

I. Einleitung

In diesem Manuskript wird der Versuch unternommen, die Grundideen und Konzepte der modernen Elementarteilchentheorie zu vermitteln und dabei den komplizierten mathematischen Formalismus weitgehend zu vermeiden. Eine Ausnahme bilden die Anhänge A und B, dort werden die Dirac-Gleichung und die Eichtheorie der elektromagnetischen Wechselwirkung mit allen notwendigen Formeln und Rechenschritten hergeleitet. Vorausgesetzt werden Grundkenntnisse in der Quantenmechanik und der Speziellen Relativitätstheorie. Eine systematische Einführung in die Elementarteilchentheorie mit ausführlichen mathematischen Berechnungen findet man in dem Lehrbuch “Feynman-Graphen und Eichtheorien für Experimentalphysiker” [1], das aus meinen zweisemestrigen Vorlesungen im Rahmen des Graduiertenkollegs an der Universität Hamburg hervorgegangen ist.

Leptonen und Quarks

Nach dem heutigen Stand des Wissens ist alle Materie aus Fermionen mit Spin $1/2$ aufgebaut: Leptonen und Quarks. Die geladenen Leptonen (Elektron e^- , Myon μ^- , Tauon τ^-) haben innerhalb der heutigen Messgenauigkeit von 10^{-18} m keine Ausdehnung und auch keine innere Struktur. Jedem geladenen Lepton ist ein eigenes Neutrino zugeordnet, und zu jedem gibt es ein Antiteilchen. Die Hadronen (Teilchen mit starker Wechselwirkung) sind dagegen ausgedehnt mit typischen Radien von 1 fm. Die vergleichsweise großen Radien deuten auf eine innere Struktur hin. Schon Anfang der 1960er Jahre stellte sich heraus, dass viele Eigenschaften der Hadronen und ihre Systematik durch die Annahme von Bausteinen mit drittelzahliger Ladung und Spin $1/2$ erklärt werden können. In den 1970er Jahren konnte dann experimentell gezeigt werden, dass diese als Quarks bezeichneten Bausteine tatsächlich in den Hadronen vorhanden sind. Sechs verschiedene Quarksorten sind inzwischen nachgewiesen worden u (up), d (down), s (strange), c (charm), b (bottom), t (top). Baryonen wie Proton und Neutron bestehen aus drei Quarks (qqq). Dies erklärt ihren halbzahligen Spin. Mesonen sind gebundene Zustände von Quark-Antiquark-Paaren, sie haben ganzzahligen Spin.

Wechselwirkungen und Feldquanten

In der Teilchenphysik sind drei Wechselwirkungen von Bedeutung. Die bekannteste ist die *elektromagnetische Wechselwirkung*, die die Kräfte zwischen geladenen Teilchen beschreibt. Wesentlich stärker ist die *starke Wechselwirkung* zwischen den Nukleonen im Atomkern, und wesentlich schwächer erscheint uns die *schwache Wechselwirkung*, die radioaktive Betazerfälle bewirkt. Alle Wechselwirkungen werden durch Austausch von Feldquanten vermittelt. Die elektromagnetische Wechselwirkung wirkt auf alle geladenen Teilchen, das Feldquant ist das Photon. Die schwache Wechselwirkung wirkt auf alle Quarks und Leptonen, die Feldquanten sind die $W^\pm$ - und Z^0 -Bosonen. Die eigentliche starke Wechselwirkung gibt es nur zwischen den Quarks; sie wird durch acht Gluonen vermittelt. Alle genannten Feldquanten haben den Spin 1 und sind Bosonen. Die um viele Zehnerpotenzen schwächere Gravitation wird meistens ignoriert, aber das könnte sich als Defizit der heutigen Elementarteilchentheorie erweisen. Die kürzlich entdeckten Gravitationswellen sind ein Beweis für die Existenz des Gravitons, das den Spin 2 haben sollte. Photon, Graviton und die acht Gluonen sind masselos, während die W - und Z -Bosonen eine sehr große Masse (80 - 90 GeV/c²) haben.

Eichtheorien und Higgs-Mechanismus

Die Eichinvarianz spielt in den neueren Theorien eine fundamentale Rolle. Im Rahmen der Eichtheorien ist es gelungen, die elektromagnetischen und die schwachen Wechselwirkungen auf eine gemeinsame theoretische Basis zu stellen. Sämtliche Kopplungen der Fermionen (Neutrinos, geladene Leptonen und Quarks) an die Feldquanten der elektro-schwachen Wechselwirkung – die Photonen, die $W^\pm$ - und die Z^0 -Bosonen – ergeben sich in eindeutiger Weise aus der eichtheoretischen Formulierung. Ein grundsätzliches Problem besteht jedoch darin, dass die Eichinvarianz nur für masselose Felder gültig ist. Masselosigkeit der Feldquanten bedeutet unendliche Reichweite der Wechselwirkung. Die sehr kurze Reichweite der schwachen Wechselwirkung und die damit gekoppelte extrem hohe Masse der Feldquanten $W^\pm$ und Z^0 lassen sich zur Zeit nur über den Higgs-Mechanismus mit dem Prinzip der Eichinvarianz in Einklang bringen. Es wird dabei die Annahme gemacht, dass die Feldquanten a priori masselos sind, dass ihnen jedoch durch die Wechselwirkung mit einem Hintergrundfeld eine (effektive) Masse verliehen wird¹.

¹Effektive Massen treten immer dann auf, wenn man in der Newton’schen Bewegungsgleichung *Masse mal Be-*

Die kurze Reichweite der Kernkräfte hat ganz andere Ursachen. In der Quantenchromodynamik (QCD) wird postuliert, dass die als Gluonen bezeichneten Feldquanten masselos sind, aber nur an die Quarks ankoppeln. Die Hadronen werden als neutral hinsichtlich der starken Ladung angesehen. Die experimentell messbaren Kernkräfte entsprechen den kurzreichweitigen van-der-Waals-Kräften zwischen elektrisch neutralen Molekülen. Das Higgs-Problem entfällt hier also. Dafür taucht ein neues Phänomen auf: trotz der im Prinzip unendlich großen Reichweite der Quark-Gluon-Kräfte ist es bisher nie gelungen, freie Quarks zu beobachten. In der QCD erklärt man dies mit der zusätzlichen Annahme des Quark-“Confinement”, für das es plausible Argumente sowie auch Hinweise aus der Gitter-Eichtheorie, aber keinen strikten Beweis gibt.

II. Die Dirac-Gleichung

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 geführt². Die Schrödinger-Gleichung beruht auf der nichtrelativistischen Mechanik und ist auf Teilchenreaktionen bei hohen Energien nicht anwendbar. Es hat schon kurz nach der Vollendung der Quantenmechanik Versuche gegeben, eine relativistische Verallgemeinerung zu finden, man ist dabei aber auf unerwartete begriffliche Schwierigkeiten gestoßen.

Die Klein-Gordon-Gleichung und das Problem der negativen Energien

Die Operatoren für Energie und Impuls haben in der Quantenmechanik die Gestalt

$$\hat{E} = i\hbar \frac{\partial}{\partial t}, \quad \hat{\mathbf{p}} = -i\hbar \nabla. \quad (1)$$

Es wird postuliert, dass die Operatoren diese Form auch in der relativistischen Quantenelektrodynamik und den anderen relativistischen Quantenfeldtheorien (Standard-Modell der vereinheitlichten elektro-schwachen Wechselwirkung, Quantenchromodynamik) beibehalten, wobei der Energieoperator dann allerdings die Summe von kinetischer Energie und Ruheenergie umfasst.

Zunächst soll daran erinnert werden, wie man zur Schrödinger-Gleichung kommt. Ausgangspunkt ist die nichtrelativistische Energie-Impuls-Beziehung, in der wir Energie und Impuls durch die Operatoren (1) ersetzen und auf die Wellenfunktion anwenden. Für ein freies Elektron ergibt das

$$E = \frac{\mathbf{p}^2}{2m_e}, \quad i\hbar \frac{\partial \psi}{\partial t} = -\frac{\hbar^2}{2m_e} \nabla^2 \psi.$$

Diese Vorgehensweise wird nun auf den relativistischen Fall übertragen. Wir beginnen mit der relativistischen Energie-Impuls-Beziehung im kräftefreien Fall

$$E^2 = \mathbf{p}^2 c^2 + m_0^2 c^4$$

und substituieren die Operatoren (1). Dies führt uns zur *Klein-Gordon-Gleichung*

$$-\hbar^2 \frac{\partial^2 \psi}{\partial t^2} = -\hbar^2 c^2 \left(\frac{\partial^2 \psi}{\partial x^2} + \frac{\partial^2 \psi}{\partial y^2} + \frac{\partial^2 \psi}{\partial z^2} \right) + m_0^2 c^4 \psi. \quad (2)$$

Es gibt einen ganz wesentlichen Unterschied zur Schrödinger-Gleichung: die Klein-Gordon-Gleichung ist eine Differentialgleichung von der zweiten Ordnung in der Zeit und ähnelt damit der klassischen Wellengleichung. Als Folge davon existieren zwei Typen von Lösungen:

$$\psi_p(\mathbf{r}, t) = A \exp(i\mathbf{k} \cdot \mathbf{r}) \exp(-i\omega t), \quad \psi_n(\mathbf{r}, t) = A \exp(i\mathbf{k} \cdot \mathbf{r}) \exp(+i\omega t). \quad (3)$$

Für beide Typen von Lösungen finden wir durch Einsetzen in (2)

$$\hbar^2 \omega^2 = c^2 \hbar^2 \mathbf{k}^2 + m_0^2 c^4 = c^2 \mathbf{p}^2 + m_0^2 c^4. \quad (4)$$

Wählen wir $\hbar \omega = +\sqrt{c^2 \mathbf{p}^2 + m_0^2 c^4}$ als positiv definit, so ergibt sich bei Anwendung des Energieoperators

$$\hat{E} \psi_p = i\hbar \frac{\partial \psi_p}{\partial t} = +\hbar \omega \psi_p, \quad \hat{E} \psi_n = i\hbar \frac{\partial \psi_n}{\partial t} = -\hbar \omega \psi_n.$$

schleunigung = Summe aller Kräfte eine oder mehrere Kräfte ignoriert. Ein einfaches Beispiel ist ein mit Helium gefüllter Ballon. Infolge des Auftriebs in der Luft steigt er nach oben. Wenn man aber die Existenz der Lufthülle ignoriert und nur die nach unten weisende Schwerkraft berücksichtigt, ist man gezwungen, dem Ballon eine negative effektive Masse zuzuordnen, da er ja scheinbar von der Erde abgestoßen wird.

²Es ist bis heute nicht gelungen, die Allgemeine Relativitätstheorie und die Quantentheorie zu vereinigen.

Dies ist ein sehr befremdliches Resultat: die Energie E kann sowohl positive wie negative Werte annehmen. Die hier betrachtete Energie enthält aber keinen potentiellen Energieanteil (potentielle Energien sind ja häufig negativ), sie ist vielmehr die Summe der positiven Ruheenergie und der positiven kinetischen Energie und somit eine Größe, die in jedem Fall größer als null sein muss.

Die Lösungen vom Typ ψ_p sind physikalisch akzeptabel, es sind die Wellenfunktionen mit positiver Energie

$$E_p = +\hbar\omega = +\sqrt{c^2\mathbf{p}^2 + m_0^2c^4}.$$

Ein ernsthaftes Problem stellen dagegen die Lösungen vom Typ ψ_n dar, diese Wellenfunktionen haben negative Energie-Eigenwerte

$$E_n = -\hbar\omega = -\sqrt{c^2\mathbf{p}^2 + m_0^2c^4}.$$

Die Klein-Gordon-Gleichung führt also zu Resultaten, die lange Zeit unverständlich blieben. Eine zweite Schwierigkeit ist, dass auch die Wahrscheinlichkeitsdichte negative Werte annehmen kann. Dies waren die Gründe, dass die Klein-Gordon-Gleichung zunächst aufgegeben wurde.

Probleme mit negativen kinetischen Energien gibt es bei der Schrödinger-Gleichung nicht, da sie von der ersten Ordnung in der Zeit ist. Die Definitionsgleichung (1) des Energie-Operators stellt sicher, dass nur ψ_p eine Lösung der Schrödinger-Gleichung sein kann, nicht aber ψ_n . Alle Schrödinger-Wellenfunktionen haben in der Tat eine Zeitabhängigkeit der Form $\exp(-i\omega t)$, und bei freien Teilchen sind die Energien stets ≥ 0 .

Die Dirac-Gleichung

Die negativen Energiewerte sind offensichtlich mit der zweiten Zeitableitung in der Klein-Gordon-Gleichung verknüpft, die zur Folge hat, dass die Funktionen ψ_n Lösungen sind. Um diese Schwierigkeit zu umgehen, suchte *Paul Dirac* nach einer Gleichung, die wie die Schrödinger-Gleichung nur von der ersten Ordnung in der Zeit ist. Die Lorentz-Invarianz erfordert dann, dass auch die räumlichen Ableitungen nur in erster Ordnung vorkommen. Für ein freies Elektron machte Dirac daher den Ansatz

$$i\hbar\frac{\partial\psi}{\partial t} = -i\hbar c\left(\alpha_1\frac{\partial\psi}{\partial x} + \alpha_2\frac{\partial\psi}{\partial y} + \alpha_3\frac{\partial\psi}{\partial z}\right) + m_e c^2 \beta\psi \quad (5)$$

mit noch unbekannten Koeffizienten α_j und β . Es stellt sich heraus, dass die Koeffizienten nicht einfach als komplexe Zahlen gewählt werden dürfen, sondern 4×4 -Matrizen sind, und dass ψ ein vierkomponentiger Wellenfunktionsvektor ist, der Dirac-Spinor genannt wird. Warum das so ist, wird in Anhang A erklärt. Der tiefere Grund liegt in der Erkenntnis, dass jede Lösung ψ der Dirac-Gleichung gleichzeitig auch eine Lösung der Klein-Gordon-Gleichung sein muss, da diese Gleichung eine notwendige Konsequenz der relativistischen Energie-Impuls-Beziehung $E^2 = (p_x^2 + p_y^2 + p_z^2)c^2 + m_0^2c^4$ ist. In Anhang A werden die mathematischen Details besprochen, und es wird die Form der Matrizen α_j und β und der Dirac-Spinoren hergeleitet.

Es stellte sich heraus, dass auch die Dirac-Gleichung Lösungen mit negativer Energie besitzt (s. Anhang A). Diracs ursprüngliche Hoffnung, die unerwünschten negativen Energiewerte mit seiner relativistischen Verallgemeinerung der Schrödinger-Gleichung vermeiden zu können, ging also nicht in Erfüllung. Das war zunächst eine herbe Enttäuschung, aber wirklich große Theoretiker - und Paul Dirac gehörte sicherlich zu den herausragenden theoretischen Physikern des 20. Jahrhunderts - akzeptieren das Unvermeidliche und suchen nach einer neuen Erklärung. Da Dirac sich außerstande sah, die negativen Energiezustände zu eliminieren, machte er die kühne Annahme, dass diese Zustände 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-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. 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 nennt man Positron, es hat die gleiche Masse m_e , aber die entgegengesetzte Ladung $+e$.

Die Elektron-Positron-Paarerzeugung wird schematisch in Abb. 1a gezeigt. Das auf dem positiven Energieniveau befindliche Elektron kann auch wieder in das Loch zurückfallen. In der Sprache

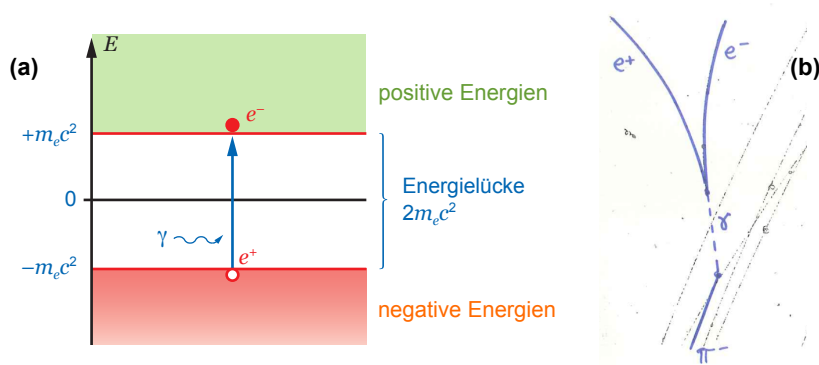


Abbildung 1: **(a)** Dirac-Bild der Antiteilchen. Die positiven Energieniveaus der Elektronen mit $E \geq +m_e c^2$ befinden sich in dem grün schattierten Bereich, die negativen Energieniveaus mit $E \leq -m_e c^2$ sind im rot schattierten Bereich. Nach der Dirac-Interpretation sind die negativen Niveaus 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 geladenes π^- -Meson trifft auf ein ruhendes Proton im flüssigen Wasserstoff und löst die Reaktion $\pi^- + p \rightarrow \pi^0 + n$ aus. Das neutrale π^0 -Meson zerfällt in zwei Gamma-Quanten, von denen eines ein Elektron-Positron-Paar erzeugt. In der Blasenkammer sind nur geladene Teilchen sichtbar.

der Teilchenphysik ist dies die Elektron-Positron-Annihilation. Die Ruheenergie der beiden Teilchen wird dabei auf γ -Quanten übertragen.

Das Positron wurde 1932 von Anderson in der Höhenstrahlung entdeckt. Die Blasenkammeraufnahme in Abb. 1b zeigt die Erzeugung eines Elektron-Positron-Paars. Das Antiproton wurde 23 Jahre später an einem eigens dafür gebauten Beschleuniger (Bevatron in Berkeley, USA) in Proton-Kern-Stößen entdeckt. Die Erzeugung von Antiprotonen durch hochenergetische γ -Strahlung gelang erstmals in einem Experiment bei DESY.

Analogie zwischen Positronen und Löchern im Halbleiter

Das Dirac-Bild der Antiteilchen ist später mit großem Erfolg auf Halbleiter übertragen worden. Hier wollen wir nur sehr kurz auf das Bändermodell des Festkörpers eingehen. In einem periodischen Kristall befinden sich die Elektronen in Energiebändern $E_n(k)$, die den Energieniveaus E_n der Atome entsprechen. Dabei ist n der sog. Bandindex und $k = 2\pi/\lambda$ die Wellenzahl. Das höchste voll besetzte Band nennt man das Valenzband. Ein voll besetztes Band kann keinen elektrischen Strom leiten, denn dazu müssten die Elektronen Energie im elektrischen Feld aufnehmen können. Das Pauli-Prinzip verbietet dies aber, weil sämtliche Energieniveaus im Band besetzt sind. Ein anschauliches Bild des Valenzbandes ist ein langer Stau auf einer Straße mit einer Baustelle. Kein Auto kann fahren, weil vor ihm schon ein anderes Auto steht.

Das nächsthöhere Band über dem Valenzband heißt Leitungsband. In Metallen ist das Leitungsband teilweise mit Elektronen gefüllt. Diese können Energie im elektrischen Feld aufnehmen, weil ihnen viele freie Energieniveaus zu Verfügung stehen, und auf diese Weise einen Strom transportieren (daher der Name “Leitungsband”). In Isolatoren ist das Leitungsband leer, und es kann kein Strom fließen. Das gilt im Prinzip auch für Halbleiter, doch ist bei diesen die Energielücke zwischen Valenz- und Leitungsband relativ klein, sie beträgt 1,1 eV in Silizium. Durch ein Lichtquant kann ein Elektron vom Valenzband in das Leitungsband gehoben werden. Die verbleibende Lücke im Valenzband wird “Loch” (engl. hole) genannt, sie verhält sich wie ein positiver Ladungsträger. In Analogie zur Teilchen-Antiteilchen-Paarerzeugung durch γ -Quanten (Abb. 1) kann man im Halbleiter Elektron-Loch-Paare durch infrarote oder optische Photonen erzeugen. Das in das Leitungsband gehobene Elektron kann einen Strom transportieren; das Loch kann dies ebenfalls.

Wie kann man verstehen, dass sich das Loch im Valenzband wie ein Teilchen mit positiver Ladung verhält? Wir bemühen wieder unsere Analogie des Staus und betrachten eine Autobrücke mit zwei übereinander liegenden Fahrbahnen (Abb. 2). Auf der unteren Fahrbahn herrscht ein totaler Stau, und die Autos stehen alle still (dies ist das volle Valenzband). Die obere Fahrbahn ist völlig leer (das leere Leitungsband). Jetzt heben wir mit einem Kran ein Auto von unten nach oben. Dazu brauchen wir Energie (beim Halbleiter ist dies die Energie des Photons). Das Auto auf der oberen Fahrbahn hat freie Fahrt und fährt nach rechts. Was macht die Lücke in der Autoschlange auf der

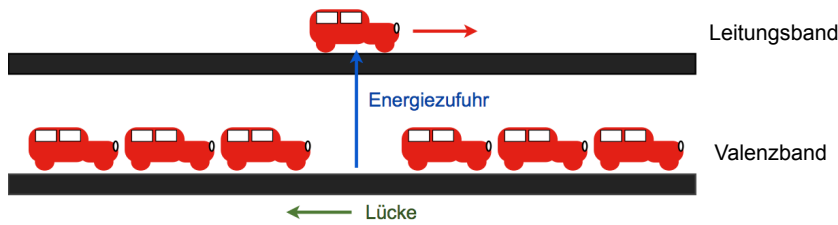


Abbildung 2: Eine Autobrücke mit zwei übereinander liegenden Fahrbahnen als Modell eines Halbleiters.

unteren Fahrbahn? Jeder Autofahrer weiß, was passiert: die Autos schieben sich sukzessive nach rechts vor, und dabei wandert die Lücke immer weiter nach links. Das Gleiche tut ein Loch im Valenzband eines Halbleiters: es wandert zum Minuspol der Batterie, da die Elektronen sich in Richtung Pluspol bewegen. Löcher im Halbleiter sind positive Ladungsträger.

Spin und magnetisches Moment des Elektrons

Die Dirac-Gleichung sagt nicht nur Antiteilchen voraus, sondern macht noch weitere tiefgehende Aussagen über Eigenschaften, die in der nichtrelativistischen Quantenmechanik nicht begründet werden konnten. Dies betrifft den Spin und das magnetische Dipolmoment. Der Begriff "Spin" ist geübten Tischtennispielern vertraut, auch die Erdkugel hat einen Spin aufgrund ihrer Rotation um ihre Achse. Das magnetische Dipolmoment bedeutet, dass sich ein Elektron wie ein kleiner Magnet mit zwei Polen verhält, Nordpol und Südpol. In Anhang A wird gezeigt, dass die Lösungen der Dirac-Gleichung Teilchen mit einem Spin (Eigendrehimpuls) von $\hbar/2$ beschreiben. Ein weiteres, selbst von Dirac ursprünglich nicht erwartetes Resultat seiner Gleichung ist der anomale Wert des magnetischen Moments des Elektrons. Aus der Dirac-Gleichung mit elektromagnetischem Potential folgt ohne jede Zusatzannahme, dass das magnetische Moment den Wert

$$\mu_e = -\frac{e\hbar}{2m_e} \equiv -\mu_B \quad (6)$$

hat, in Übereinstimmung mit dem Experiment (μ_B wird Bohr'sches Magneton genannt). Der Beweis ist mathematisch recht aufwändig und kann hier nicht gebracht werden. 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. Ein mechanisches Modell des Elektrons (Abb. 3) ergibt also falsche Resultate.

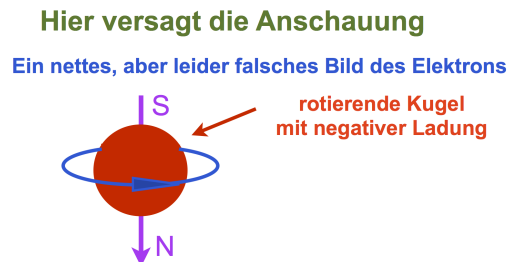


Abbildung 3: Eine negativ geladene Kugel, die um die eigene Achse rotiert, hat einen Spin und ein magnetisches Dipolmoment, d.h. sie verhält sich wie ein Kreisel und wie ein Stabmagnet. Bei dem gezeigten Drehsinn ist der Südpol oben, der Nordpol unten. Dies "klassische" Modell des Elektrons war lange Zeit populär, aber es hat zwei Probleme: die Geschwindigkeit am Äquator müsste größer als die Lichtgeschwindigkeit sein, und das Magnetfeld wäre nur halb so groß wie der Messwert.

Ein brauchbares anschauliches Modell des Elektrons gibt es meines Wissens nicht. Dies ist ein Beispiel, wie die Quantentheorie oft unser Vorstellungsvermögen überfordert. Ein anschauliches Bild des "Dirac-Elektrons" ist mir nicht bekannt.

Das Positron hat als Antiteilchen des Elektrons eine positive Ladung und ein positives magnetisches Moment der Größe μ_B . Die elementaren Bausteine der Materie, Elektronen und Quarks, sind sämtlich Spin-1/2-Fermionen und gehorchen der Dirac-Gleichung. Protonen und Neutronen haben zwar auch den Spin $\hbar/2$, aber ihre Wellenfunktionen sind keine Lösungen der Dirac-Gleichung.

Dies folgt schon aus den ungewöhnlichen Werten der magnetischen Momente:

$$\mu_p = 2,79 \mu_K, \quad \mu_n = -1,91 \mu_K \quad (\mu_K = \frac{e\hbar}{2m_p} \text{ Kernmagneton}). \quad (7)$$

Diese experimentell bestimmten magnetischen Momente waren ein erster Hinweis darauf, dass Proton und Neutron eine innere Struktur besitzen. Heute wissen wir, dass sie aus drei Quarks aufgebaut sind und eine Ausdehnung von rund $1 \text{ fm} = 10^{-15} \text{ m}$ haben, im Gegensatz zum Elektron oder den Quarks, die elementar sind und einen Radius von weniger als $0,001 \text{ fm}$ haben, sofern sie überhaupt eine Ausdehnung besitzen.

III. Die Quantenelektrodynamik

Dirac-Gleichung mit elektromagnetischem Feld

Die Dirac-Gleichung eines freien Elektrons lautet

$$i\hbar \frac{\partial \psi}{\partial t} = -i\hbar c \left(\alpha_1 \frac{\partial \psi}{\partial x} + \alpha_2 \frac{\partial \psi}{\partial y} + \alpha_3 \frac{\partial \psi}{\partial z} \right) + m_e c^2 \beta \psi. \quad (8)$$

Um die Wechselwirkung relativistischer Elektronen mit hochenergetischen Photonen (γ -Quanten) zu beschreiben, benötigt man die Dirac-Gleichung mit einem elektromagnetischen Feldterm. Aus dem Aharonov-Bohm-Effekt ergibt sich, wie dieser Zusatzterm auszusehen hat. Die de-Broglie-Wellenlänge eines Elektrons hat bei Anwesenheit eines Vektorpotentials³ A die Gestalt

$$\lambda = \frac{2\pi\hbar}{m_e v + q A} \quad \text{mit } q = -e.$$

Übersetzt in der Operatorformalismus der Quantentheorie folgt daraus, dass der Impulsoperator

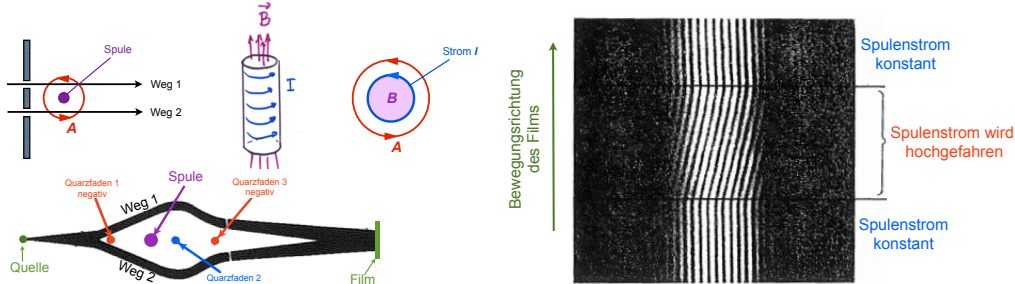


Abbildung 4: Experimentelle Verifikation des Aharonov-Bohm-Effekts. Links oben: Schema des Experiments zur Beobachtung von Elektronen-Interferenzen am Doppelspalt mit einer Solenoid-Spule zur Erzeugung eines magnetischen Vektorpotentials. Links unten: Skizze des Experiments von Möllenstedt und Bayh. Rechts: Die Interferenzstreifen bei konstantem und variablem Strom in der Spule.

eines Teilchens der Ladung q , das sich in einem Vektorpotential A befindet, die folgende Gestalt hat

$$\hat{p} = -i\hbar \nabla - q A. \quad (9)$$

In der feldfreien Dirac-Gleichung muss man also die folgende Substitutionen machen

$$-i\hbar \frac{\partial}{\partial x} \rightarrow -i\hbar \frac{\partial}{\partial x} - q A_x, \quad -i\hbar \frac{\partial}{\partial y} \rightarrow -i\hbar \frac{\partial}{\partial y} - q A_y, \quad -i\hbar \frac{\partial}{\partial z} \rightarrow -i\hbar \frac{\partial}{\partial z} - q A_z.$$

Wenn man noch den potentiellen Energie-Term $q\Phi\psi$ hinzufügt, ergibt sich die Dirac-Gleichung eines Elektrons im elektromagnetischen Feld⁴.

$$i\hbar \frac{\partial \psi}{\partial t} = -i\hbar c \left(\alpha_1 \frac{\partial \psi}{\partial x} + \alpha_2 \frac{\partial \psi}{\partial y} + \alpha_3 \frac{\partial \psi}{\partial z} \right) + m_e c^2 \beta \psi - q c (\alpha_1 A_x \psi + \alpha_2 A_y \psi + \alpha_3 A_z \psi) + q \Phi \psi. \quad (10)$$

³Die Definition des Vektorpotentials und der Eichtransformationen der Elektrodynamik findet man in dem Lehrbuch "Theoretische Physik 2 für Studierende des Lehramts: Elektrodynamik und Spezielle Relativitätstheorie" [3]. Dort wird auch das Experiment zum Aharonov-Bohm-Effekt ausführlich studiert, und es wird plausibel gemacht, wie man die Dirac-Gleichung herleiten kann.

⁴Wenn man die Feinstruktur des H-Atoms berechnen will, muss man $q = -e$ einsetzen und den potentiellen Energie-Term in der Form $q\Phi\psi = -e^2/(4\pi\epsilon_0 r) \psi$ schreiben.

Bei der Wechselwirkung relativistischer Elektronen mit hochenergetischen Photonen ist der potentielle Energie-Term $q\Phi\psi$ vernachlässigbar. Der vorletzte Term in dieser Differentialgleichung enthält das Produkt des elektromagnetischen Vektorpotentials und der Wellenfunktion des Elektrons, er beschreibt also die Elektron-Photon-Wechselwirkung. Die Dirac-Gleichung mit Wechselwirkungsterm kann nur näherungsweise gelöst werden. Die Feynman-Graphen sind eine bildliche und sehr anschauliche Darstellung der Näherungslösung.

Als einfaches mathematisches Beispiel für eine Näherungslösung betrachten wir die Taylorreihe der Funktion $f(x) = \sqrt{1+x}$ für kleine Werte von x :

$$\sqrt{1+x} = 1 + \frac{x}{2} - \frac{x^2}{8} + \frac{x^3}{16} \dots$$

Wenn x sehr klein ist, genügt die erste Ordnung, $\sqrt{1+x} \approx 1 + x/2$. Will man genauer werden, nimmt man die Terme höherer Ordnung hinzu ($-x^2/8$, $x^3/16$..).

Die weiter unten gezeigten Feynman-Graphen repräsentieren die erste Ordnung der Näherungslösung der Dirac-Gleichung mit Feld. Sie sind mit moderatem Aufwand berechenbar und liefern schon recht genaue Werte für den Wirkungsquerschnitt oder die Wahrscheinlichkeit einer QED-Reaktion. Um hochpräzise Vorhersagen zu machen, muss man die Terme höherer Ordnung hinzunehmen. Die zugehörigen Graphen enthalten mehr Vertizes und zusätzliche innere Linien. Ihre Berechnung kostet sehr viel Mühe.

Elementare Prozesse der QED

Alle QED-Reaktionen können aus vier elementaren Prozessen aufgebaut werden, die in Abb. 5 skizziert sind:

- (1) die Emission eines Photons durch ein geladenes Fermion,
- (2) die Absorption eines Photons durch ein geladenes Fermion,
- (3) die Erzeugung eines Fermion-Antifermion-Paars durch ein Photon,
- (4) die Annihilation eines Fermion-Antifermion-Paars in ein Photon.

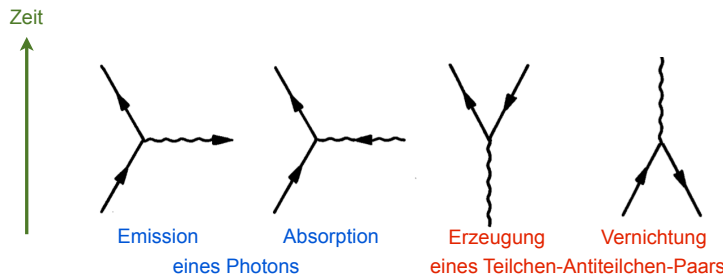


Abbildung 5: Die vier elementaren Prozesse der Quantenelektrodynamik.

Die beteiligten Teilchen oder Photonen sind *reell*, wenn sie ein- oder auslaufen, und sie sind *virtuell*, wenn sie nur im Zwischenzustand als innere Linien auftreten. Ein reelles Teilchen erfüllt die relativistische Energie-Impuls-Beziehung $E^2 = \mathbf{p}^2 c^2 + m_0^2 c^4$, ein virtuelles Teilchen verletzt diese Beziehung und kann daher nicht als freies Teilchen existieren.

In den Elementarprozessen sind Energie und Impuls immer erhalten, aber mindestens eines der Teilchen oder Quanten ist virtuell. Beispielsweise ist die Elektron-Positron-Annihilation in ein reelles Photon kinematisch unmöglich, denn im Ruhesystem des Paares müsste das Photon die Energie $E_\gamma = 2E$, aber den Impuls $p_\gamma = 0$ haben. Dies Photon muss also virtuell sein, sein Massenquadrat ist $m_\gamma^2 = 4E^2/c^2 > 0$.

Reale Prozesse

Alle realen Prozesse lassen sich durch Kombination der Elementarprozesse aufbauen. In der elastischen Elektron-Proton-Streuung werden die beiden ersten Elementarprozesse kombiniert, das ausgetauschte Photon ist virtuell. Bei der Reaktion $e^- + e^+ \rightarrow \mu^- + \mu^+$ werden die Elementarprozesse (4) und (3) kombiniert, das Photon im Zwischenzustand ist virtuell. Zur Beschreibung der Comptonstreuung kombiniert man die Elementarprozesse (2) und (1), aber in anderer Weise als bei der elastischen Streuung; im Compton-Graphen tritt ein virtuelles Elektron auf.

Am Beispiel der Elektron-Proton-Streuung möchte ich kurz die beiden wichtigsten Feynman-Regeln zur Auswertung der Graphen nennen. Als Vertexfaktor tritt die Ladung des Teilchens auf: $-e$ beim Elektron, $+e$ beim Proton. Bei einem d-Quark wäre der Faktor $-e/3$. Das virtuelle Photon

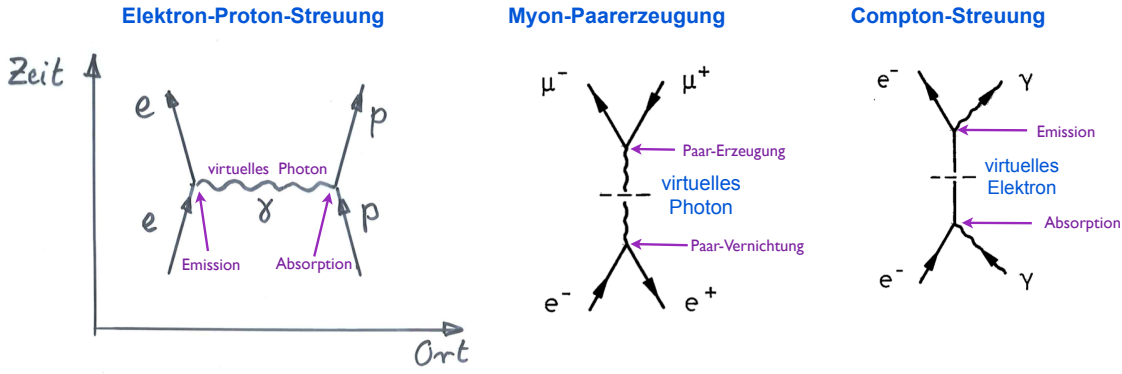


Abbildung 6: Kombination der elementaren Prozesse zu realen Prozessen. Gezeigt werden die Feynman-Graphen für: Elektron-Proton-Streuung, Elektron-Positron-Annihilation in ein virtuelles Photon und die nachfolgende Erzeugung eines $\mu^- \mu^+$ -Paars, Compton-Streuung.

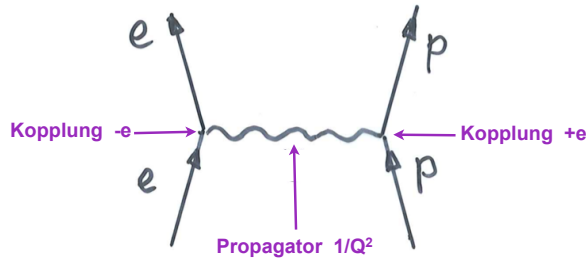


Abbildung 7: Feynman-Regeln zur Auswertung des Streu-Graphen.

“propagiert” die Wirkung vom linken zum rechten Vertex. Mathematisch wird das Photon durch den Propagator P_γ beschrieben, der in vereinfachter Form (Vernachlässigung des Photon-Spins) lautet:

$$P_\gamma(Q^2) = \frac{1}{Q^2} \quad \text{mit} \quad Q^2 = |E_\gamma^2 - p_\gamma^2 c^2|. \quad (11)$$

Die Übergangsamplitude ist proportional zum Produkt der Vertexfaktoren und des Propagators. Die Übergangswahrscheinlichkeit (bzw. den Wirkungsquerschnitt σ) erhält man durch Quadrieren der Übergangsamplitude:

$$\sigma(e + p \rightarrow e + p) \sim \frac{e^4}{Q^4}.$$

IV. Die vereinigte elektromagnetische und schwache Wechselwirkung

Elementare Prozesse und Reaktionen der schwachen Wechselwirkung

Ich zeige in Abb. 8 eine Auswahl der Elementarprozesse mit Leptonen und den Feynman-Graphen für den Myonzerfall, mit dem die Myon-Lebensdauer genau berechnet werden kann. Die schwache WW kann ein Neutrino in ein geladenes Lepton umwandeln, wobei ein W -Boson emittiert wird. Bei der Emission eines Z -Bosons bleibt die Teilchensorte erhalten.

Die schwache WW kann auch die Quark-Sorte ändern, einige Elementarprozesse werden in Abb. 9 gezeigt. Der Betazerfall des Neutrons

$$n \rightarrow p + e^- + \bar{\nu}_e$$

wird durch das rechts gezeigte Quarkdiagramm beschrieben. Das Neutron besteht aus zwei d-Quarks und einem u-Quark. Eines der d-Quarks geht unter Emission eines virtuellen W^- -Bosons in ein u-Quark über, und das virtuelle W^- -Boson zerfällt anschließend in ein Elektron und ein Antineutrino. Die Feynman-Regeln der schwachen Wechselwirkung sollen am Beispiel der Reaktion $\nu_\mu + e^- \rightarrow \mu^- + \nu_e$ erklärt werden. Der Vertexfaktor ist g für beide Vertizes, man kann g als “schwache Ladung” interpretieren. (Hinweis: Es ist üblich, den Vertexfaktor in der Form $g/\sqrt{2}$ zu schreiben. Diese kleine Modifikation ist ohne Belang für unsere qualitative Betrachtung). Der Propagator des ausgetauschten W -Bosons hat eine andere Gestalt als der Photonpropagator. Bei

Leptonische Prozesse der schwachen Wechselwirkung

(Auswahl)

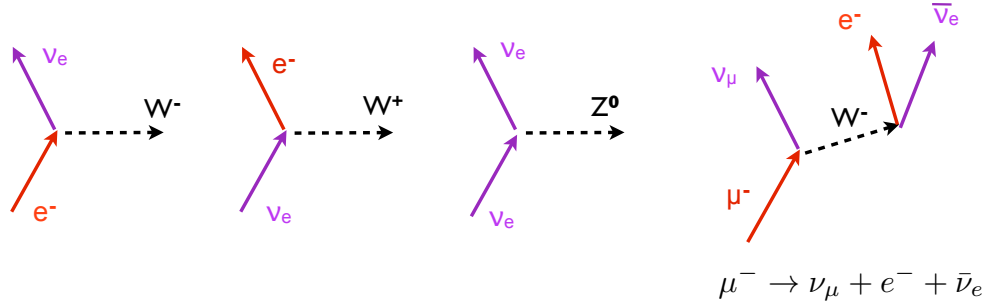


Abbildung 8: Leptonische Elementarprozesse der schwachen Wechselwirkung und der Feynman-Graph für den Myon-Zerfall.

Umwandlung von Quarks durch schwache WW

(Auswahl)

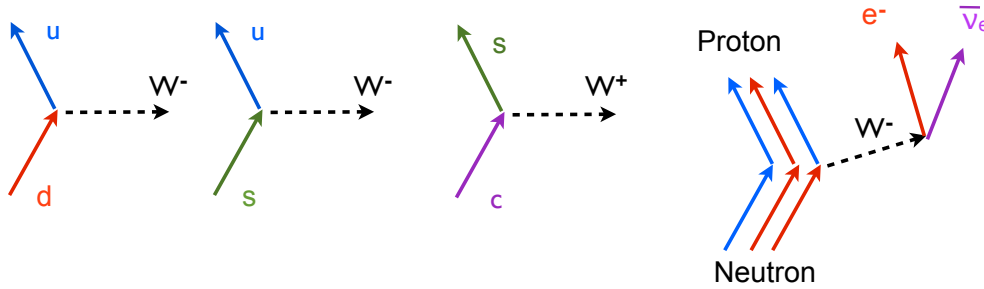


Abbildung 9: Elementarprozesse der schwachen Wechselwirkung mit Quarks und der Feynman-Graph für den Neutron-Zerfall $n \rightarrow p + e^- + \bar{\nu}_e$.

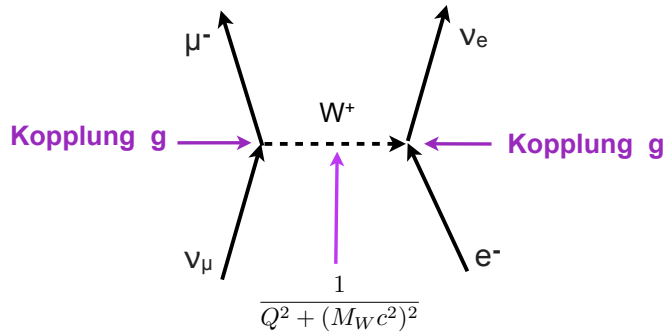


Abbildung 10: Feynman-Regeln zur Auswertung des Graphen der Reaktion $\nu_\mu + e^- \rightarrow \mu^- + \nu_e$.

Vernachlässigung des W -Spins lautet er

$$P_W(Q^2) = \frac{1}{Q^2 + (M_W c^2)^2}. \quad (12)$$

Der Zusatzterm im Nenner, das Quadrat der Ruheenergie des W -Bosons, erweist sich als sehr bedeutsam. Die Übergangsamplitude ist auch hier wieder proportional zum Produkt der Vertexfaktoren und des Propagators, und der Wirkungsquerschnitt ist proportional zu Quadrat der Übergangsamplitude:

$$\sigma(\nu_\mu + e^- \rightarrow \mu^- + \nu_e) \sim \frac{g^4}{(Q^2 + (M_W c^2)^2)^2}.$$

Theorie der vereinigten elektro-schwachen Wechselwirkung

Die erste erfolgreiche Theorie des Betazerfalls wurde von Enrico Fermi in den 1930er Jahren geschaffen und behielt ihre Gültigkeit bis etwa 1970, obwohl es in der Fermi-Theorie ernsthafte

Divergenzprobleme gibt, wenn man zu höheren Ordnungen der Näherungsrechnung geht. In dieser Zeit glaubte man, dass die “schwache Wechselwirkung” intrinsisch um viele Zehnerpotenzen schwächer als die elektromagnetische Wechselwirkung sei. In der Kernphysik ist das tatsächlich der Fall, aber auch in der Teilchenphysik, die mit den Proton- oder Elektron-Beschleunigern bis Ende der 1960er Jahre zugänglich war. Umso erstaunlicher ist die Erkenntnis, dass man beide Wechselwirkungen auf ein gemeinsames theoretisches Fundament stellen kann.

In der von Glashow, Weinberg und Salam entwickelten Theorie⁵ werden der Aharonov-Bohm-Effekt und die Eichtransformationen der QED geeignet verallgemeinert. Daraus ergibt sich die Existenz von drei Feldern W^+, W^-, W^0 mit dem Kopplungsparameter g und einem Feld B^0 mit dem Kopplungsparameter g' . Die W^+ und W^- sind identisch mit den im vorigen Abschnitt diskutierten geladenen Feldquanten der schwachen WW. Die neutralen Felder W^0 und B^0 koppeln beide an die elektrisch neutralen Neutrinos, daher kann keines von beiden identisch mit dem Photonfeld A sein. Man bildet deswegen die Linearkombinationen

$$A = B^0 \cos \theta_W + W^0 \sin \theta_W, \quad Z^0 = -B^0 \sin \theta_W + W^0 \cos \theta_W \quad (13)$$

und wählt den sog. *Weinberg-Winkel* θ_W so, dass das Photonfeld nicht an Neutrinos koppelt.

Die Glashow-Weinberg-Salam-Theorie verknüpft die *a priori* unbekannten Kopplungsparameter g, g' sind mit der Elementarladung

$$e = g \sin \theta_W = g' \cos \theta_W. \quad (14)$$

Der Weinberg-Winkel ist ein unbekannter Parameter der Theorie, der experimentell bestimmt werden muss. Aus den Experimenten am Elektron-Positron-Collider LEP und anderen Beschleunigern ergab sich $\sin^2 \theta_W = 0.23120 \pm 0.00015$. Die kleinen Fehlergrenzen sind ein Beweis für die Präzision der Experimente.

Die Theorie macht also die überraschende Vorhersage, dass die Kopplungen der schwachen Wechselwirkung keineswegs extrem klein sind (wie man bis in die 1960er Jahre vermutete), sondern sogar größer als die elektrische Elementarladung:

$$g = 2.1 e, \quad g' = 1.1 e.$$

Dann erhebt sich natürlich die Frage, warum man überhaupt von der *schwachen* Wechselwirkung spricht und aus welchen Grund diese Wechselwirkung fast immer als so klein erscheint. Die Lösung des Problems liegt beim Propagator. In den allermeisten Fällen ist das Quadrat des Energie-Impuls-Übertrags Q^2 viel kleiner als das Quadrat der Ruheenergie der W - oder Z -Bosonen. Für $Q^2 \ll (M_W c^2)^2$ wird der W -Propagator sehr viel kleiner als der Photonpropagator

$$P_W(Q^2) = \frac{1}{Q^2 + (M_W c^2)^2} \approx \frac{1}{(M_W c^2)^2} \ll P_\gamma(Q^2) = \frac{1}{Q^2},$$

und der zu $|P_W(Q^2)|^2$ proportionale “schwache” Wirkungsquerschnitt ist noch um ein Vielfaches kleiner als der elektromagnetische Wirkungsquerschnitt.

Experimentelle Bestätigung der elektro-schwachen Vereinigung

Wenn die Energie des Beschleunigers ausreicht, zu sehr hohen Werten von Q^2 vorzustößen, sollten Reaktionen der schwachen Wechselwirkung ebenso große Wirkungsquerschnitte haben wie elektromagnetische Reaktionen. Die Messungen der tief-inelastischen Elektron-Proton-Streuung am Proton-Elektron-Collider HERA bestätigen diese Vorhersage in eindrucksvoller Weise wie Abb. 11 zeigt. (Anmerkung: “tief-inelastisch” bedeutet, dass das Feldquant eine derart hohe Energie auf das Proton überträgt, dass dieses zertrümmert wird und in einen Jet hochenergetischer Hadronen übergeht). Die über die elektromagnetische Wechselwirkung verlaufende Reaktion

$$e^- + p \rightarrow e^- + \text{Hadronen}$$

hat einen großen Wirkungsquerschnitt bei $Q^2 \rightarrow 0$, der mit wachsendem Q^2 rasch abfällt. Die über die schwache Wechselwirkung verlaufende Reaktion

$$e^- + p \rightarrow \nu_e + \text{Hadronen}$$

hat einen kleinen Wirkungsquerschnitt bei $Q^2 \rightarrow 0$, der bis $Q^2 = (M_W c^2)^2$ ungefähr konstant bleibt. Oberhalb von $Q^2 = (M_W c^2)^2$ nähern sich die Kurven an und fallen gemeinsam mit wachsendem Q^2 ab. Abbildung 11 ist eine perfekte Illustration der Idee der vereinigten elektromagnetischen und schwachen Wechselwirkungen.

⁵Ein sehr wesentlicher Beitrag stammt von G. 't Hooft, der bewies, dass die Theorie “renormierbar” ist und dass höhere Ordnungen der Näherungsrechnung endliche Resultate liefern.

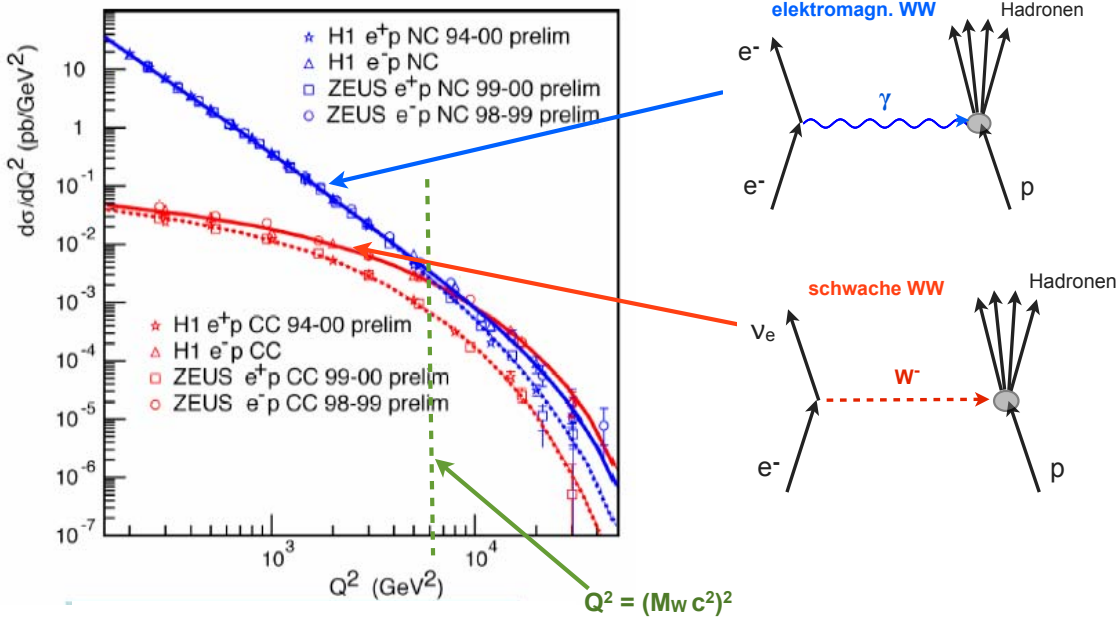


Abbildung 11: Experimenteller Nachweis der elektro-schwachen Vereinigung (Daten von HERA). Gezeigt werden die Wirkungsquerschnitte als Funktion von Q^2 für die Reaktion $e^- + p \rightarrow e^- + \text{Hadronen}$ mit Austausch eines Photons (blaue Kurve und Messpunkte) und für die Reaktion $e^- + p \rightarrow \nu_e + \text{Hadronen}$ mit Austausch eines W^- -Bosons (rote Kurve und Messpunkte). Für $Q^2 \geq (M_W c^2)^2$ laufen die Kurven zusammen. Rechts werden die Graphen der beiden Prozesse gezeigt.

Eigenschaften des Z^0 -Bosons

Die Glashow-Weinberg-Salam-Theorie hat eine erstaunliche Vorhersagekraft und ist hervorragend durch die Präzisionsexperimente am Elektron-Positron-Collider LEP bestätigt worden. Die Erzeugung des Z^0 in der Elektron-Positron-Annihilation und der sofortige Zerfall in Myon-Paare oder Hadronen wird in Abb. 12 gezeigt. Man beobachtet ein ausgeprägtes Resonanzmaximum, das der Resonanz in einem elektrischen Schwingkreis ähnelt. Die Breite der Resonanzkurve eines Schwingkreises hängt von der Dämpfung ab. Ganz analog ist dies bei der Z^0 -Resonanz, die Dämpfung ist hier durch die Zerfälle des Z^0 gegeben. Die Breite Γ und die mittlere Lebensdauer τ sind über die Energie-Zeit-Unschärferelation verknüpft

$$\Gamma \tau = \hbar.$$

Daraus ergibt sich eine extrem kurze Lebensdauer von $\tau = 2.6 \cdot 10^{-25}$ s. Folgende Zerfälle sind präzise vermessen und berechnet worden:

$$Z^0 \rightarrow e^- e^+, \quad \mu^- \mu^+, \quad \tau^- \tau^+, \quad Z^0 \rightarrow q \bar{q} \rightarrow 2 \text{ Hadron} - \text{Jets}$$

Die Zerfälle in Neutrino-Antineutrino sind berechenbar aber nicht messbar (Neutrinos sind *unsichtbar* in einem Speicherring-Experiment). Da sich die partiellen Zerfallswahrscheinlichkeiten zur bekannten totalen Zerfallswahrscheinlichkeit aufaddieren, ist es möglich, die Zahl der Neutrino-Familien zu bestimmen. Eines der wichtigsten LEP-Resultate ist, dass genau drei verschiedene Sorten von Neutrinos existieren: ν_e , ν_μ und ν_τ , siehe auch Abb. 13.

Der Higgs-Mechanismus

Das Standard-Modell der elektro-schwachen Wechselwirkung ist theoretisch sehr gut fundiert und beruht auf einer verallgemeinerten Eichinvarianz. Die Vorhersagen dieser Theorie sind an den Elektron-Positron-Collidern mit hervorragender Präzision bestätigt worden, es gibt jedoch ein tief liegendes 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

$$\Phi(r) \sim \frac{e^{-\mu r}}{r} \quad \text{mit} \quad \mu = \frac{Mc}{\hbar}, \quad (15)$$

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

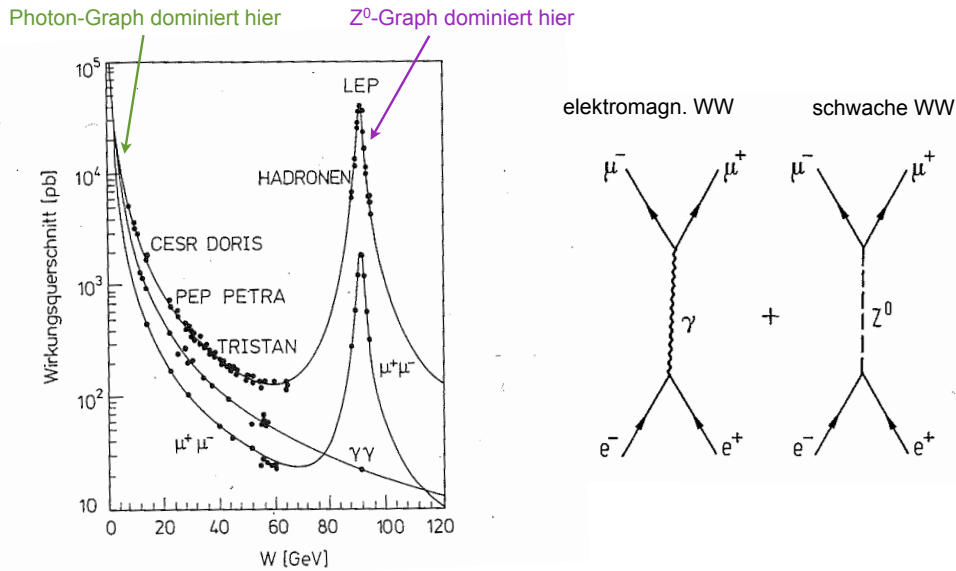


Abbildung 12: Die totalen Wirkungsquerschnitte für Positron-Elektron-Annihilation in Hadronen, Myon-Paare und zwei Gamma-Quanten. Durchgezogene Kurven: Vorhersagen des Standard-Modells.

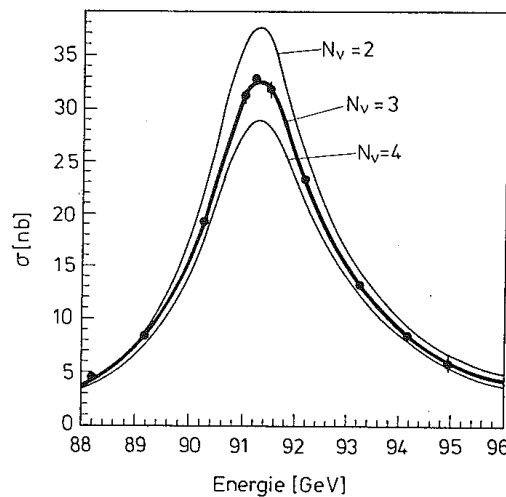


Abbildung 13: Gemessener Wirkungsquerschnitt der Reaktion $e^-e^+ \rightarrow \text{Hadronen}$ bei der Z^0 -Resonanz (Punkte) und die theoretischen Kurven für zwei, drei und vier Neutrino-Sorten.

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 dem 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 $2 \cdot 10^{-18}$ m reduziert wird. Als Folge dieser Abschirmung wird den a priori masselosen Feldquanten Z^0 und $W^\pm$ eine "effektive Masse" verliehen.

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 Innern des Supraleiters verschwindet. In dieser nur 50 nm dicken Schicht fließen widerstandsfreie Supraströme, die das Magnetfeld exponentiell abklingen lassen. Ignoriert man die Existenz des Supraleiters, so muss man den Photonen als Quanten des Feldes eine effektive Masse zuordnen ($m_{eff} \approx 5 \text{ eV}/c^2$).

Auch die Massen der Quarks und Leptonen werden auf den Higgs-Mechanismus zurückgeführt. Für mich ist dies ein sehr unbefriedigender Aspekt des Modells, denn die Masse des jeweiligen Teilchens ist proportional zu seiner Kopplung an das Higgs-Feld. Es gibt demnach genauso viele

verschiedene Kopplungen wie es verschiedene Massen gibt, vom sehr leichten Elektron bis hin zum extrem schweren Top-Quark. Meines Wissens kann niemand die vielen verschiedenen Kopplungen erklären. Da Masse eng mit Gravitation verknüpft ist, könnte man spekulieren, dass das Rätsel der Teilchenmassen sich erst dann auflösen lässt, wenn es gelingen sollte, die Quantentheorie und die Allgemeine Relativitätstheorie zu vereinen.

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.

V. Die Quantenchromodynamik

Für die Physik der Atomkerne und der Hadronen (der stark wechselwirkenden Teilchen) hat es bis in die 1960er Jahre keine fundierte Theorie sondern nur phänomenologische Modelle gegeben. Eine theoretische Basis ist erst mit der Quantenchromodynamik (QCD) geschaffen worden durch Murray Gell-Mann und andere. Die QCD beruht auf dem Quarkmodell: alle Hadronen sind aus kleinsten Objekten aufgebaut, die Spin $1/2$ und drittelzahlige Ladungen haben. Heute sind sechs Quarks und die zugehörigen Antiquarks bekannt. Ihre Massen reichen von $0.34 \text{ GeV}/c^2$ bei den leichten u - und

Tabelle 1: Die sechs bekannten Quarks.

Abkürzung (Name)	Ladung	ungefähre Masse in GeV/c^2
u (up)	$+2/3 e$	0.34
d (down)	$-1/3 e$	0.34
s (strange)	$-1/3 e$	0.51
c (charm)	$+2/3 e$	1.6
b (bottom)	$-1/3 e$	4.8
t (top)	$+2/3 e$	172

d -Quarks, die Proton und Neutron aufbauen, bis $172 \text{ GeV}/c^2$ bei dem extrem massiven t -Quark. Die Hadronen mit halbzahligem Spin werden Baryonen genannt, sie bestehen aus drei Quarks. Die Hadronen mit ganzzahligem Spin heißen Mesonen, sie sind gebundene Quark-Antiquark-Zustände.

Die Ladungen der starken Wechselwirkung

Kurz nach Aufstellung des Quarkmodells wurde eine bei Hadronen unbekannte innere Quantenzahl der Quarks postuliert, die den Namen "Farbe" (color) erhielt. Die Motivation war, dass andernfalls das Pauliprinzip bei Baryonen mit Spin $3/2$ verletzt gewesen wäre. In der Quantenchromodynamik gewinnen die Farben eine tiefe physikalische Bedeutung: sie sind die *Ladungen der starken Wechselwirkung*. Im Unterschied zur QED gibt es drei verschiedene Farbladungen R (rot), G (grün), B (blau) und zu jeder die entsprechende Antiladung.

u – Quark	$\bar{u}$ – Antiquark
elektrische Ladung $+ (2/3)e$	elektrische Ladung $- (2/3)e$
Farbladung R, G oder B	Farbladung $\bar{R}, \bar{G}$ oder $\bar{B}$.

Die starken Kräfte wirken auf die Farbladungen, unterscheiden aber nicht zwischen den verschiedenen Quark-Sorten u, d, s, c, b, t . Die elektromagnetischen und schwachen Wechselwirkungen hingegen unterscheiden zwischen den verschiedenen Quark-Sorten, aber die starke Ladung können sie nicht wahrnehmen, sie sind "farbblind".

Eine sehr bemerkenswerte Aussage der QCD ist, dass die Hadronen "farbneutral" sind, in der Art wie Atome elektrisch neutral sind, da die positive Ladung des Kerns durch die negative Ladung der Hülle kompensiert wird. Mesonen sind gebundene Quark-Antiquark-Systeme, es treten dabei

die Kombinationen Rot-Antirot, Grün-Antigrün und Blau-Antiblau auf. Baryonen sind gebundene Drei-Quark-Systeme in der Farbkombination Rot-Grün-Blau. Wie in der Optik ergibt die additive Mischung der drei Grundfarben "Weiss", in diesem Fall bedeutet das Ladungsneutralität. Die Gluonen koppeln daher nicht direkt an die Hadronen. Das wäre auch unerwünscht, denn sonst gäbe es Kernkräfte mit unendlicher Reichweite.

Die Gluonen tragen die starken Ladungen in der Kombination Farbe-Antifarbe. Es sollte daher neun Gluonen geben. Nur acht davon sind in der QCD realisiert, die in Abb. 14 graphisch aufgetragen sind. Das neunte Gluon muss ausgeschlossen werden, weil es Kernkräfte mit unendlicher Reichweite zwischen den farbneutralen Nukleonen vermitteln würde. Die mathematische Struktur der QCD ist so gewählt worden, dass ein neuntes Gluon nicht vorkommt. Im Unterschied zu

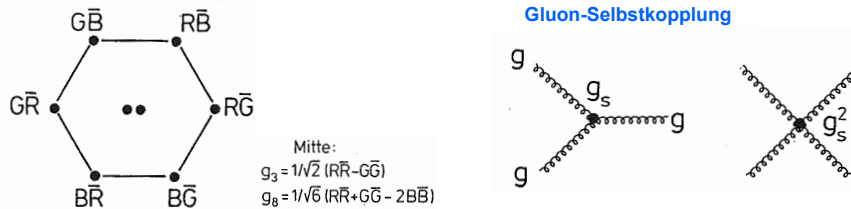


Abbildung 14: Links: Die acht Gluonen der QCD. Rechts: Die Gluon-Selbstkopplung.

den Photonen können auch drei oder vier Gluonen untereinander koppeln, man nennt dies die Gluon-Selbstkopplung (Abb. 14).

In Abb. 15 wird gezeigt, wie zwei verschiedenfarbige Quarks (rot und blau) ihre Farbladungen wechseln können, indem sie ein Rot-Antiblau-Gluon austauschen.

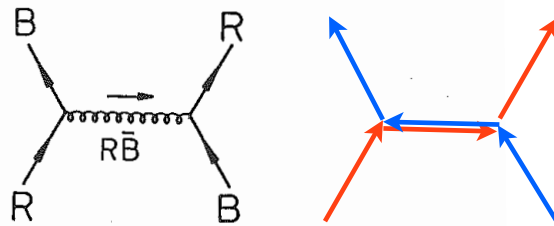


Abbildung 15: Diagramm der Streuung eines roten und eines blauen Quarks mit Austausch der Farbladungen. Rechts wird der Fluss der Farbladungen gezeigt.

Das gezeigte Feynman-Diagramm ist nützlich zur Veranschaulichung der physikalischen Vorgänge. Im Unterschied zu den Feynman-Diagrammen der QED (oder der schwachen WW) sind die Quark-Gluon-Diagramme aber völlig ungeeignet, Wirkungsquerschnitte zu berechnen. Die Farbladungen sind sehr viel stärker als die Elementarladung e , und eine näherungsweise Lösung der Feldgleichungen ist unmöglich. In der sog. Gitter-Eichtheorie versucht man mit riesigem Rechenaufwand, nicht-approximative Lösungen der QCD-Feldgleichungen zu finden.

Confinement, Nichtexistenz freier Quarks

Seit Einführung des Quarkmodells hat man viele Suchen nach drittelzahlig geladenen Teilchen durchgeführt, mit der Idee freie Quarks zu finden. All diese Suchen blieben erfolglos, und so ist man auf das Konzept des "confinement" gekommen: Quarks und Gluonen können nicht als freie Teilchen existieren, sondern sind in ein Volumen eingesperrt, das etwa die Größe eines Protons hat.

Es ist lehrreich, die anziehende Coulomb-Kraft zwischen Proton und Elektron mit der anziehenden starken Kraft zwischen Quark und Antiquark zu vergleichen. Wie in der Elektrostatik kann man sich die Kraft durch Feldlinienbilder veranschaulichen (Abb. 16), je dichter die Feldlinien sind, umso größer wird die Kraft. Mit wachsendem Abstand der Ladungen laufen die elektrischen Feldlinien immer weiter auseinander, die Coulomb-Kraft wird immer kleiner. Die chromo-elektrischen Feldlinien hingegen bilden einen engen Fluss-Schlauch von ca. 1 fm Durchmesser. Sie werden durch magnetische Gluonen zusammengehalten, ähnlich wie parallel fließende elektrische Ströme durch magnetische Kräfte zusammengehalten werden. Da alle Feldlinien vom Quark zum Antiquark gehen, ist die Kraft unabhängig vom Abstand zwischen beiden, und das Potential wächst linear mit dem Abstand r an.

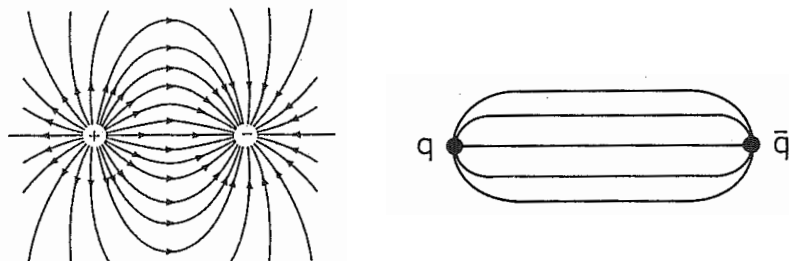


Abbildung 16: Links: Die elektrischen Feldlinien zwischen Proton und Elektron breiten sich über den ganzen Raum aus; die Coulombkraft fällt deswegen mit $1/r^2$ ab, das Potential wie $1/r$. Rechts: Die chromoelektrischen Feldlinien zwischen Quark und Antiquark verlaufen infolge der Gluon-Selbstkopplung in einem engen Schlauch; die Kraft ist unabhängig vom Abstand, das Potential wächst linear mit r an.

Wenn man versucht, Quark und Antiquark in einem Meson immer weiter voneinander zu entfernen mit dem Ziel, sie schliesslich zu trennen, so muss man wegen des linear anwachsenden Potentials soviel Energie aufwenden, dass der Feldlinienschlauch zerreißt und sich an der Bruchstelle ein neues Antiquark-Quark-Paar bildet. Letztendlich hat man das eine Meson durch Energiezufuhr in zwei Mesonen umgewandelt. Einzelne freie Quarks sind nicht dabei herausgekommen. Etwas Ähnliches passiert, wenn man versucht, magnetische Einzelpole zu bekommen, indem man einen Stabmagneten zerreißt. An der Bruchstelle entsteht ein neues Nordpol-Südpol-Paar, und aus dem einen Stabmagneten werden zwei Stabmagnete⁶.

Quarks und Gluonen “sichtbar” gemacht

Quarks kann man nicht als freie Teilchen beobachten, man kann sie aber indirekt “sichtbar” machen. An Elektron-Positron-Collidern hoher Energie kann man die Erzeugung von Teilchen-Antiteilchen-Paaren wie $e^- + e^+ \rightarrow \mu^- + \mu^+$ oder $e^- + e^+ \rightarrow \tau^- + \tau^+$ sehr schön beobachten. Die erzeugten Leptonen/Antileptonen laufen diametral auseinander. Auch die Quark-Antiquark-Paarerzeugung findet statt, mit dem Unterschied, dass diese Objekte sich am Rand des Confinement-Bereichs in enge Bündel von Hadronen umwandeln, die man “Jets” nennt. Ein solches Zwei-Jet-Ereignis wird in Abb. 17 gezeigt.

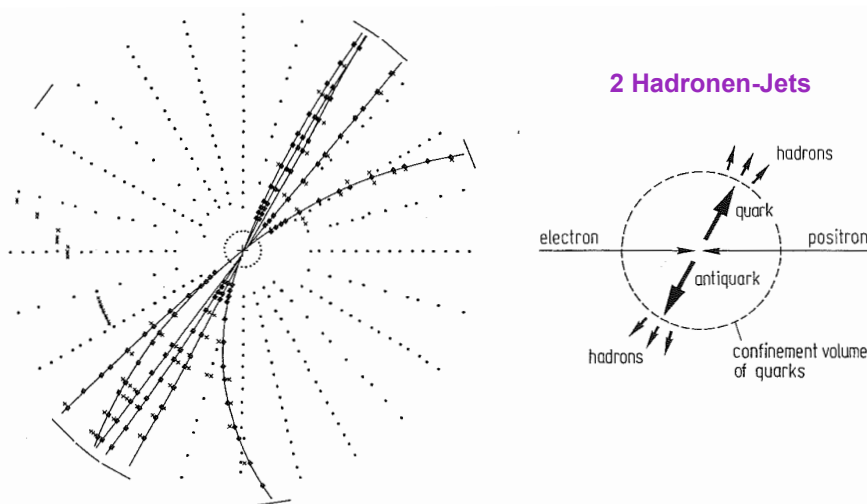


Abbildung 17: Quarks “sichtbar” gemacht: Beobachtung der Reaktion $e^- + e^+ \rightarrow q + \bar{q} \rightarrow 2 \text{ Hadron-Jets}$ am Collider PETRA.

Das wichtigste Ergebnis der PETRA-Experimente ist die Entdeckung der Gluonen. Gelegentlich wird nach der Quark-Antiquark-Paarerzeugung noch ein Gluon emittiert, welches dann einen dritten Hadron-Jet hervorruft. Eines der ersten Ereignisse dieser Art zeige ich in Abb. 18.

⁶Aus theoretischer Sicht gibt es kein Argument, was gegen die Existenz magnetischer Monopole spricht. Man hat sie trotz intensiver Suche nur nie gefunden. Es spricht auch nichts gegen die Existenz einzelner freier Quarks, doch auch die wurden nie gefunden.

Entdeckung der Gluonen bei PETRA

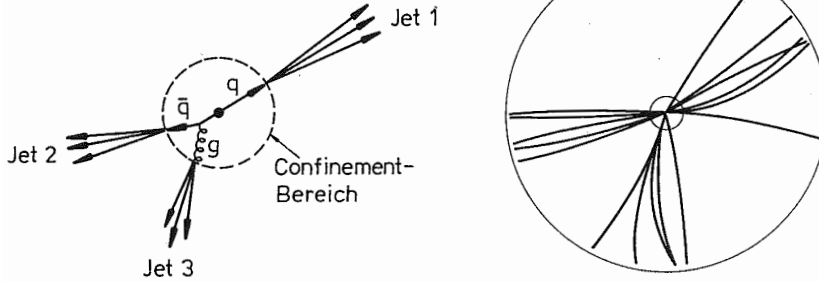


Abbildung 18: Entdeckung der Gluonen am Speicherring PETRA in Hamburg. Gezeigt wird die Reaktion $e^- + e^+ \rightarrow q + \bar{q} + g \rightarrow 3 \text{ Hadron-Jets}$. Links: Schema der Reaktion. Rechts: Spuren der Hadronen in einer zylindrischen Driftkammer.

Anhang A: Herleitung und Eigenschaften der Dirac-Gleichung

Berechnung der Matrizen α_j und β

Eine relativistisch invariante Differentialgleichung, die nur die erste Ableitung nach der Zeit enthält, muss auch von erster Ordnung in den Ortskoordinaten $(x, y, z) \equiv (x_1, x_2, x_3)$ sein. Für die Wellengleichung eines freien Elektrons machte Dirac den Ansatz (5). Dividiert durch $i\hbar$ lautet dieser Ansatz

$$\frac{\partial \psi}{\partial t} = -c \left(\alpha_1 \frac{\partial \psi}{\partial x_1} + \alpha_2 \frac{\partial \psi}{\partial x_2} + \alpha_3 \frac{\partial \psi}{\partial x_3} \right) - \frac{i}{\hbar} m_e c^2 \beta \psi. \quad (16)$$

Wir werden gleich sehen, dass die Wellenfunktion als vierkomponentiger Vektor und die Koeffizienten α_j und β als 4×4 -Matrizen angesetzt werden müssen. Wie kommt man zu diesem Ergebnis? Erstens muss die relativistische Energie-Impuls-Beziehung

$$E^2 = (p_x^2 + p_y^2 + p_z^2)c^2 + m_0^2 c^4$$

erfüllt werden, und zweitens bleibt die Form (1) der Energie- und Impulsoperatoren weiterhin gültig:

$$\hat{E} = i\hbar \frac{\partial}{\partial t}, \quad \hat{p}_j = -i\hbar \frac{\partial}{\partial x_j}.$$

Die Konsequenz ist: jede Lösung ψ der Dirac-Gleichung muss auch eine Lösung der Klein-Gordon-Gleichung sein.

Um auf die Klein-Gordon-Gleichung zu kommen, bilden wir die zweite Zeitableitung. Dazu wird (16) nach t differenziert

$$\frac{\partial^2 \psi}{\partial t^2} = -c \left(\alpha_1 \frac{\partial}{\partial x_1} \left(\frac{\partial \psi}{\partial t} \right) + \alpha_2 \frac{\partial}{\partial x_2} \left(\frac{\partial \psi}{\partial t} \right) + \alpha_3 \frac{\partial}{\partial x_3} \left(\frac{\partial \psi}{\partial t} \right) \right) - \frac{i}{\hbar} m_e c^2 \beta \frac{\partial \psi}{\partial t}.$$

Auf der rechten Seite wird $\partial \psi / \partial t$ aus (16) eingesetzt:

$$\begin{aligned} \frac{\partial^2 \psi}{\partial t^2} &= c^2 \sum_{j=1}^3 \alpha_j^2 \frac{\partial^2 \psi}{\partial x_j^2} - \left(\frac{m_e c^2}{\hbar} \right)^2 \beta^2 \psi \\ &+ c^2 \sum_{j \neq k} \frac{1}{2} (\alpha_j \alpha_k + \alpha_k \alpha_j) \frac{\partial^2 \psi}{\partial x_j \partial x_k} + \frac{i m_e c^3}{\hbar} \sum_{j=1}^3 (\alpha_j \beta + \beta \alpha_j) \frac{\partial \psi}{\partial x_j}. \end{aligned} \quad (17)$$

In der Klein-Gordon-Gleichung gibt es keine gemischten Ableitungen, daher muss die zweite Zeile der Gleichung (17) identisch null sein:

$$c^2 \sum_{j \neq k} \frac{1}{2} (\alpha_j \alpha_k + \alpha_k \alpha_j) \frac{\partial^2 \psi}{\partial x_j \partial x_k} + \frac{i m_e c^3}{\hbar} \sum_{j=1}^3 (\alpha_j \beta + \beta \alpha_j) \frac{\partial \psi}{\partial x_j} \equiv 0.$$

Da dies für beliebige Funktionen ψ erfüllt sein muss, folgt daraus

$$\alpha_j \alpha_k + \alpha_k \alpha_j = 0 \quad \text{für } j \neq k, \quad \alpha_j \beta + \beta \alpha_j = 0. \quad (18)$$

Mit reellen oder komplexen Zahlen kann man diese Relationen nicht erfüllen, da deren Multiplikation dem Kommutativgesetz gehorcht. Es werden daher Matrizen mit komplexen Koeffizienten versucht. Die erste Zeile der Gleichung (17) ist identisch mit der Klein-Gordon-Gleichung, wenn noch folgende weitere Bedingungen gelten:

$$\alpha_j^2 = I, \quad \beta^2 = I,$$

wobei I die Einheitsmatrix ist. Die Matrizen müssen folgende Bedingungen erfüllen:

- (1) Der Hamilton-Operator $\hat{H}$ ist hermitesch, also sind es auch die Matrizen α_j, β .
- (2) $\alpha_j^2 = \beta^2 = I$ (Einheitsmatrix). Daraus folgt, dass die Eigenwerte ± 1 sind.
- (3) Die Matrizen haben alle die Spur 0 (die Spur ist die Summe der Diagonalelemente). Dies folgt aus den Regeln (18):

$$\text{Spur}(\alpha_j) = \text{Spur}(\underbrace{\beta\beta}_I \alpha_j) = \text{Spur}(\underbrace{\beta\alpha_j}_{-\alpha_j\beta} \beta) = -\text{Spur}(\alpha_j).$$

Nun ist die Spur die Summe der Eigenwerte. Da diese die Werte $+1$ und -1 haben, muss die Dimension N der Matrizen gerade sein. Die Dimension $N = 2$ reicht nicht aus, denn es gibt nur drei linear unabhängige hermitesche Matrizen mit Spur 0. Dies sind die Pauli-Matrizen:

$$\sigma_1 = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \sigma_2 = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}, \quad \sigma_3 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}.$$

Die kleinste mögliche Dimension ist $N = 4$. Eine spezielle Darstellung lautet

$$\begin{aligned} \alpha_1 &= \begin{pmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{pmatrix}, & \alpha_2 &= \begin{pmatrix} 0 & 0 & 0 & -i \\ 0 & 0 & i & 0 \\ 0 & -i & 0 & 0 \\ i & 0 & 0 & 0 \end{pmatrix}, \\ \alpha_3 &= \begin{pmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \\ 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \end{pmatrix}, & \beta &= \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}. \end{aligned} \quad (19)$$

Die Wellenfunktion ist ein vierkomponentiger Spaltenvektor, der *Dirac-Spinor* genannt wird.

Lösungen der Dirac-Gleichung für ruhende Teilchen

Hier soll die Dirac-Gleichung für einen sehr einfachen Spezialfall gelöst werden, an dem man aber bereits wesentliche Charakteristika relativistischer Wellengleichungen und ihrer Lösungen erkennen kann, nämlich für ein freies ruhendes Elektron⁷. Da der Impuls und der Wellenvektor des Teilchens null sind, verschwindet die Ortsabhängigkeit der Wellenfunktion. Zur Erinnerung: ein freies Teilchen beschreiben wir durch eine ebene Welle mit der Ortsabhängigkeit $\exp(i\mathbf{k} \cdot \mathbf{r}) = \exp(i(k_x x + k_y y + k_z z))$. Für $\mathbf{k} = \mathbf{p}/\hbar = 0$ sind die Ableitungen $\partial\psi/\partial x = \partial\psi/\partial y = \partial\psi/\partial z = 0$, und die Dirac-Gleichung nimmt die einfache Gestalt an

$$i\hbar \frac{\partial\psi}{\partial t} = m_e c^2 \beta \psi. \quad (20)$$

Gleichung (20) hat vier unabhängige Lösungen. Es gibt zwei Funktionen mit der Zeitabhängigkeit $\exp(-i\omega_0 t)$:

$$\psi_1 = \exp(-i\omega_0 t) \begin{pmatrix} 1 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \quad \psi_2 = \exp(-i\omega_0 t) \begin{pmatrix} 0 \\ 1 \\ 0 \\ 0 \end{pmatrix} \quad \text{mit} \quad \omega_0 = \frac{m_e c^2}{\hbar}, \quad (21)$$

und zwei Funktionen mit der Zeitabhängigkeit $\exp(+i\omega_0 t)$:

$$\psi_3 = \exp(+i\omega_0 t) \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \end{pmatrix}, \quad \psi_4 = \exp(+i\omega_0 t) \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}. \quad (22)$$

⁷Die Dirac-Gleichung ist als Verallgemeinerung der Schrödinger-Gleichung für relativistische Teilchen konstruiert worden, aber sie gilt selbstverständlich auch für nichtrelativistische Teilchen.

Wendet man den Energie-Operator auf ψ_1 oder ψ_2 an, so erhält man positive Energiewerte:

$$\hat{E}\psi_1 = E_1\psi_1, \quad \hat{E}\psi_2 = E_2\psi_2 \quad \text{mit} \quad E_1 = E_2 = +\hbar\omega_0 = +m_e c^2.$$

Die dritte und vierte Funktion ergeben bei Anwendung des Energie-Operators jedoch einen negativen Wert:

$$\hat{E}\psi_3 = E_3\psi_3, \quad \hat{E}\psi_4 = E_4\psi_4 \quad \text{mit} \quad E_3 = E_4 = -\hbar\omega_0 = -m_e c^2.$$

Die Funktionen ψ_3 und ψ_4 sind offensichtlich ungeeignet, reale Teilchen zu beschreiben. Eine sehr wichtige Erkenntnis ist, dass die konjugiert komplexen Funktionen ψ_3^* und ψ_4^* dafür geeignet sind. Sie beschreiben das Positron mit seinen zwei Spineinstellungen⁸. Ihre Zeitabhängigkeit ist $\exp(-i\omega_0 t)$. Wendet man den Energie-Operator auf ψ_3^* und ψ_4^* , so ergeben sich daher positive Energie-Eigenwerte.

$$\hat{E}\psi_3^* = +\hbar\omega_0\psi_3^*, \quad \hat{E}\psi_4^* = +\hbar\omega_0\psi_4^* \quad (23)$$

Man erkennt, dass die von Dirac konstruierte relativistische Verallgemeinerung der Schrödinger-Gleichung nicht nur das Elektron, sondern auch sein Antiteilchen beschreibt. Aus dem Vergleich der Formeln (21) und (23) lernen wir, dass Elektron und Positron exakt die gleiche Ruhemasse m_e haben. Das gilt ganz generell für alle Teilchen/Antiteilchen-Paare. Die elektrischen Ladungen von Elektron und Positron haben exakt den gleichen Betrag, aber verschiedenes Vorzeichen. Das kann mathematisch bewiesen werden, indem man die Dirac-Gleichung mit elektromagnetischem Feld analysiert.

Spin des Elektrons und Positrons

Der Operator der z -Komponente des Spins hat in der Dirac-Theorie die Form

$$\hat{S}_z = \frac{\hbar}{2} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}. \quad (24)$$

Wendet man $\hat{S}_z$ auf die Spinoren ψ_1 und ψ_2 an, so folgt

$$\hat{S}_z\psi_1 = +\frac{\hbar}{2}\psi_1, \quad \hat{S}_z\psi_2 = -\frac{\hbar}{2}\psi_2. \quad (25)$$

Die Spinoren ψ_1 und ψ_2 beschreiben demnach die beiden Spineinstellungen des Elektrons. Ganz entsprechend gilt

$$\hat{S}_z\psi_3^* = +\frac{\hbar}{2}\psi_3^*, \quad \hat{S}_z\psi_4^* = -\frac{\hbar}{2}\psi_4^*.$$

Auch das Positron hat den Spin 1/2.

Anhang B: Das Konzept der Eichtheorien

In diesem Anhang für mathematisch Interessierte möchte ich die Eichtheorie der elektromagnetischen Wechselwirkung Schritt für Schritt erläutern. Die Eichtheorien der vereinheitlichten elektroschwachen und der starken Wechselwirkung erfordern wesentlich mehr mathematischen Aufwand.

Das skalare und das Vektorpotential in der klassischen Elektrodynamik

Um die Darstellung des elektrischen und magnetischen Feldes durch Potentiale zu finden, gehen wir von den Maxwell-Gleichungen aus.

$$\begin{aligned} (1) \quad & \nabla \cdot \mathbf{E} = \rho/\varepsilon_0 \\ (2) \quad & \nabla \cdot \mathbf{B} = 0 \\ (3) \quad & \nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t} \\ (4) \quad & \nabla \times \mathbf{B} = \mu_0 \mathbf{J} + \varepsilon_0 \mu_0 \frac{\partial \mathbf{E}}{\partial t} \end{aligned}$$

Die Divergenz des Magnetfeldes $\mathbf{B}$ ist immer null, das Feld kann daher als Rotation eines geeigneten Vektorfeldes $\mathbf{A}$ geschrieben werden, das man "Vektorpotential" nennt.

$$\mathbf{B} = \nabla \times \mathbf{A}.$$

⁸Bei genauer Betrachtung zeigt sich, dass noch eine Matrixtransformation angewandt werden muss, um die Wellenfunktionen des Positrons zu erhalten [1].

Das Vektorpotential ist nicht eindeutig. Wenn $\chi(x, y, z)$ eine beliebige skalare Funktion ist, so ergibt das neue Vektorpotential

$$\mathbf{A}' = \mathbf{A} + \nabla\chi$$

das gleiche Magnetfeld $\mathbf{B}$, denn die Rotation des Gradientenfeldes $\nabla\chi$ verschwindet. Diese Transformation des Vektorpotentials nennt man eine *Eichtransformation*.

Die Potentialdarstellung des elektrischen Feldes ergibt sich, indem man den Ausdruck $\mathbf{B} = \nabla \times \mathbf{A}$ in die 3. Maxwell-Gleichung einsetzt. Es folgt

$$\nabla \times \mathbf{E} + \frac{\partial}{\partial t}(\nabla \times \mathbf{A}) = 0 \Rightarrow \nabla \times \left(\mathbf{E} + \frac{\partial \mathbf{A}}{\partial t} \right) = 0.$$

Wie man sieht, ist das Vektorfeld $\mathbf{E} + \frac{\partial \mathbf{A}}{\partial t}$ rotationsfrei und kann deshalb als Gradientenfeld geschrieben werden

$$\mathbf{E} + \frac{\partial \mathbf{A}}{\partial t} = -\nabla\Phi.$$

Damit erhalten wir die allgemeinen, auch bei zeitabhängigen Feldern gültigen Darstellungen

$$\mathbf{E} = -\nabla\Phi - \frac{\partial \mathbf{A}}{\partial t}, \quad \mathbf{B} = \nabla \times \mathbf{A}. \quad (26)$$

Eichinvarianz der klassischen Elektrodynamik

Gegeben seien das skalare Potential $\Phi(\mathbf{r}, t)$ und das Vektorpotential $\mathbf{A}(\mathbf{r}, t)$. Wenn $\chi(\mathbf{r}, t)$ eine beliebige skalare Funktion von Raum und Zeit ist, so lautet die verallgemeinerte Eichtransformation

$$\mathbf{A}' = \mathbf{A} + \nabla\chi, \quad \Phi' = \Phi - \frac{\partial \chi}{\partial t}. \quad (27)$$

Die neuen Potentiale $\mathbf{A}'$ und Φ' ergeben die gleichen Felder $\mathbf{E}$ und $\mathbf{B}$ wie die alten Potentiale. Diese Eigenschaft wird als *Eichinvarianz der Elektrodynamik* bezeichnet.

Eichinvarianz in der Quantentheorie

Die Dirac-Gleichung eines Elektrons in einem elektromagnetischen Potential ist durch Formel (10) gegeben. Die Koeffizienten α_1 , α_2 und α_3 werden zu einem Vektor $\boldsymbol{\alpha}$ zusammengefasst, dann kann die Gleichung in kompakter Form geschrieben werden

$$i\hbar \frac{\partial \psi}{\partial t} = -i\hbar c(\boldsymbol{\alpha} \cdot \nabla \psi) + m_e c^2 \beta \psi - q c(\boldsymbol{\alpha} \cdot \mathbf{A} \psi) + q \Phi \psi. \quad (28)$$

Der Term der potentiellen Energie $q \Phi \psi$ ist hier explizit hingeschrieben worden. Für ein Elektron muss man $q = -e$ einsetzen.

Jetzt kommt ein kniffliges Problem: was passiert bei einer Eichtransformation der elektromagnetischen Potentiale? Zwar bleiben elektrische und magnetische Felder bei der Eichtransformation (27) invariant, aber die Potentiale Φ und $\mathbf{A}$ ändern sich, und damit ändern sich auch die Terme der Dirac-Gleichung. Bedeutet das etwa, dass Elektrodynamik und Quantentheorie inkompatibel sind? Dies ist glücklicherweise nicht der Fall, wenn man parallel zur Eichtransformation der Potentiale eine geeignete Phasentransformation der Wellenfunktion vornimmt.

Bei den folgenden Rechnungen ist es sinnvoll, alle Terme der Gleichung (28) auf die linke Seite zu bringen und sie in folgender Form zu gruppieren

$$\left\{ i\hbar \frac{\partial \psi}{\partial t} - q \Phi \psi \right\} + \left\{ i\hbar c(\boldsymbol{\alpha} \cdot \nabla \psi) + q c(\boldsymbol{\alpha} \cdot \mathbf{A} \psi) \right\} - m_e c^2 \beta \psi = 0. \quad (29)$$

Sei ψ eine Lösung dieser Gleichung. Nun wird die Eichtransformation (27) durchgeführt. Die transformierte Dirac-Gleichung mit der noch unbekannten, geeignet transformierten Wellenfunktion ψ' lautet

$$\left\{ i\hbar \frac{\partial \psi'}{\partial t} - q \Phi' \psi' \right\} + \left\{ i\hbar c(\boldsymbol{\alpha} \cdot \nabla \psi') + q c(\boldsymbol{\alpha} \cdot \mathbf{A}' \psi') \right\} - m_e c^2 \beta \psi' = 0. \quad (30)$$

Behauptung:

$$\psi'(\mathbf{r}, t) = \exp\left(\frac{iq}{\hbar} \chi(\mathbf{r}, t)\right) \psi(\mathbf{r}, t)$$

ist eine Lösung der Gl. (30)

Zum Beweis berechnen wir die zeitlichen und räumlichen Ableitungen von ψ'

$$\begin{aligned}\frac{\partial \psi'}{\partial t} &= \exp\left(\frac{iq}{\hbar}\chi\right) \left[\frac{iq}{\hbar} \frac{\partial \chi}{\partial t} \psi + \frac{\partial \psi}{\partial t} \right] \\ \nabla \psi' &= \exp\left(\frac{iq}{\hbar}\chi\right) \left[\frac{iq}{\hbar} (\nabla \chi) \psi + \nabla \psi \right]\end{aligned}$$

und setzen sie in die linke Seite von (30) ein. Die erste geschweifte Klammer ergibt unter Benutzung der Formeln (27)

$$\left\{ i \hbar \frac{\partial \psi'}{\partial t} - q \Phi' \psi' \right\} = \exp\left(\frac{iq}{\hbar}\chi\right) \left\{ i \hbar \frac{\partial \psi}{\partial t} - q \left(\Phi' + \frac{\partial \chi}{\partial t} \right) \right\} = \exp\left(\frac{iq}{\hbar}\chi\right) \left\{ i \hbar \frac{\partial \psi}{\partial t} - q \Phi \right\}$$

Die zweite geschweifte Klammer ergibt

$$\begin{aligned}\{ i \hbar c (\boldsymbol{\alpha} \cdot \nabla \psi') + q c (\boldsymbol{\alpha} \cdot \mathbf{A}' \psi') \} &= \exp\left(\frac{iq}{\hbar}\chi\right) \{ i \hbar c (\boldsymbol{\alpha} \cdot \nabla \psi) - q c (\boldsymbol{\alpha} \cdot \nabla \chi) \psi + q c (\boldsymbol{\alpha} \cdot \mathbf{A}' \psi) \} \\ &= \exp\left(\frac{iq}{\hbar}\chi\right) \{ i \hbar c (\boldsymbol{\alpha} \cdot \nabla \psi) + q c (\boldsymbol{\alpha} \cdot \mathbf{A} \psi) \}\end{aligned}$$

Nach diesen Umrechnungen werden die beiden geschweiften Klammern wieder in die linke Seite von (30) eingesetzt, und der Phasenfaktor $\exp(iq\chi/\hbar)$ wird ausgeklammert:

$$\exp\left(\frac{iq}{\hbar}\chi\right) \left[\left\{ i \hbar \frac{\partial \psi}{\partial t} - q \Phi \right\} + \{ i \hbar c (\boldsymbol{\alpha} \cdot \nabla \psi) + q c (\boldsymbol{\alpha} \cdot \mathbf{A} \psi) \} - m_e c^2 \beta \psi \right] = 0$$

Die rechte Seite wird null, weil ψ eine Lösung der untransformierten Dirac-Gleichung (29) ist. Damit ist der Beweis erbracht, dass die phasentransformierte Wellenfunktion ψ' eine Lösung der eichtransformierten Dirac-Gleichung (30) ist.

Die Eichinvarianz der Quanten-Elektrodynamik ist somit gewährleistet, wenn man die kombinierten Transformationen macht

$$\begin{aligned}\mathbf{A} &\rightarrow \mathbf{A}' = \mathbf{A} + \nabla \chi, \\ \Phi &\rightarrow \Phi' = \Phi - \frac{\partial \chi}{\partial t}, \\ \psi &\rightarrow \psi' = \exp\left(\frac{iq}{\hbar}\chi\right) \psi.\end{aligned}$$

Lokale Phaseninvarianz

Die absolute Phase der Wellenfunktion ist nicht messbar. Wenn man eine *globale Phasentransformation* der Art $\psi' = \exp(i\alpha) \psi$ mit einem konstanten Phasenfaktor vornimmt, erfüllt ψ' die gleiche Wellengleichung wie ψ , und alle Erwartungswerte bleiben invariant.

Ganz anders wird es, wenn man zulässt, dass sich die Phase von Ort zu Ort ändert. Eine *lokale Phasentransformation* ist charakterisiert durch einen orts- und zeitabhängigen Phasenfaktor

$$\psi(\mathbf{r}, t) \rightarrow \psi'(\mathbf{r}, t) = \exp\left(\frac{iq}{\hbar}\chi(\mathbf{r}, t)\right) \psi(\mathbf{r}, t). \quad (31)$$

Sie ändert die Wellenfunktion nachhaltig und führt im allgemeinen zu einer geänderten physikalischen Situation. Bevor wir dies für Lösungen der Dirac-Gleichung genauer untersuchen, soll der Sachverhalt an einem Beispiel illustriert werden, nämlich der Interferenz von Elektronenwellen hinter einem Doppelspalt (vgl. Abb. 4). Eine lokale Phasentransformation bedeutet beispielsweise, dass wir die Phase der ersten Teilwelle willkürlich um π verschieben, die Phase der zweiten Teilwelle aber unverändert lassen. Das ändert das Interferenzmuster, helle Streifen werden dunkel, dunkle Streifen werden hell. Lokale Phasentransformationen ändern demnach messbare Größen (hier das Interferenzmuster), und es sieht ganz so aus, als könne es keine *lokale Phaseninvarianz* (Invarianz gegenüber lokalen Phasentransformationen) geben. Jetzt kommt eine sehr interessante und folgenreiche Beobachtung: eine lokale Phasentransformation der Wellenfunktion kann durch ein geeignetes elektromagnetisches Vektorpotential rückgängig gemacht werden durch Ausnutzen des Aharonov-Bohm-Effekts.

Um einen anschaulichen Eindruck dieser Gedankengänge zu vermitteln, werden in Abb. 19 globale und lokale Phasentransformationen von Wasserwellen skizziert.

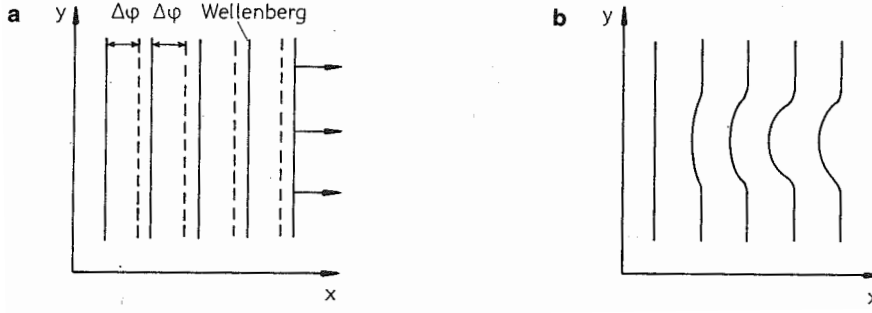


Abbildung 19: **(a)** In einem flachen Behälter wandert eine ebene Wasserwelle in die positive x -Richtung. Die Wellenberge sind als durchgezogene Linien angedeutet. Eine globale Phasentransformation ändert an jedem Ort (x, y) die Phase um den gleichen Betrag $\Delta\varphi$. Dadurch verschieben sich zwar die Wellenberge zu den gestrichelten Positionen, es bleibt aber eine ebene Welle, und im zeitlichen Mittel hat die globale Transformation keinen Effekt. **(b)** Bei einer lokalen Phasentransformation ist die Phasenänderung von Ort zu Ort verschieden: $\Delta\varphi = f(x, y)$. Die transformierte Welle ist keine ebene Welle mehr. Man könnte eine solche Änderung durch ein Hindernis unter der Wasseroberfläche bewirken. Lokale Phasentransformationen erfordern also die Existenz äußerer Kräfte, und aus der Form des Phasenverschiebung kann man Rückschlüsse auf die Natur der Kräfte ziehen.

Wir kehren zur Dirac-Gleichung zurück und nehmen an, dass am Anfang kein elektromagnetisches Feld vorhanden ist. Ein Spin $1/2$ Teilchen der Ladung q genügt dann der Dirac-Gleichung (29) mit $\Phi = 0$ und $\mathbf{A} = 0$.

$$i\hbar \frac{\partial \psi}{\partial t} + i\hbar c (\boldsymbol{\alpha} \cdot \nabla \psi) - m_e c^2 \beta \psi = 0. \quad (32)$$

Jetzt führen wir die lokale Phasentransformation (31) durch. Die phasentransformierte Wellenfunktion ψ' ist sicherlich keine Lösung der feldfreien Dirac-Gleichung (32), aber gemäss dem Postulat der lokalen Phaseninvarianz muss es möglich sein, die Dirac-Gleichung so abzuändern, dass ψ' eine Lösung wird. Aus den im vorigen Abschnitt gemachten Rechnungen wissen wir bereits, wie die abgewandelte Dirac-Gleichung auszusehen hat: nach Gl. (30) sie hat die Gestalt

$$\left\{ i\hbar \frac{\partial \psi'}{\partial t} - q\Phi' \psi' \right\} + \left\{ i\hbar c (\boldsymbol{\alpha} \cdot \nabla \psi') + qc (\boldsymbol{\alpha} \cdot \mathbf{A}' \psi') \right\} - m_e c^2 \beta \psi' = 0, \quad (33)$$

wobei die Potentiale in Einklang mit (27) wie folgt berechnet werden

$$\mathbf{A}' = \nabla \chi, \quad \Phi' = -\frac{\partial \chi}{\partial t} \quad (\Phi = 0, \quad \mathbf{A} = 0).$$

Die wichtige Erkenntnis ist: Die lokal phasentransformierte Wellenfunktion erfüllt nicht mehr die feldfreie Dirac-Gleichung, sondern die Dirac-Gleichung (10) mit den elektromagnetischen Potentialen. Dabei ergeben sich genau die den Feynman-Regeln zugrundeliegenden Kopplungen zwischen den Potentialen und der Wellenfunktion.

Skalares und Vektorpotential lassen sich zu einem Vierervektor $A^\mu = (\Phi/c, \mathbf{A})$ zusammenfassen (siehe [3]). A^μ kann als Wellenfunktion der Photonen interpretiert werden. Das Postulat der lokalen Phaseninvarianz impliziert also die Existenz des Photonen-Feldes.

Schlussbemerkungen

Wenn man es genau nimmt, ist die Eichtheorie zur Beschreibung der elektromagnetischen Wechselwirkung entbehrlich, denn man kennt die klassische Elektrodynamik seit langer Zeit, und mit Hilfe der Aharonov-Bohm-Relation kann man die Form der Dirac- oder Schrödinger-Gleichung mit elektromagnetischen Feldtermen bestimmen. Der Wert der obigen Betrachtungen liegt darin, dass die eichtheoretische Methode damit auf ihre Richtigkeit überprüft werden kann.

Der wahre Wert der Eichtheorien zeigt sich bei den schwachen und starken Wechselwirkungen, weil die W - und Z -Bosonen sowie die Gluonen vorher unbekannt waren und ihre Existenz und ihre Eigenschaften sich erst aus der Eichtheorie erschlossen haben. Auch die Theorie der vereinigten elektro-schwachen Wechselwirkung beruht wesentlich auf der Eichinvarianz.

Die mathematische Gruppentheorie spielt hierbei eine wichtige Rolle. Die obigen Phasentransformationen bilden eine kommutative (abelsche) $U(1)$ -Gruppe (Unitäre Transformationen in einer

Dimension): es gilt $e^a \cdot e^b = e^b \cdot e^a$. In der elektro-schwachen bzw. der starken Wechselwirkung braucht man die nicht-abelschen Gruppen SU(2) und SU(3), das macht den mathematischen Formalismus viel komplizierter. Auch hier wird das Prinzip der lokalen Phaseninvarianz angewandt und impliziert die Existenz der $W^\pm$ - und Z^0 -Feldquanten bzw. der Gluonen.

Literatur und Hinweise

Lehrbücher des Verfassers

- [1] *Feynman-Graphen und Eichtheorien für Experimentalphysiker*, Springer 1995
- [2] *Theoretische Physik 1 für Studierende des Lehramts: Quantenmechanik*, Springer 2012
- [3] *Theoretische Physik 2 für Studierende des Lehramts: Elektrodynamik und Spezielle Relativitätstheorie*, Springer 2013

Wenn es Fragen, Unklarheiten oder Anregungen zu diesem Manuskript gibt oder wenn Sie Fehler entdecken, würde ich mich über eine e-mail oder einen Anruf freuen.

peter.schmueser@desy.de

Tel. 04187 6113