

Unverkäufliche Leseprobe



Scurrile Quantenwelt von Silvia Arroyo Camejo

ISBN: 3540297200
© Springer Verlag 2006

amazon.de
and you're done.™

Gedruckte Verlagsausgabe dieses Titels bestellen

4

Das Doppelspaltexperiment

Was ist das Doppelspaltexperiment?

Die verzwickte Frage, ob das Licht aus *Teilchen* oder *Wellen* besteht, bedeutete seit jeher, sowohl für die antike Philosophie als auch die frühen Naturwissenschaften, ein unlösbares Rätsel. Tatsächlich konnte sie, wie wir im Detail später noch sehen werden, nicht einmal bis zum heutigen Tag explizit entschieden werden. Sie war schon immer Gegenstand heftiger Diskussionen und Auslöser leidenschaftlicher Theorienschlachten. Ein historischer Rückblick auf die zur jeweiligen Zeit vorherrschenden physikalischen Ansichten der Naturwissenschaftler gleicht dabei beinahe einer Art Tennisspiel der Theorien.

Der 1801 vom englischen Universalgenie THOMAS YOUNG (1773–1829) entworfene *Doppelspaltversuch* jedoch sollte eine eindeutige Entscheidung zu Gunsten der *Wellentheorie des Lichts* bringen. Einer amüsanten Anekdote nach, soll YOUNG mehr unbedarft durch Naturbeobachtungen zu der Idee gekommen sein sich mit der Interferenzfähigkeit des Lichts zu befassen. Als er schwimmende Enten in einem Teich beobachtete, bemerkte er die sich ungestört überlagernden, durch die sich bewegenden Entenkörper verursachten Wasserwellen. Durch eben diese Entdeckung inspiriert, konzipierte er schließlich sein *Doppelspaltexperiment mit Licht*.

Vor diesem bedeutenden Experiment war das physikalische Weltbild maßgeblich durch die erfolgreichen Theorien des Engländers ISAAC NEWTON (1643–1727) bestimmt, welcher seinerseits durch die Formulierung der klassischen Mechanik einen, wenn nicht *den* bedeutendsten Grundstein der klassischen Physik legte und somit in den Kanon der Physik einging. Er war es auch, der in seinen Abhandlungen über die Optik die *Korpuskulartheorie* des Lichts geprägt hatte, mit Hilfe derer er die optischen Gesetze der Brechung und Reflexion erklären konnte. Seiner Theorie nach sollte das weiße Licht aus verschiedenfarbigen Teilchen, genannt *Korpuskeln*, bestehen. Ein weißer Lichtstrahl stellte folglich einen Fluss von Korpuskeln dar, der aus Lichtteilchen verschiedenster Farbe besteht.

Zwar existierte zu jener Zeit auch schon eine Wellentheorie des Lichts, die, wie der niederländische Physiker CHRISTIAN HUYGENS (1629–1695) zeigte, die Brechung und Beugung des Lichts zu erklären vermochte, doch ließ die außerordentliche Autorität des erfolgreichen (und jähzornigen) NEWTON anderen Theorien als seiner eigenen keine Chance in der Fachwelt, pflegte er sich doch stets bei Fachkollegen gegen andere Meinungen besserwisserisch durchzusetzen oder, im anderen Extremfall, bei übereinstimmender Meinung zänkisch erboht den Anspruch zu erheben als erster

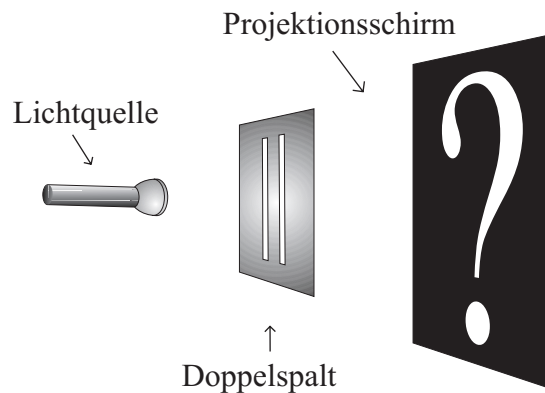


Abb. 4.1. Das Doppelspaltexperiment mit Licht aus der Seitenansicht

die Theorie erdacht zu haben. So kam es, dass NEWTONs Korpuskulartheorie ca. 100 Jahre nahezu unangefochten Bestand genoss.

Jedoch sollte nun YOUNGs Doppelspaltexperiment zur neuen Wellentheorie des Lichts führen. Jenes überaus wichtige Experiment ist, wie in Abb. 4.1 dargestellt, folgendermaßen aufgebaut:

Eine Lichtquelle strahlt möglichst monochromatisches, kohärentes Licht auf einen Doppelspalt. Der Doppelspalt besteht aus einer das Licht abschirmenden Platte, die zwei schmale Spalte besitzt. Hinter diesem Doppelspalt befindet sich eine Projektionswand, die zur Versuchsanalyse den Teil des Lichts auffängt, welcher den Doppelspalt passieren konnte.

Was passiert beim Doppelspaltversuch mit Licht?

Stellen wir uns zunächst vor, bei der Quelle handle es sich erst einmal nicht um etwas derart Unanschauliches wie Licht, sondern um uns wesentlich besser vertraute, handfeste Dinge wie beispielsweise einen möglichst schlechten Fußballspieler, der Fußbälle (idealerweise) willkürlich und ziellos auf einen Doppelspalt schießt, welcher wiederum eine Mauer mit zwei länglichen Durchbrechungen bzw. Schlitzen sei.

Unter diesen Bedingungen kann man recht leicht sagen, wie die *Ankunftswahrscheinlichkeitsverteilung* der Fußbälle aussehen wird: Hinter jedem Loch wird sich ein Haufen mit Fußbällen bilden, egal ob nur das jeweils andere Loch existiert oder nicht. Formeller ausgedrückt gilt für den Doppelspaltversuch mit Fußbällen demgemäß:

$$P_{1+2} = P_1 + P_2, \quad (4.1)$$

d. h. die Ankunfts-wahrscheinlichkeitsverteilung P_{1+2} der Fußbälle bei der Öffnung beider Spalte 1 und 2 ist gleich der Summe der einzelnen Ankunfts-wahrscheinlichkeitsverteilungen P_1 bzw. P_2 bei der Öffnung jeweils nur einer der Spalte 1 oder 2.

Natürlich erwarten wir intuitiv nach einer in der Physik üblichen Verallgemeinerung die Gültigkeit eben dieser Relation (4.1), auch für Tennisbälle, Pingpong-Bälle, Murmeln etc., denn schließlich handelt es sich physikalisch gesehen bei all diesen unterschiedlichen Objektgruppen um prinzipiell wenig unterschiedliche Dinge.

Einzelspalt:

Handelt es sich bei unserer Quelle nun um Licht, und gehen wir von einer Teilchenvorstellung des Lichts aus, so erwarten wir, dass auf der Projektionswand eine ähnliche Ankunfts-wahrscheinlichkeitsverteilung der Photonen wie bei den Fußbällen besteht.

Der Querschnitt des Doppelspaltexperimentes mit Licht wird hier für die zwei folgenden möglichen Fälle eines Einzelspalts schematisch dargestellt. Die Graphik in Abb. 4.2 zeigt dabei das Doppelspaltexperiment mit der ausschließlichen Öffnung von Spalt 1, die in Abb. 4.3 zeigt daraufhin das mit der exklusiven Öffnung von Spalt 2. In beiden Fällen befinden sich links die Strahlungsquelle, in der Mitte der Doppelspalt und rechts der Projektionsschirm mit einer üblichen Darstellungsweise der Ankunfts-wahrscheinlichkeitsverteilung der Photonen.

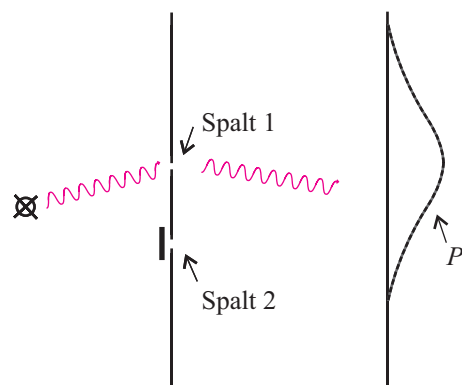


Abb. 4.2. Die Ankunfts-wahrscheinlichkeitsverteilung der Photonen bei der ausschließlichen Öffnung von Spalt 1

Führen wir also das Doppelspaltexperiment zunächst nur mit geöffnetem Spalt 1 durch, wie in Abb. 4.2 zu sehen, müssten wir auf der Projektionswand hinter dem Spalt theoretisch einen hellen Streifen von etwa der Breite des Spalts beobachten.

Die Durchführung jenes Experiments zeigt uns genau diese erwartete *Ankunftswahrscheinlichkeit* P_1 der Photonen, also eine klassische Lichtintensitätsverteilung auf der Photoplatte. Allerdings ist der Streifen auf Grund eines Beugungseffektes, welcher im Übrigen von HUYGENS Wellentheorie des Lichts vorausgesagt wird, augenscheinlich ein wenig nach den Seiten hin verwaschen bzw. verbreitert. Diese *Beugung am Spalt* tritt immer dann auf, wenn die Spaltbreite in der Größenordnung der Wellenlänge der elektromagnetischen Strahlung liegt. Es ergibt sich resultierend für die Ankunftswahrscheinlichkeitsverteilung der Photonen eine ideale Gaußsche Glockenkurve.

Hierbei ist natürlich selbsterklärend, dass die Durchführung des Experiments, wenn nur Spalt 2 geöffnet ist, zu demselben Ergebnis hinter Spalt 2 führt. Bei der alleinigen Öffnung von Spalt 2 folgt ebenfalls eine Intensitätsverteilung des Lichts entsprechend einer Gaußschen Glockenkurve, so wie sie in Abb. 4.3 abgebildet ist.

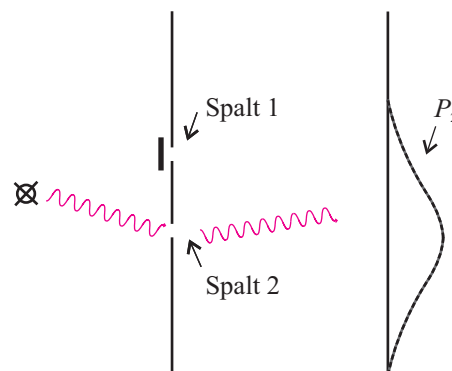


Abb. 4.3. Die Ankunftswahrscheinlichkeitsverteilung der Photonen bei der ausschließlichen Öffnung von Spalt 2

Doppelspalt:

Unseren Überlegungen bezüglich des Doppelspaltexperiments mit Fußbällen und anderen „päckchenhaften“ Objekten zu Folge dürften wir an dieser Stelle vermuten, dass die Ankunftswahrscheinlichkeitsverteilung der Photonen im Doppelspaltexperiment mit Licht bei der Öffnung beider Spalte

1 und 2 gleich der Summe der einzelnen Ankunfts-wahrscheinlichkeitsverteilungen bei der Öffnung jeweils nur einer der Spalte 1 oder 2 ist. Würde dies doch unzweifelhaft mit der Alltagserfahrung aus den Experimenten mit Fußbällen und Murmeln etc. übereinstimmen.

Doch die Durchführung des realen Experiments zeigt uns: *Dem ist nicht so!*

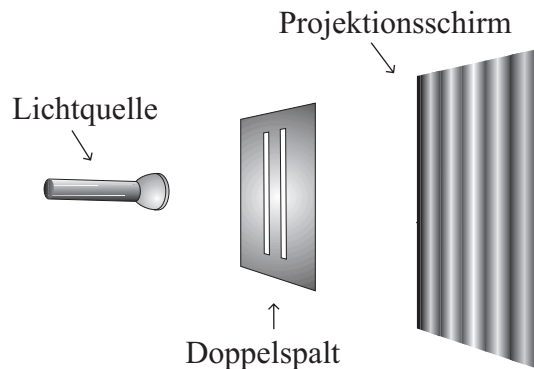


Abb. 4.4. Das beim realen Experiment entstehende Interferenzmuster

Die praktische Durchführung des Doppelspaltexperiments mit Licht zeigt uns eine völlig andere Intensitätsverteilung des Lichts, nämlich eine solche wie sie in Abb. 4.4 zu sehen ist. Die experimentell ermittelte Intensitätsverteilung, welche sich auf dem Projektionsschirm abzeichnet, ist ein auf den ersten Blick unerklärliches Streifenmuster: Es lässt sich eine regelmäßige Abfolge von hellen und dunklen Streifen ausmachen. Es muss demzufolge für den Doppelspaltversuch mit Licht offenkundig

$$P_{1+2} \neq P_1 + P_2 \quad (4.2)$$

gelten. Die Photonen, welche durch Spalt 1 gelangen, und jene, welche Spalt 2 passieren, können nicht einfach auf die Weise addiert werden, wie dies mit Fußbällen etc. der Fall ist. Der tatsächliche Mechanismus ist offensichtlich subtiler.

Wie lässt sich das Streifenmuster erklären?

Genau genommen ist das Phänomen, dass unter Umständen „*Etwas* + *Etwas* = *Nichts*“ sein kann, keineswegs absolut undenkbar. Lässt man z. B.

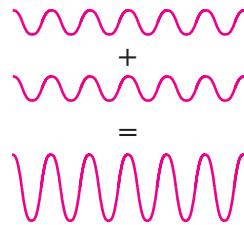
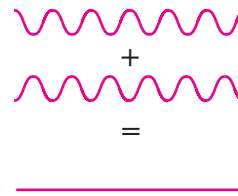
**konstruktive
Interferenz:****destruktive
Interferenz:**

Abb. 4.5. Das Prinzip der Interferenz: Die Addition der Einzelelongationen führt zur resultierenden Elongation der Welle

nebeneinander zwei Steine ins Wasser fallen, bilden sich um jeden Stein konzentrische Oberflächenwellen aus, die sich gegenseitig durchdringen (*interferieren*). An den Stellen, an denen nun Wellental und Wellental bzw. Wellenberg und Wellenberg aufeinander treffen (siehe Abb. 4.5 links), verstärken sich die Auslenkungen von der Nulllinie des Wasserspiegels in ruhendem Wasser. Dies nennt man *konstruktive Interferenz*. Treffen jedoch Wellenberg und Wellental aufeinander (siehe 4.5 rechts), so löschen sich die Auslenkungen des Wassers gegenseitig aus, d. h. es wird an jenen Stellen überhaupt keine Welle registriert. Dies nennt man *destruktive Interferenz*.

Das Phänomen der Interferenz lässt sich für Wasserwellen demnach sehr leicht nachvollziehen. Durch weitere Experimente lässt sich nachweisen, dass jede Art von Wellen diese Eigenschaft der *Interferenzfähigkeit* besitzt. So interferieren auch die longitudinalen Schallwellen, Wellen auf einem Seil, Wellen in einem Federwurm usw.

Teilchenobjekte hingegen wie Fußbälle, Pingpong-Bälle, Murmeln und Ähnliches können nicht interferieren, wie wir oben festgestellt haben. Demnach lässt sich schlussfolgern, dass Interferenzfähigkeit eine Eigenschaft ist, die ausschließlich Wellen zuzuschreiben ist.

Stellen wir uns also das Licht als Welle vor, können wir sein Verhalten am Doppelspalt mit Hilfe der Interferenz erklären, die allgemein bei der Überlagerung von Wellen auftritt:

An den Stellen, an denen Wellental und Wellental bzw. Wellenberg und Wellenberg aufeinander treffen, tritt konstruktive Interferenz auf, so dass sich auf dem Projektionsschirm helle Streifen ergeben. An den Stellen,

wo Wellental auf Wellenberg trifft, tritt destruktive Interferenz auf, welche dunkle Streifen auf dem Projektionsschirm zur Folge haben.

Will man nun die resultierende *Elongation* $y_{1+2}(t)$, also die Auslenkung von der Ruhelage zur Zeit t , der interferierenden Wellen an einem konkreten Ort x bestimmen, so werden, wie oben schon erwähnt, die Einzelelongationen $y_1(t)$ bzw. $y_2(t)$ der Wellen am Ort x unter Berücksichtigung des Vorzeichens addiert (siehe dazu Abb. 4.5). Die resultierende Elongation am Ort x zur Zeit t ist somit

$$y_{1+2}(x; t) = y_1(x; t) + y_2(x; t) . \quad (4.3)$$

Für die resultierende *Amplitude* (= maximale Elongation) am Ort x ergibt sich somit

$$\bar{y}_{1+2}(x) = \bar{y}_1(x) + \bar{y}_2(x) . \quad (4.4)$$

Da die *Intensität* einer Welle am Ort x als das Amplitudenquadrat

$$I_1(x) = \bar{y}_1^2(x) \quad \text{bzw.} \quad I_2(x) = \bar{y}_2^2(x) \quad (4.5)$$

definiert ist, folgt für die Intensität der resultierenden Welle am Ort x aus den Gleichungen (4.4) und (4.5)

$$I_{1+2}(x) = (\bar{y}_1(x) + \bar{y}_2(x))^2 . \quad (4.6)$$

Hieran erkennen wir nun den definitiven Unterschied zwischen der Wahrscheinlichkeitsverteilung von Teilchenobjekten beim Doppelspaltexperiment entsprechend (4.1) und der Intensitätsverteilung von Wellen beim Doppelspaltexperiment nach (4.6), denn für letztere gilt ja, wie wir gerade gesehen haben, dass die resultierende Intensitätsverteilung I_{1+2} der Welle eben nicht der Summe der einzelnen Intensitätsverteilungen I_1 und I_2 entspricht, also

$$I_{1+2}(x) \neq I_1 + I_2 = \bar{y}_1^2(x) + \bar{y}_2^2(x) . \quad (4.7)$$

Die Intensitätsverteilung der elektromagnetischen Strahlung ist quantentheoretisch gesehen jedoch nichts anderes als die Ankunftsverteilung P der Photonen. Da Photonen aber (Licht-) *Teilchen* sind, müsste für sie eigentlich Gleichung (4.1) gelten. Das Experiment zeigt uns nun, dass wir bei der theoretischen Erklärung dieses Doppelspaltexperimentes ein Wellenmodell des Lichts annehmen müssen, um die gefundene Intensitätsverteilung berechnen zu können.

Ist Licht also doch eine Welle?

Eine solche Aussage ist problematisch, denn was ein derart unanschauliches physikalisches Objekt wie die elektromagnetische Strahlung ist, lässt sich nicht ohne weiteres festlegen. Eindeutig lässt sich jedoch sagen, dass elektromagnetische Strahlung etwas ist, das sich nicht mit unserem klassischen, makroskopischen „Welle-oder-Teilchen-Modell“ Verständnis beschreiben lässt.

Es ist nun mal so, dass elektromagnetische Strahlung in bestimmten Experimenten (z. B. im Doppelspaltexperiment) einen Wellencharakter und in gewissen anderen Experimenten (z. B. beim photoelektrischen Effekt) einen Teilchencharakter zeigt. Jedoch sind unsere klassischen Bilder einer Welle oder eines Teilchens hier nur *formale Modelle*, sprich *Arbeitsmodelle*, um etwas zu beschreiben, das sich unserer Alltagserfahrung und Vorstellungskraft derartig entzieht, wie es die Objekte der Mikrowelt nun einmal tun. Das Licht war schon immer etwas Faszinierendes und wird es nicht zuletzt wegen seiner „quantenmechanischen Zwiespaltenheit“ auch immer bleiben.