etwas allgemeiner gefragt: an welcher Stelle setzt der Übergang vom Quantenverhalten zum klassischen Verhalten ein (Engl. *quantum-to-classical transition*)?

Interferenzen mit "Fußbällen" kann man in der Tat beobachten, allerdings sind dies Miniaturfußbälle, die Fulleren-

Moleküle C_{60} . Experimentelle Resultate werden in Abb. 6 gezeigt. Diese Messung ist aus zwei Gründen wichtig. Ers-

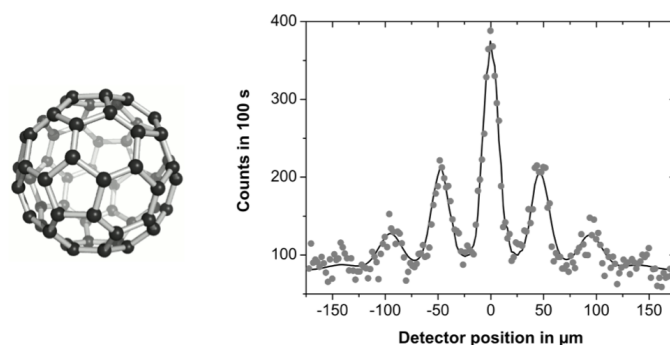


Abbildung 6: Links: Das Fulleren C_{60} ähnelt einem Fußball. Es ist innen hohl und hat einen Durchmesser von etwa 0,8 nm. Rechts: Beobachtete Interferenzen mit C_{60} -Molekülen der mittleren Geschwindigkeit 136 ± 3 m/s [6]. Abdruck mit freundlicher Genehmigung von M. Arndt, Univ. Wien. © 2003 American Association of Physics Teachers, reproduced with permission.

tens ist die *de Broglie*-Wellenlänge der Moleküle viel kleiner als ihr Durchmesser, und zweitens darf man die aus 60 C-Atomen bestehenden Bälle nicht als Punktteilchen ansehen, da die Atome innerhalb des Moleküls Schwingungen durchführen und Strahlung emittieren. Die in der Quelle hoch erhitzten Fulleren-Moleküle senden auf ihrem Flug vom Beugungsgitter zum Detektor 2 bis 3 Infrarot-Photonen aus. Deren Wellenlänge von einigen $10 \mu m$ und ist allerdings weit größer als der Spaltabstand im Gitter von $0,1 \mu m$. Daher ist es prinzipiell unmöglich, durch Messung dieser Photonen den Weg des C_{60} -Moleküls zu bestimmen, und aus diesem Grund wird das Interferenzmuster auch nicht beeinträchtigt.

2.4 Elektronenwellen in Magnetfeldern

Ein Magnetfeld übt auf bewegte geladene Teilchen die Lorentzkraft aus, die in der Elektrotechnik vielfach ausgenutzt wird. Es wird außerdem die *de Broglie*-Wellenlänge beeinflusst, doch dieser Effekt ist kaum bekannt und viel diffiziler. Man kann das Magnetfeld als Rotation eines Vektorpotentials darstellen $\mathbf{B} = \text{rot } \mathbf{A} = \nabla \times \mathbf{A}$. Während das Vektorpotential in der klassischen Elektrodynamik den Charakter einer Hilfsgröße hat, gewinnt es in der Quantentheorie eine zentrale Bedeutung: zusammen mit dem skalaren Potential bildet es die Wellenfunktion der Photonen. Die *de Broglie*-Wellenlänge eines Elektrons wird durch das Vektorpotential geändert

$$\lambda = \frac{2\pi\hbar}{m_e v - e A} . \quad (4)$$

Diese Beziehung wurde von Ehrenberg und Siday sowie von Aharonov und Bohm vorhergesagt und ist als *Aharonov-Bohm-Effekt* bekannt.

Wie kann man diese Vorhersage experimentell testen? Die Idee ist, ein geeignetes Doppelspaltexperiment durchzuführen (Abb. 7a). Hinter einem Schirm mit zwei Spalten befindet sich eine Solenoidspule, die so klein ist, dass sie zwischen den Spalten Platz hat. Das Vektorpotential

umgibt die Spule mit kreisförmigen Feldlinien. Auf dem Weg 1 läuft das Elektron antiparallel zu den $\mathbf{A}$ -Feldlinien, auf dem Weg 2 läuft es parallel dazu. Am Beobachtungsebene

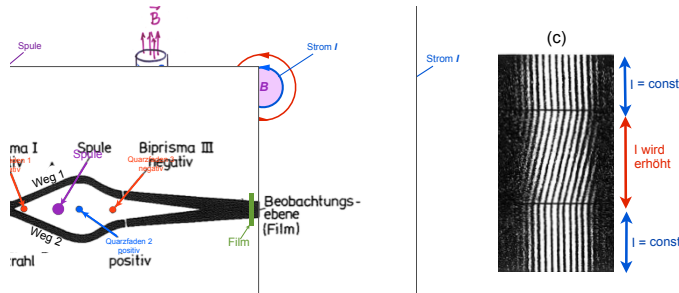


Abbildung 7: (a) Gedankenexperiment zur Beobachtung von Doppelspalt-Interferenzen mit einer Solenoid-Spule zur Erzeugung eines magnetischen Vektorpotentials. (b) Schema des Aufbaus von Möllenstedt und Bayh [7]. (c) Das beobachtete Interferenz-Muster [7] bei konstantem und bei gleichförmig anwachsendem Strom in der Spule. Der Film zur Aufnahme der Interferenzen wird in der vertikalen Richtung bewegt.

schirm haben die beiden Teilwellen deswegen eine Phasendifferenz, die proportional zum Vektorpotential A und damit zum Strom in der Spule ist. Dies wurde durch Messungen von Möllenstedt und Bayh bestätigt. Die benutzte Versuchsanordnung wird in Abb. 7b gezeigt. Ein Elektronenstrahl wird durch einen metallisierten Quarzfaden in zwei kohärente Teilstrahlen aufgespalten. Der Quarzfaden befindet sich auf negativem Potential und wirkt wie ein optisches Biprisma. Zwei weitere Quarzfäden mit passenden Potentialen sorgen dafür, dass die beiden Teilstrahlen auf einem Film zur Interferenz kommen. Hinter dem ersten Quarzfaden wird eine Solenoidspule (ca. $15 \mu\text{m}$ Durchmesser) angebracht, die aus Wolframdraht von $4 \mu\text{m}$ Dicke gewickelt ist. Die Messresultate werden in Abb. 7c gezeigt. Man beobachtet sehr scharfe Interferenzlinien. Durch Verändern des Spulenstroms I kann das Interferenzmuster kontinuierlich verschoben werden. Dies wird sichtbar gemacht, indem der fotografische Film synchron dazu bewegt wird, so dass die Interferenzstreifen schräg verlaufen. Die gemessene Phasenverschiebung stimmt quantitativ mit der theoretischen Vorhersage überein.

Der Aharonov-Bohm-Effekt kann dafür genutzt werden, magnetische Feldlinien direkt "sichtbar" zu machen, und zwar mit viel besserer Auflösung als durch Eisenfeilspäne. Die vom Vektorpotential verursachten Phasendifferenzen werden mit Hilfe der Elektronen-Holografie in beobachtbare Interferenzmuster umgewandelt. Abbildung 8 ist eine eindrucksvolle Demonstration dieser Methode.

Abschließend soll noch erwähnt werden, dass die durch den Aharonov-Bohm-Effekt modifizierte *de Broglie*-Beziehung von fundamentaler Bedeutung in der Quantenelektrodynamik ist. Aus ihr ergibt sich die korrekte Kopplung zwischen geladenen Teilchen und Photonen in der QED [8].

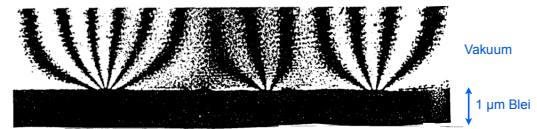


Abbildung 8: Sichtbarmachung magnetischer Feldlinien mit Hilfe der Elektronen-Holografie [2]. Gezeigt werden die Feldlinien oberhalb eines supraleitenden Bleifilms von $1 \mu\text{m}$ Dicke, der sich in einem schwachen Magnetfeld befindet. Wiedergabe mit freundlicher Genehmigung von A. Tonumura und des Springer-Verlags.

3 Bosonen und Fermionen

In der Quantentheorie sind identische Teilchen grundsätzlich ununterscheidbar, ganz anders als im täglichen Leben. Zwei exakt baugleiche Autos sind leicht unterscheidbar, wenn man einem von ihnen einen Farbfleck auf die Motorhaube malt. Die Elektronen in einem Eisenatom hingegen sind nicht unterscheidbar, wir können sie weder durch Farbflecke markieren noch mit einem Nummernschild versehen. Auch die Protonen im Atomkern sind nicht voneinander zu unterscheiden. Die grundsätzliche Ununterscheidbarkeit elementarer Teilchen hat weitreichende Konsequenzen für Vielteilchen-Wellenfunktionen: die Wahrscheinlichkeitsdichte muss invariant sein gegenüber der Vertauschung irgend zweier Teilchen.

Von Wolfgang Pauli wurde bewiesen, dass ein System von Teilchen mit ganzzahligem Spin (Bosonen) eine symmetrische Gesamtwellenfunktion haben muss, bei Teilchen mit halbzahligem Spin (Fermionen) dagegen eine antisymmetrische Gesamtwellenfunktion. Daraus resultieren zwei völlig konträre Verhaltensweisen. Bosonen (Photonen, Heliumatome etc.) gehen vorzugsweise alle in exakt den gleichen Quantenzustand. Fermionen (Elektronen, Protonen, Neutronen etc.) tun das genaue Gegenteil: jedes Fermion beansprucht seinen eigenen Quantenzustand und geht niemals in einen bereits besetzten Zustand.

3.1 Der Bose-Verstärkungseffekt

Stimulierte Emission und Laser

Die stimulierte Emission ist eine der wichtigsten Folgen des Bose-Prinzips und Grundlage des Lasers (LASER = "Light Amplification by Stimulated Emission of Radiation"). Ein angeregtes Atom kann auf zwei Weisen in den Grundzustand übergehen: durch spontane oder durch stimulierte Emission eines Photons (Abb. 9). Die spontane Emission läuft

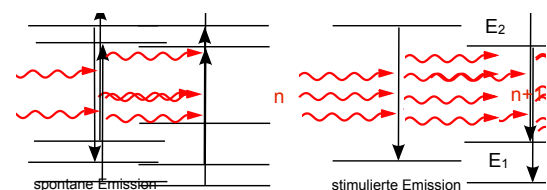


Abbildung 9: Spontane und stimulierte Emission von Strahlung.

"von selbst" ab (siehe aber Abschnitt 4.4), die stimulierte Emission wird durch ein externes Strahlungsfeld induziert. Wenn bereits n Photonen in einem bestimmten Quantenzu-

stand vorhanden sind (charakterisiert durch wohldefinierte Frequenz, Richtung, Polarisation und Phase der Lichtwelle), wird die Wahrscheinlichkeit, dass das nächste Photon durch stimulierte Emission ebenfalls in diesen Zustand geht, n -mal größer als die Wahrscheinlichkeit w_{spont} für die spontane Emission in einen anderen Zustand, bei dem z.B. die Lichtwelle eine andere Richtung hat.

$$w_{\text{stim}} = n w_{\text{spont}}. \quad (5)$$

Dieser *Bose-Verstärkungseffekt* ist entscheidend für die Funktion des Lasers und bewirkt, dass die Zahl der Photonen, die exakt parallel zur Laserachse laufen, exponentiell anwächst, bis schließlich eine Sättigung erreicht wird. Laserlicht ist nahezu monochromatisch, eng gebündelt, hat eine hohe Brillanz und ist optisch kohärent: all diese bemerkenswerten Eigenschaften sind eine direkte Konsequenz des Bose-Verstärkungseffekts.

Bose-Einstein-Kondensation

Mit dem Laser lassen sich Interferenzen sehr einfach demonstrieren. Es handelt sich dabei um die Überlagerung von makroskopischen Wellen, bestehend aus Millionen von Photonen, die alle den gleichen Quantenzustand einnehmen und kohärent sind. Im Bereich der Teilcheninterferenzen ist dies lange Zeit nicht möglich gewesen. Wie wir oben gesehen haben, sind die Elektroneninterferenzen an Doppelspalten, aber auch die Neutroneninterferenzen in Kristallen grundsätzlich als Interferenz jedes einzelnen Teilchens mit sich selber zu deuten. Kohärenz zwischen verschiedenen Elektronen kann man aufgrund des Pauli-Prinzips auch nicht erwarten.

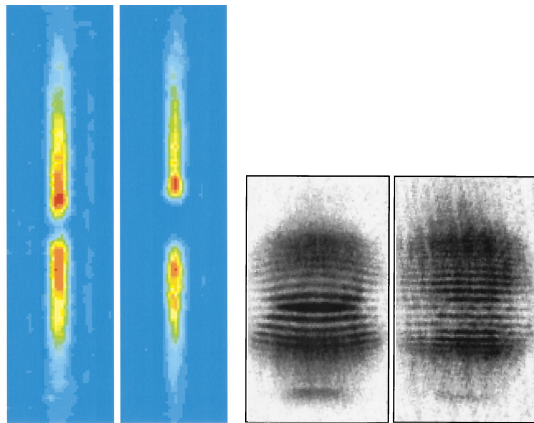


Abbildung 10: Links: Phasenkontrast-Aufnahmen von zwei getrennten Bose-Einstein-Kondensaten, bestehend aus Natrium-Atomen [9]. Rechts: beobachtete Interferenzen bei der Überlappung der beiden Kondensate. Bildwiedergabe mit freundlicher Erlaubnis von W. Ketterle und Science. © 1997 AAAS.

Anders ist dies bei Atomen mit ganzzahligem Gesamtdrehimpuls, die Bose-Teilchen sind und sehr wohl den gleichen Quantenzustand einnehmen dürfen. Seit langer Zeit ist ein solches System bekannt: es ist die supraflüssige Phase des Heliums, die bei Temperaturen unter 2,17 K auftritt. In den 1990er Jahren ist es gelungen, in Teilchenfallen bei extrem tiefen Temperaturen sehr viele Alkali-Atome in einem einzigen Quantenzustand zu kondensieren. Diese *Bose-Einstein-Kondensate* haben ungewöhnliche Eigenschaften. Wolfgang

Ketterle und Mitarbeitern gelang es, ein solches Kondensat aufzuteilen und die beiden Atomwolken zur Überlappung zu bringen (Abb. 10). Dabei sind erstmals makroskopische Materiewellen-Interferenzen beobachtet worden.

Supraleitung

Die Supraleitung ist ein drittes Phänomen, das auf dem Bose-Prinzip beruht. Träger des Suprastroms sind Paare von Elektronen mit entgegengesetztem Spin, die sog. *Cooperpaare*, die eine makroskopische Wellenfunktion im Supraleiter bilden.

3.2 Das Pauli-Ausschließungsprinzip

Um das Periodische System der Elemente und Regelmäßigkeiten in den Eigenschaften der Atome zu erklären, postulierte Wolfgang Pauli 1925 das Ausschließungsprinzip: *Zwei beliebige Elektronen in einem Atom müssen sich in mindestens einer der vier Quantenzahlen n , ℓ , m_ℓ und m_s unterscheiden.*

Eine Welt ohne Pauli-Prinzip? Lieber nicht!

In der Nachtszene von Goethes Tragödie *Faust 1* sagt Faust in seinem großen Eingangsmonolog *„Dass ich erkenne, was die Welt im Innersten zusammenhält“*. Dieses Streben ist das Leitmotiv der Elementarteilchenphysik geworden. Die Erforschung der kleinsten Bestandteile der Materie und der zwischen ihnen wirkenden Kräfte und Wechselwirkungen hat zu tiefgreifenden Erkenntnissen geführt. Wir wissen, dass Materie aus Elektronen, Protonen und Neutronen besteht. Wir kennen die elektrischen Kräfte, die Elektronen und Atomkerne zusammen halten, und die starken Wechselwirkungen, welche die Bindung der Nukleonen im Atomkern bewirken. Da ergibt sich die naheliegende Frage: wenn die elektrische Anziehung zwischen Protonen und Elektronen so groß ist, warum fallen die Elektronen nicht in den Kern, um zusammen mit den Protonen eine extrem dichte neutrale Materie zu bilden? Die klassische Physik hat keine Antwort auf diese Frage, wohl aber die Quantentheorie. Wenn wir versuchen, Elektronen auf das kleine Kernvolumen einzuschränken, müssten sie aufgrund der Heisenbergschen Unschärferelation eine so hohe kinetische Energie haben, dass die Coulombkraft nicht ausreichte, sie an die Protonen zu binden.

Jetzt kommt das Ausschließungsprinzip hinzu, das einen weiteren Hinderungsgrund darstellt, Materie auf ein winziges Raumvolumen zu konzentrieren. Wir machen ein Gedankenexperiment und nehmen an, die Elektronen seien Bosonen. Das Wasserstoff- und Heliumatom bleiben unverändert, aber bereits beim Lithium ergeben sich Unterschiede. Das dritte Elektron wird in die K-Schale gehen, falls Li bosonisch ist, es muss aber in die L-Schale gehen, wenn Li fermionisch ist (Abb. 11a). Als Fermion ist es weiter vom Kern entfernt und relativ locker gebunden (man braucht 25 eV, um He zu ionisieren, und nur 5 eV, um Li zu ionisieren). Lithium als fermionisches System hat also völlig andere chemische Eigenschaften als ein hypothetisches Li-Atom mit bosonischen Elektronen.

Die Unterschiede werden viel krasser bei schweren Atomen. Wenn man alle 82 Elektronen eines Bleiatoms in die K-Schale stecken könnte, wäre die Ausdehnung des Atoms viel geringer als sie in Wahrheit ist. Blei hätte eine unglaublich hohe Massendichte. Das Schalenmodell der Atome, das die

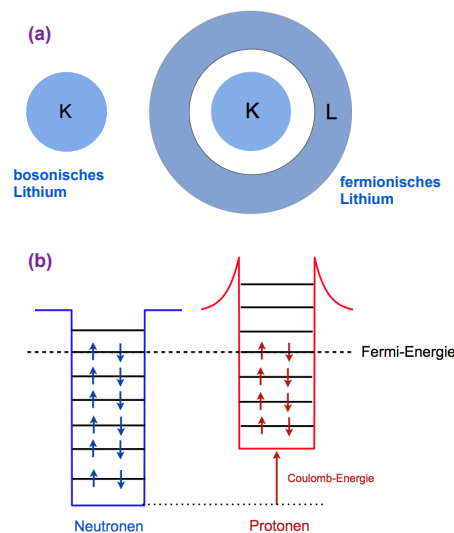


Abbildung 11: (a) Schema der Elektronenwolken für bosonisches und fermionisches Lithium. (b) Die Potentialtöpfe für Neutronen und Protonen im Atomkern. Der Topf der Protonen ist wegen der positiven elektrostatischen Energie angehoben. Die Niveaus werden von unten her mit jeweils zwei Neutronen bzw. Protonen besetzt, bis die Fermi-Energie erreicht ist. Das Proton-Neutron-Verhältnis ergibt sich aus der Forderung, dass das Fermi-Niveau in beiden Töpfen die gleiche Höhe haben muss. Das ist wie bei zwei Wasserbehältern, die durch eine kommunizierende Röhre verbunden sind.

Grundlage des Periodischen Systems der Elemente, der Chemie und der Biologie darstellt, würde zusammenbrechen.

Die Auswirkungen auf Atomkerne wären sogar noch dramatischer, wenn Protonen und Neutronen nicht dem Pauli-Prinzip gehorchten. Bei mittelschweren Kernen ist das Proton-Neutronverhältnis etwa 1:1, bei schweren Kernen ist es etwa 2:3. Eine qualitative Erklärung liefert ein Potentialtopfmodell der Kernkräfte, unter Berücksichtigung der elektrostatischen Abstoßung der Protonen (siehe Abb. 11b).

Auf jedes Energieniveau des Neutron-Potentialtopfs können jeweils zwei Neutronen mit entgegengesetztem Spin gesetzt werden, entsprechendes gilt auch für die Protonen. Ein Bleikern hat 126 Neutronen und 82 Protonen. Man muss die Niveaus des Neutrontopfes bis zu einer bestimmten Höhe besetzen, um all diese Teilchen unterzubringen, und entsprechend die Niveaus des Protontopfes. Was würde passieren, wenn die Nukleonen Bosonen wären? Dann würden sich alle 126 Neutronen auf das tiefste Energieniveau des Neutron-Potentialtopfs begeben. Der Protontopf liegt wegen der Coulombenergie viel höher. Nichts in der Welt könnte die Protonen daran hindern, sich durch Betazerfall oder Einfang von Hüllenelektronen in Neutronen umzuwandeln und dann auch auf das Grundniveau des Neutrontopfes zu fallen. Die Konsequenz: schwere Atomkerne würden ihre Ladung verlieren, und es gäbe keine schweren Atome.

In etwas lockerer Sprechweise können wir sagen, dass das Pauli-Prinzip zwei Fermionen mit parallelem Spin daran hindert, sich am gleichen Ort aufzuhalten. Das sieht so aus, als ob es eine abstoßende Kraft gäbe, die bei sehr kleinen Abständen wirksam wird. Eine solche Kraft ist jedoch unbe-

kannt, und wir müssen uns damit begnügen, dass es nur das abstrakte Ausschließungsprinzip ist, welches die Atome daran hindert, kleiner zu werden, als sie in Wirklichkeit sind. Trotzdem können wir in Abwandlung des Faust-Monologs den Wunsch aussprechen:

Dass ich erkenne, was die Welt im Innersten "auseinanderhält"

Wir dürfen gespannt darauf sein, ob jemand eines Tages eine dynamische Erklärung für das Ausschließungsprinzip findet. Auch für den *modernen Faust* gibt es noch viel zu erforschen.

Eine anschauliche, auf der täglichen Erfahrung oder der klassischen Physik beruhende Erklärung des Pauli- oder Bose-Prinzips ist mir nicht bekannt. Versuche, eine solche Erklärung zu finden, sind vermutlich auch zwecklos, weil die zwei Typen identischer Teilchen sich so konträr verhalten.

4 Relativistische Quantentheorie

Die Vereinigung von zwei der wichtigsten Theorien des frühen 20. Jahrhunderts, der Speziellen Relativitätstheorie und der Quantenmechanik, zu einer relativistischen Quantentheorie hat sich als sehr fruchtbar erwiesen und zu radikal neuen Einsichten geführt [10]. Die erste relativistische Wellengleichung war die Klein-Gordon-Gleichung, eine Differentialgleichung zweiter Ordnung in x , y , z und t . Sie hat zwei Typen von Lösungen mit unterschiedlichen Zeitabhängigkeiten: $\psi_p \sim \exp(-i\omega t)$ und $\psi_n \sim \exp(+i\omega t)$. Bei Anwendung des Energieoperators $\hat{E} = i\hbar \frac{\partial}{\partial t}$ auf ψ_p erhält man einen positiven Wert für die relativistische Gesamtenergie eines freien Teilchens (Ruhe-Energie plus kinetische Energie), was ja auch den Erwartungen entspricht. Beim zweiten Typ ψ_n dagegen werden die Eigenwerte des Energieoperators negativ. Dies unsinnige Ergebnis blieb lange Zeit unverständlich.

4.1 Dirac-Gleichung und Antiteilchen

Probleme mit negativen kinetischen Energien gibt es bei der Schrödinger-Gleichung nicht, die von der ersten Ordnung in der Zeit ist. Daher konstruierte Paul Dirac eine relativistische Gleichung, die ebenfalls von der ersten Ordnung in der Zeit ist. Wegen der relativistischen Invarianz muss sie dann auch von der ersten Ordnung in den Ortskoordinaten sein. Die Dirac-Gleichung für freie Elektronen lautet

$$i\hbar \frac{\partial \psi}{\partial t} = -i\hbar c \left(\alpha_1 \frac{\partial \psi}{\partial x} + \alpha_2 \frac{\partial \psi}{\partial y} + \alpha_3 \frac{\partial \psi}{\partial z} \right) + \beta m_e c^2 \psi. \quad (6)$$

Die Koeffizienten α_j und β sind nicht einfach komplexe Zahlen, sondern 4×4 -Matrizen, und die Wellenfunktionen sind vierkomponentige Spaltenvektoren, die man Dirac-Spinoren nennt [11].

Dirac erkannte schnell, dass auch seine Gleichung zwei Typen von Lösungen mit positiven und negativen Eigenwerten des Energieoperators hat. Als herausragender theoretischer Physiker akzeptierte er das Unvermeidliche und suchte nach einer neuen Erklärung. Er machte die kühne Annahme, dass die Zustände mit negativer Energie tatsächlich existieren, aber normalerweise sämtlich mit Elektronen besetzt sind. Nach dieser Deutung ist der Grundzustand, oft

auch das *Vakuum* genannt, nicht leer, sondern enthält unendlich viele Elektronen mit negativer Energie. Das Pauli-

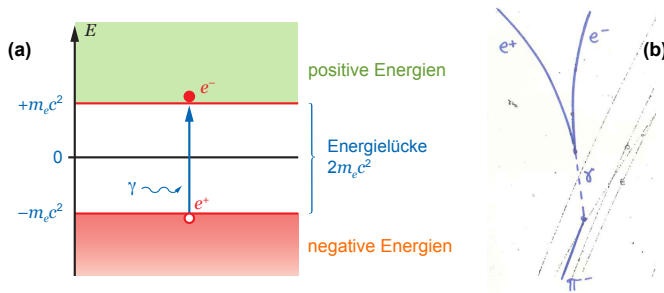


Abbildung 12: **(a)** Dirac-Bild der Antiteilchen. Die positiven Energieniveaus der Elektronen mit $E \geq +m_e c^2$ befinden sich in dem grün schattierten Bereich, die negativen Energieniveaus mit $E \leq -m_e c^2$ im rot schattierten Bereich. Die negativen Niveaus sind sämtlich mit Elektronen besetzt. Ein γ -Quant der Energie $E_\gamma > 2m_e c^2$ kann ein Elektron e^- von einem Zustand negativer Energie in einen Zustand positiver Energie anheben. Das verbleibende ‐Loch‐ im See der Elektronen negativer Energie verhält sich wie eine positive Ladung mit positiver Energie. Dies ist das Positron e^+ . **(b)** Erzeugung eines Elektron-Positron-Paars in einer Flüssig-Wasserstoff-Blasenkammer. Ein π^- -Meson trifft auf ein ruhendes Proton und löst die Reaktion $\pi^- + p \rightarrow \pi^0 + n$ aus. Das neutrale π^0 -Meson zerfällt in zwei γ -Quanten, von denen eines ein Elektron-Positron-Paar erzeugt. In der Blasenkammer sind nur geladene Teilchen sichtbar.

Prinzip verbietet den Übergang eines ‐normalen‐ Elektrons von seinem Zustand positiver Energie in ein negatives Energieniveau, so dass man normalerweise von den vielen negativen Niveaus nichts merkt. Das ist ein Glück, denn sonst würden die atomaren Elektronen in negative Energiezustände übergehen und die gesamte Materie zum Verschwinden bringen.

Durch ein γ -Quant mit einer Energie von mehr als $2m_e c^2$ könnte jedoch ein Elektron von einem negativen auf ein positives Energieniveau angehoben werden (siehe Abb. 12a). Das verbleibende Loch im ‐See‐ der Elektronen mit $E < 0$ sollte sich wie ein Teilchen mit positiver Ladung und positiver Energie verhalten. Aufgrund dieser Überlegungen hat Paul Dirac die Existenz von *Antiteilchen* vorhergesagt und damit eine der revolutionärsten Ideen der theoretischen Physik hervorgebracht. Das Antiteilchen des Elektrons, das Positron, wurde 1932 von Anderson in der Höhenstrahlung entdeckt. Es hat die gleiche Masse m_e , aber die entgegengesetzte Ladung $+e$. Die Blasenkammeraufnahme in Abb. 12b zeigt die Erzeugung eines Elektron-Positron-Paars durch γ -Strahlung.

Das Dirac-Bild ist später erfolgreich auf Halbleiter übertragen worden. Ein nicht dotierter Halbleiter mit komplett gefülltem Valenzband und leerem Leitungsband ist ein Isolator. Durch ein Photon, dessen Energie größer als die Energielücke des Halbleiters ist, kann ein Elektron-Loch-Paar erzeugt werden. Das Loch im Valenzband entspricht dem Positron (in [11] wird erklärt, weshalb Löcher positive Ladungsträger sind).

In der Teilchenphysik ist eine alternative Deutung der Wellenfunktionen vom Typ ψ_n vorteilhaft, die von Feynman

stammt. Die konjugiert-komplexen Funktionen ψ_n^* kann man als Wellenfunktionen der Positronen interpretieren, und diese haben einen positiven Energie-Eigenwert [8].

4.2 Spin und magnetisches Moment

Die Dirac-Gleichung macht noch andere tiefgehende Aussagen über Eigenschaften, die in der nichtrelativistischen Quantenmechanik nicht begründet werden können. Dies betrifft zunächst den Spin (Eigendrehimpuls) des Elektrons von $\hbar/2$. Die erste und zweite Komponente eines Dirac-Spinors beschreiben das Elektron mit seinen zwei Spineinstellungen, die dritte und vierte beschreiben (nach dem Übergang zum konjugiert-Komplexen) das Positron mit seinen zwei Spineinstellungen. Die mathematischen Details findet man in [8].

Ein weiteres, selbst von Dirac nicht erwartetes Resultat ist das anomale magnetische Moment des Elektrons. Aus der Dirac-Gleichung mit elektromagnetischem Potential folgt ohne jede Zusatzannahme, dass das magnetische Moment den Wert

$$\mu_e = -\frac{e\hbar}{2m_e} \equiv -\mu_B \quad (7)$$

hat, in sehr guter Übereinstimmung mit dem Experiment. Dieser Wert ist doppelt so groß wie erwartet, wenn man sich das Elektron als geladene Kugel vorstellt, die mit dem Drehimpuls $\hbar/2$ um die eigene Achse rotiert. Das magnetische Moment dieser Kugel sollte nämlich ein halbes Bohr-Magneton betragen. Das ‐magneto-mechanische‐ Modell des Elektrons ergibt also falsche Resultate. Ein anschauliches Bild des ‐Dirac-Elektrons‐ existiert meines Wissens nicht.

4.3 Antisymmetrie der Dirac-Spinoren

Bei einer Rotation des Koordinatensystems um den Winkel $2\pi \simeq 360^\circ$ geht eine Schrödinger-Wellenfunktion in sich selbst über. Dies erscheint so selbstverständlich, dass man man sich darüber wundert, dass es auch anders sein könnte. Aus der Dirac-Theorie folgt jedoch die eigentümliche Vorhersage, dass die Spinoren bei einer 2π -Rotation in ihr Negatives übergehen. Erst nach einer 4π -Rotation, also nach zwei kompletten Umdrehungen, gehen sie in sich selbst über.

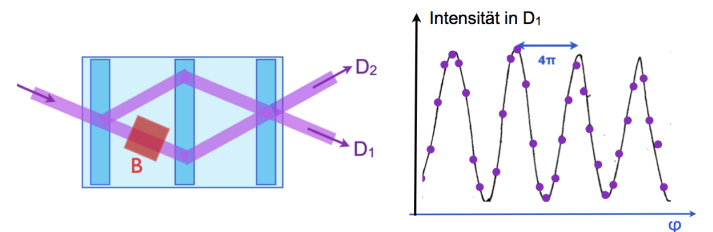


Abbildung 13: Untersuchung der Spin-Präzession in einem Perfektkristall-Interferometer. In einem Interferometerarm durchläuft die Neutronenteilwelle ein Magnetfeld B . Gezeigt wird die Intensität im Detektor D_1 als Funktion des Präzessionswinkels φ . Bei einer Rotation um 2π geht ein Interferenzmaximum in ein Minimum über. Gezeichnet mit freundlicher Genehmigung nach einer Abbildung von H. Rauch.

Mit Neutron-Interferenzen ist diese befremdliche Antisymmetrie verifiziert worden. Die experimentellen Daten und eine schematische Darstellung des verwendeten Perfektkristall-Interferometers werden in Abb. 13 gezeigt. In einem der Interferometerarme durchläuft das Neutron ein transversales Magnetfeld, in welchem das magnetische Moment des Neutrons eine Präzessionsbewegung durchführt. Der Rotationswinkel φ des magnetischen Momentenvektors, der natürlich identisch mit dem Rotationswinkel des Spinvektors ist, kann aus dem Feld B und der Länge ℓ des Magneten berechnet werden. Das gemessene Interferenzmuster beweist, dass eine 2π -Rotation ein Interferenzmaximum in ein Minimum überführt. Erst bei einer 4π -Rotation gehen Maxima in Maxima über und Minima in Minima.

4.4 Die komplexe Natur des Vakuums

Im 19. Jahrhundert wurde der *Äther* postuliert, da man sich transversale elektromagnetische Wellen ohne irgendein Trägermedium nicht vorstellen konnte. Der Äther wurde als fest im Raum verankert angesehen, so dass Bewegungen im Raum einen Mitführungseffekt des Lichts hervorrufen müssten. Im Michelson-Morley-Experiment wurde vergeblich nach diesem Mitführungseffekt gesucht. Der experimentelle Befund, dass Licht unabhängig von seiner Richtung relativ zur Erdbahn immer die gleiche Geschwindigkeit c hat, war der Ausgangspunkt der Speziellen Relativitätstheorie. Albert Einstein hatte damit dem Konzept des Äthers ein Ende bereitet, und nach 1905 wurde das "Vakuum", der materie- und feldfreie Raum, als vollkommen "leer" angesehen.

Ironischerweise änderte die zwei Jahrzehnte später entwickelte relativistische Quantenfeldtheorie die Sichtweise, und das Vakuum wurde allmählich wieder "bevölkert" [13]. Im Jahr 1927 publizierte Paul Dirac eine bahnbrechende Arbeit mit dem Titel *The Quantum Theory of the Emission and Absorption of Radiation*. Das elektromagnetische Strahlungsfeld wird als ein Ensemble von unendlich vielen quantisierten harmonischen Oszillatoren behandelt. Im "Vakuum", definiert als der tiefste Energiezustand des Strahlungsfeldes, existieren immer noch die Nullpunktsschwingungen für alle erlaubten Frequenzen. Die Energie-Zeit-Unschärferelation erlaubt einem Oszillator, für eine sehr kurze Zeit in den Anregungszustand E_1 überzugehen und danach in den Grundzustand E_0 zurückzukehren. Die bei diesen Vakuumfluktuationen auftretenden *virtuellen Photonen* haben vielfältige Auswirkungen.

Spontane Emission: gar nicht so spontan

Die Schrödinger-Gleichung hat einen prinzipiellen Mangel: sie ist außerstande, die spontane Emission von Strahlung zu erklären. Wenn kein externes Strahlungsfeld vorhanden ist, sind der Grundzustand und alle angeregten Zustände eines Atoms *stationär*. Ein Elektron, das sich auf dem 2p-Niveau des H-Atoms befindet, wird für ewige Zeiten dort bleiben, es denkt nicht daran, in den 1s-Zustand zu springen und dabei ein Photon zu emittieren. Damit scheint eines der wichtigsten Postulate des Bohrschen Atommodells unerfüllbar zu sein.

Die Quantenfeldtheorie ist imstande, diesen prinzipiellen Mangel zu beheben. Es sind die bei Vakuumfluktuationen auftretenden virtuellen Photonens, die die angeregten Atome zu ihren "spontanen" Übergängen veranlassen.

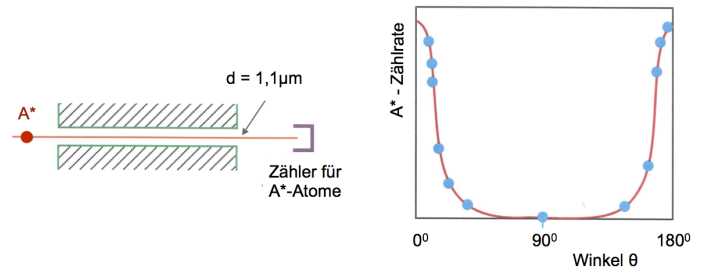


Abbildung 14: Unterdrückung der spontanen Emission durch "Abschirmung" der Vakuumfluktuationen. Links wird ein Schema des Experiments gezeigt, rechts ist die Zählrate der angeregten Cs-Atome A^* als Funktion des Winkels θ zwischen dem $\mathbf{E}$ -Vektor und den Platten aufgetragen. Gezeichnet mit freundlicher Genehmigung nach einer Skizze von S. Haroche.

Wie kann man diese Aussage testen? Die Idee ist, angeregte Atome in eine Art "Faraday-Käfig" zu sperren, der die Vakuumfluktuationen abschirmt, und dann zu überprüfen, ob sie länger als normal im Anregungszustand bleiben. Ein Experiment dieser Art ist von S. Haroche und Mitarbeitern mit einem Cäsium-Atomstrahl durchgeführt worden. Als Abschirmung dienten zwei sehr ebene parallele Metallplatten oberhalb und unterhalb des Atomstrahls (Abb. 14), deren Abstand $d = 1,1 \mu\text{m}$ kleiner als die halbe Wellenlänge der beim Übergang $5d \rightarrow 6s$ emittierten Strahlung war ($\lambda = 3,5 \mu\text{m}$). Eine elektromagnetische Welle mit $\lambda = 3,5 \mu\text{m}$ kann sich im Raum zwischen den Platten ungehindert ausbreiten, sofern ihr elektrischer Vektor senkrecht auf den Platten steht. Ist aber der $\mathbf{E}$ -Vektor parallel zu den Platten, so wird die Welle exponentiell abgeschwächt und dringt nur wenige μm in den Zwischenraum ein. Mit einem Detektor wurde die Zahl der angeregten Cs-Atome als Funktion des Winkels θ zwischen dem $\mathbf{E}$ -Vektor und den Platten gemessen. Bei paralleler Ausrichtung ($\theta = 0^\circ, 180^\circ$) wurde eine hohe Zählrate gemessen, aus der man eine 13-fache Verlängerung der natürlichen Lebensdauer herleiten konnte. Bei senkrechter Ausrichtung ($\theta = 90^\circ$) erreichte kein einziges angeregtes Atom den Detektor.

Es sind auch Experimente mit angeregten Atomen in metallischen Hohlräumen (*cavities*) durchgeführt worden, die Vakuumfluktuationen jeglicher Polarisation abschirmten. Der spontane Zerfall konnte komplett zum Erliegen gebracht werden [12].

Die Vakuumfluktuationen haben noch andere messbare Konsequenzen: die Energieaufspaltung zwischen den $2s_{1/2}$ - und $2p_{1/2}$ -Niveaus des H-Atoms (*Lamb shift*) und die geringfügige Abweichung des magnetischen Moments des Elektrons vom Bohr'schen Magneton. Beide Effekte sind mit hoher Präzision gemessen und berechnet worden, die Übereinstimmung zwischen Messungen und QED-Rechnungen ist fantastisch.

Der Higgs-Mechanismus

Die nächste Population des physikalischen Vakuums kam mit der Vereinheitlichung der elektromagnetischen und schwachen Wechselwirkungen in den 1970er Jahren. Das *Standard-Modell der elektro-schwachen Wechselwirkung* ist theoretisch sehr gut fundiert und beruht auf einer verallgemeinerten Eichinvarianz (eine Einführung findet man in

[8]). Die Vorhersagen dieser Theorie sind an den Elektron-Positron-Collidern mit hervorragender Präzision bestätigt worden, es gibt jedoch ein tiefliegendes Problem: die Eichinvarianz ist nur gültig, wenn alle Feldquanten und Teilchen masselos sind. Für die Photonen trifft das zu, aber die Feldquanten der schwachen Wechselwirkung, die Z^0 - und $W^\pm$ -Bosonen, haben extrem hohe Massen, sie sind fast hundertmal so schwer wie das Proton. Massive Feldquanten erzeugen ein Yukawa-Potential

$$\Phi(r) \sim \frac{e^{-\mu r}}{r} \quad \text{mit} \quad \mu = \frac{Mc}{\hbar}, \quad (8)$$

dessen Reichweite für Massen von $90 \text{ GeV}/c^2$ extrem kurz ist, etwa $1/1000$ des Protonradius.

Bis heute ist nur ein einziger Weg bekannt, den Z^0 - und $W^\pm$ -Bosonen eine Masse zu geben und gleichzeitig die Eichinvarianz zu bewahren: dies ist der berühmte Higgs-Mechanismus, benannt nach den schottischen Physiker Peter Higgs. Die grundlegende Idee ist folgende: die schwache Wechselwirkung hat "an sich" eine unendliche Reichweite und Feldquanten mit Masse null, genau wie die elektromagnetische Wechselwirkung, aber die schwachen Kräfte werden durch ein Hintergrundfeld abgeschirmt, so dass die beobachtete Reichweite auf $2 \cdot 10^{-18} \text{ m}$ reduziert wird. Als Folge dieser Abschirmung wird den a priori masselosen Feldquanten Z^0 und $W^\pm$ eine "effektive Masse" verliehen [14].

Ein Abschirmvorgang dieser Art ist der Meissner-Ochsenfeld-Effekt in Supraleitern, der dafür sorgt, dass ein externes Magnetfeld nur in eine dünne Oberflächenschicht eindringen kann und im Inneren des Supraleiters verschwindet. In dieser nur 50 nm dicken Schicht fließen widerstands-freie Supraströme, die das Magnetfeld exponentiell abklingen lassen.

Auch die Massen der Quarks und Leptonen werden auf den Higgs-Mechanismus zurückgeführt. Das Higgs-Feld muss überall im Raum vorhanden sein, denn auch in der Sonne und anderen Sternen verlaufen Prozesse der schwachen Wechselwirkung mit den in irdischen Labors gemessenen Wahrscheinlichkeiten.

In den Experimenten am Large Hadron Collider des Forschungszentrums CERN wurde ein neues Teilchen gefunden, dessen Masse von $125 \text{ GeV}/c^2$ in dem Bereich liegt, wo man aufgrund zahlreicher früherer Experimente und theoretischer Analysen das Higgs-Teilchen erwartet. Vermutlich handelt es sich dabei wirklich um das Higgs-Teilchen, doch es sind noch weitere Messungen nötig, um diese Vermutung abzusichern.

5 Der Quanten-Klassik-Übergang

5.1 Schrödingers Katze

Um gewisse absurde Aspekte der Quantenmechanik zu demonstrieren, präsentierte Erwin Schrödinger 1935 ein Gedankenexperiment, in dem ein makroskopisches Objekt (die Katze) mit einem mikroskopischen System (angeregtes Atom) verschränkt ist. Beide befinden sich in einem verschlossenen Kasten. Beim Übergang in den Grundzustand emittiert das Atom ein Photon, welches mittels einer geeigneten Apparatur Giftgas freisetzt und dadurch die Katze tötet.

Das Atom befindet sich in einer Superposition des Anregungszustand ψ_a und des Grundzustands ψ_g , und erst nach Öffnen des Kastens entscheidet es sich, ob das Atom in ψ_a oder ψ_g ist. Wenn die quantenmechanischen Regeln auch für die Katze gelten, sollte sie sich in einer Superposition von lebendig und tot befinden, wobei die Entscheidung, ob die Katze lebendig oder tot ist, solange offen gehalten wird, bis ein Beobachter den Deckel des Kastens öffnet. Diese absurde Situation wird natürlich nie beobachtet.

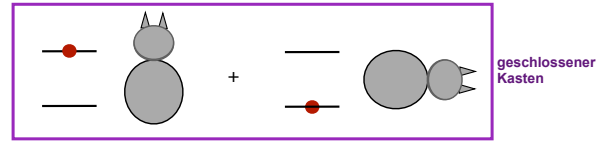


Abbildung 15: Die Idee der Schrödingers Katze. Die lebende Katze ist mit dem angeregten Atom verschränkt, die tote Katze mit dem Atom im Grundzustand.

Andererseits muss man bedenken, dass makroskopische Objekte wie Katzen aus Atomen aufgebaut sind, die jedes für sich der Quantenmechanik gehorchen und auch der Verschränkung unterworfen sein können. Es bleibt demnach die wichtige Frage: warum und auf welche Weise verschwinden die Skurrilitäten der Quantenmechanik in großen Systemen?

5.2 Dekohärenz

In den letzten 25 Jahren sind große Fortschritte im theoretischen Verständnis großer Systeme gemacht worden. Dort tritt eine Dekohärenz auf, die die quantenmechanischen Superpositionseffekte sehr rasch zum Verschwinden bringt. Die "Umgebung" hat einen signifikanten Einfluss auf die Dekohärenzzeit (d.h. die Zeitspanne, innerhalb derer die quantenmechanische Kohärenz verschwindet). Man kann eine Katze nicht völlig von ihrer Umgebung isolieren und beispielsweise in einen Vakuumbehälter stecken, denn dann würde sie auch ohne die Giftkapsel sterben. Daher treffen die Moleküle der Luft auf die Katze und reduzieren die Interferenzeffekte. Man kann die Katze auch nicht in einen Flüssig-Helium-Kryostaten sperren, um die thermische 300-Kelvin-Strahlung zu unterdrücken, und so treffen die zahlreichen Photonen dieser Strahlung auf das Tier. Durch die Wechselwirkung mit der Umgebung sickert unweigerlich Information über das Quantensystem in die Umgebung mit der Konsequenz, dass die Quantenkohärenz beeinträchtigt wird. Genauere Analysen zeigen, dass die Dekohärenzzeit ganz rapide mit wachsender Systemgröße absinkt. In makroskopischen Systemen sind die störenden Prozesse so effektiv, dass die Quantenkohärenz in unmessbar kurzer Zeit verschwindet. Was ist die Konsequenz?

Eine Katze kann nicht in einem verschränkten Zustand existieren, vielmehr ist sie - wie in der klassischen Welt - entweder lebendig oder tot, aber nicht beides gleichzeitig.

Wir können feststellen, dass die Quantenmechanik nach heutigen Erkenntnissen das Schrödingersche Gedankenexperiment unbeschadet "überstanden" hat.

Messung einer Dekohärenzzeit

In einem hervorragenden Artikel [15] beschreibt Serge Haroche ein Experiment zur Dekohärenz in einem *mesoskopischen*

schen System. Mesoskopische Systeme liegen in ihrer Größe zwischen mikroskopischen Systemen (Atomen) und makroskopischen Systemen (Katzen). Es wird ein “Schrödinger-Kätzchen” (*Schrödinger kitten*) realisiert, das aus wenigen verschränkten Mikrowellenphotonen in einem supraleitenden Mikrowellenresonator besteht. Das Experiment ist sehr anspruchsvoll, sowohl im theoretischen Hintergrund als auch in der praktischen Realisierung, so dass wir nur sehr pauschal darauf eingehen können. Die benutzte Apparatur wurde bereits in Abb. 3 gezeigt. Ein erstes Rubidium-Atom erzeugt einen verschränkten “Schrödinger-Katzen-Zustand” in dem supraleitenden Resonator C, ein zweites, mit zeitlicher Verzögerung folgendes Atom tastet ab, ob die Verschränkung noch vorhanden ist oder ob inzwischen ein Übergang in ein inkohärentes Gemisch von Zuständen eingetreten ist.

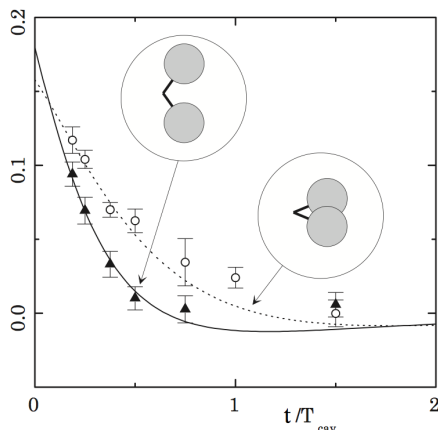


Abbildung 16: Gemessene zeitliche Abnahme der Kohärenz in einem “Schrödinger-Kätzchen-Experiment” [4], [15]. Die Dekohärenzzeit beträgt etwa $50 \mu\text{s}$. Nachdruck mit freundlicher Genehmigung von J.-M. Raimond. Figure adapted with permission from [4]. © 1996 by the American Physical Society.

Durch Variation der Verzögerung kann die Dekohärenzzeit gemessen werden, sie genügt der Formel

$$t_{\text{decoh}} = T_{\text{cav}}/\bar{n},$$

wobei $T_{\text{cav}} = 160 \mu\text{s}$ die Dämpfungszeitkonstante der Resonatorschwingung ist und $\bar{n}$ die mittlere Zahl der Mikrowellenphotonen. Eine Messung wird in Abb. 16 gezeigt. Für $\bar{n} = 3$ ergibt sich eine Dekohärenzzeit von etwa $50 \mu\text{s}$. Im makroskopischen Fall $\bar{n} \gg 1$ wird die Dekohärenzzeit unmessbar klein: die Dekohärenz einer realen Katze erfolgt instantan.

Literatur

[1] Eine systematische Einführung in die Quantenmechanik und eine ausführliche Beschreibung der hier diskutierten Experimente findet man in: P. Schmüser, *Theoretische Physik für Studierende des Lehramts, Band 1 Quantenmechanik*, Springer 2012.

[2] A. Tonomura, *Electron Holography*, Springer 1994.

- [3] R. P. Feynman, R. B. Leighton, M. Sands, *The Feynman Lectures on Physics*, Addison-Wesley 1965. Deutsche Ausgabe: *Vorlesungen über Physik*, Oldenbourg 1991.
- [4] M. Brune et al., *Observing the Progressive Decoherence of the “Meter” in a Quantum Measurement*, Phys. Rev. Lett. **77**, 4887 (1996).
- [5] H. Rauch, *Neutronen-Interferometrie: Schlüssel zur Quantenmechanik*, Physik in unserer Zeit, 29. Jahrg. 1998, Nr. 2.
- [6] O. Nairz, M. Arndt, and A. Zeilinger, *Quantum interference experiments with large molecules*, Amer. Journ. of Physics, **71** (4) 319 (2003).
- [7] G. Möllenstedt und W. Bayh, *Kontinuierliche Phasenverschiebung von Elektronenwellen im kraftfeldfreien Raum durch das magnetische Vektorpotentials eines Solenoids*, Phys. Blätter **18**, 229 (1962).
- [8] P. Schmüser, *Feynman-Graphen und Eichtheorien für Experimentalphysiker*, Springer 1995.
- [9] M. R. Andrews, C. G. Townsend, H.-J. Miesner, D. S. Durfee, D. M. Kurn, W. Ketterle, *Observation of Interference Between Two Bose Condensates*, SCIENCE Vol. 275, 637 (1997).
- [10] Es ist bis heute nicht gelungen, die Allgemeine Relativitätstheorie und die Quantentheorie zu vereinigen.
- [11] Die Dirac-Gleichung und ihre Lösungen werden in aller Ausführlichkeit in [8] besprochen. Eine kurze Diskussion findet man in: P. Schmüser, *Theoretische Physik für Studierende des Lehramts, Band 2 Elektrodynamik und Spezielle Relativitätstheorie*, Springer 2013.
- [12] S. Haroche, J.-M. Raymond, *Cavity Quantum Electrodynamics*, Scientific American, April 1993, p. 26.
- [13] Sehr lesenswerte Betrachtungen zur Entwicklung der Quantenfeldtheorie und der Vorstellungen über die Natur des Vakuums findet man in einem Artikel von Victor F. Weisskopf, einem der Pioniere dieses Gebiets: *The development of field theory in the last 50 years*, Physics Today November 1981.
- [14] Effektive Massen treten auf, wenn man in der Newton’schen Bewegungsgleichung (Masse · Beschleunigung = Summe aller Kräfte) nur eine Kraft berücksichtigt und die anderen Kräfte ignoriert. Ein triviales Beispiel ist ein mit Helium gefüllter Ballon. Der Ballon steigt nach oben, entgegengesetzt zur Schwerkraft. Ignoriert man die Auftriebskraft in der Luft, so muss man dem Ballon eine negative effektive Masse zuordnen, da er scheinbar von der Erde abgestoßen wird.
- [15] S. Haroche, *Entanglement, Decoherence and the Quantum/Classical Boundary*, Physics Today, July 1998.